

Regulação da inteligência artificial no século XXI: desafios jurídicos, éticos e sociotécnicos

Regulation of artificial intelligence in the 21st century: legal, ethical and sociotechnical challenges

Letícia Gabrielli Cordeiro Telles de Moraes

Resumo

A inteligência artificial (IA) tornou-se um dos pilares centrais das transformações tecnológicas, econômicas e sociais contemporâneas. A expansão do uso de algoritmos autônomos e sistemas capazes de aprender, prever e decidir com base em grandes volumes de dados intensificou debates sobre riscos, responsabilidades, direitos fundamentais e limites éticos. Neste contexto, a regulação da IA passou a ser um desafio global, exigindo modelos normativos que conciliem inovação tecnológica, segurança, transparência e proteção de indivíduos e instituições. Este artigo analisa abordagens internacionais de regulação da IA, com foco no AI Act da União Europeia, nos referenciais da OCDE, em iniciativas regulatórias emergentes no Brasil e nos Estados Unidos e nas discussões sobre governança algorítmica. A metodologia inclui revisão bibliográfica, análise documental e estudo comparativo entre arcabouços regulatórios. Os principais resultados indicam que estruturas baseadas em risco, transparência, accountability e supervisão contínua tendem a ser mais eficazes. Conclui-se que a regulação da IA no século XXI requer equilíbrio entre desenvolvimento econômico e a salvaguarda de direitos fundamentais, considerando impactos sociais, éticos e sociotécnicos.

Palavras-chave: inteligência artificial; regulação; governança algorítmica; ética digital; políticas públicas.

Abstract

Artificial intelligence (AI) has become one of the central pillars of contemporary technological, economic and social transformations. The expansion of autonomous algorithms and systems capable of learning, predicting and making decisions based on large volumes of data has intensified debates on risks, responsibilities, fundamental rights and ethical constraints. In this

context, AI regulation has emerged as a global challenge, requiring normative models that reconcile technological innovation, safety, transparency and the protection of individuals and institutions. This article analyzes international approaches to AI regulation, focusing on the European Union's AI Act, OECD guidelines, emerging regulatory initiatives in Brazil and the United States, and discussions on algorithmic governance. The methodology includes bibliographic review, document analysis and comparative study of regulatory frameworks. The main findings indicate that regulatory structures based on risk, transparency, accountability and continuous oversight tend to be more effective. It is concluded that AI regulation in the 21st century requires balancing economic development with the protection of fundamental rights, considering social, ethical and sociotechnical impacts.

Keywords: artificial intelligence; regulation; algorithmic governance; digital ethics; public policy.

1. Introdução

A emergência da inteligência artificial como infraestrutura crítica da vida contemporânea reconfigura de maneira profunda as bases normativas, institucionais e materiais do Estado e do Direito. O que antes era tratado como tecnologia de apoio, restrita a nichos de inovação, hoje se consolida como engrenagem estratégica de decisões que afetam direitos, liberdades e condições materiais de existência. Modelos algorítmicos modulam fluxos de informação, estruturam dinâmicas de consumo, preveem comportamentos, interferem em processos democráticos, classificam indivíduos e organizam o exercício de funções públicas essenciais. Trata-se, portanto, de um fenômeno que desloca a inteligência artificial do campo exclusivamente técnico para o campo jurídico-constitucional, exigindo instrumentos capazes de regular poderes inéditos.

Esse deslocamento aparece de forma evidente nos debates contemporâneos a respeito da governança algorítmica. A produção acadêmica aponta que algoritmos não são ferramentas neutras, mas dispositivos sociotécnicos que carregam decisões humanas e reproduzem, ampliam ou cristalizam desigualdades estruturais. Frank Pasquale descreve esse processo como “the black box society”, indicando que a opacidade estrutural dos sistemas reforça assimetrias de poder entre desenvolvedores, Estado, empresas e cidadãos. Mireille Hildebrandt, por sua vez, sustenta que a IA altera não apenas práticas, mas condições de possibilidade do próprio Direito, ao instaurar formas de decisão automatizada que escapam às categorias jurídicas tradicionais.

No contexto brasileiro, o debate assume contornos ainda mais sensíveis. Relatório do Supremo Tribunal Federal publicado em 2024 aponta que a IA reconfigura dimensões dos direitos fundamentais, afetando autodeterminação informativa, privacidade, igualdade, devido processo legal, participação democrática e até mesmo o exercício da jurisdição. A complexidade do tema se intensifica diante do Projeto de Lei 2.338/2023, que procura estabelecer bases regulatórias para sistemas de IA no país, inserindo-o em movimento internacional mais amplo que inclui o AI Act europeu, diretrizes da OCDE, recomendações da UNESCO e iniciativas regulatórias adotadas em Estados Unidos, China e América Latina.

Neste cenário, surge o principal problema jurídico que orienta este estudo: como construir um marco regulatório capaz de equilibrar a inovação tecnológica com a proteção eficaz de direitos fundamentais em um ambiente algorítmico marcado por opacidade, assimetria informacional e potenciais violações sistêmicas? A resposta exige uma reconstrução teórica e normativa que articule princípios constitucionais, doutrina jurídica, teoria da responsabilidade civil, fundamentos

de governança e critérios técnicos de segurança e transparência.

Ao mesmo tempo, a IA expõe uma tensão clássica entre autonomia privada, livre iniciativa e regulação estatal. Como aponta Danilo Doneda, a proteção de dados e a compreensão dos fluxos informacionais tornam-se pilares necessários do Estado Democrático diante de tecnologias que concentram poder decisório em operadores privados. Bruno Bioni reforça que a proteção da pessoa humana no ambiente digital depende da construção de “espaços regulatórios de responsabilidade”, nos quais agentes tecnológicos não possam agir sem accountability. Ricardo Campos acrescenta que a IA inaugura uma nova etapa do constitucionalismo digital, na qual o Direito precisa lidar com formas inéditas de normatividade produzidas por sistemas técnicos.

Portanto, mais do que um desafio regulatório, a IA instaura um desafio teórico: compreender como o Direito reconstrói suas categorias em um ambiente no qual decisões passam a ser tomadas - ou influenciadas - por sistemas não humanos. A literatura contemporânea reconhece que a IA altera noções clássicas de causalidade, culpa, previsibilidade e imputação, o que exige uma revisão da responsabilidade civil, administrativa e até constitucional.

É nesse contexto que se insere o presente artigo, cujo propósito é construir uma reflexão jurídico-acadêmica aprofundada sobre os caminhos da regulação da inteligência artificial no Brasil, dialogando com experiências internacionais e com as tensões constitucionais emergentes. Parte-se da premissa de que a regulação da IA não deve ser pensada apenas como instrumento de contenção, mas como política pública estruturante destinada a garantir condições materiais de proteção de direitos fundamentais na era digital.

2. O Impacto da Inteligência Artificial no Constitucionalismo Contemporâneo e nos Modelos Jurídicos de Regulação

2.1 A Expansão da Inteligência Artificial e a Reconfiguração dos Direitos Fundamentais

A consolidação da inteligência artificial como infraestrutura decisória - e não apenas como ferramenta instrumental - produz impactos diretos sobre o núcleo de direitos fundamentais reconhecidos no constitucionalismo brasileiro e comparado. Essa constatação aparece de forma contundente no relatório publicado pelo Supremo Tribunal Federal em 2024 sobre os impactos da IA no constitucionalismo contemporâneo, no qual se afirma que a crescente delegação de decisões a sistemas algoritmos transforma a própria arquitetura institucional das garantias. A IA, ao operar por modelos estatísticos, reconfigura pilares como igualdade, privacidade, autonomia, liberdade de expressão, devido processo legal e responsabilidade estatal.

A literatura jurídica contemporânea evidencia que o primeiro impacto relevante da IA ocorre na esfera da autodeterminação informativa, conceito resgatado pela doutrina brasileira a partir das contribuições de Danilo Doneda e da consolidação da Lei Geral de Proteção de Dados. A expansão de modelos de aprendizado de máquina intensifica a coleta massiva de dados, o cruzamento de bases informacionais e a criação de perfis comportamentais, muitas vezes sem que os indivíduos tenham conhecimento ou meios de contestação. Shoshana Zuboff já havia demonstrado que a lógica do “capitalismo de vigilância” transforma a experiência humana em matéria-prima para operações econômicas e políticas, e é precisamente nesse ambiente que a IA opera com maior eficiência.

No Brasil, esse processo produz tensões constitucionais particulares. A Constituição de 1988 atribui centralidade à proteção da dignidade humana, da privacidade e da igualdade. Entretanto, sistemas algorítmicos introduzem práticas que muitas vezes escapam do radar regulatório. O uso crescente de IA por instituições financeiras, empresas de recrutamento, órgãos públicos e plataformas digitais cria cenários nos quais decisões são tomadas com base em probabilidades geradas por modelos estatísticos treinados com dados profundamente marcados por desigualdades sociais. Essas decisões, embora automatizadas, produzem efeitos jurídicos concretos: negar crédito, selecionar candidatos, classificar indivíduos como “risco”, reorganizar fluxos de informação e até limitar a participação democrática.

Esse cenário coloca em evidência o que Ricardo Campos denomina de “normatividade algorítmica”: a ideia de que sistemas de IA criam, de fato, normas técnicas que influenciam comportamentos e decisões com impacto jurídico, mas sem qualquer procedimento democrático de produção normativa. Essa nova forma de poder difuso fundamenta a necessidade de modelos regulatórios mais incisivos, capazes de controlar e responsabilizar agentes privados cujas decisões automatizadas afetam diretamente os direitos fundamentais.

Outro impacto estrutural recai sobre o direito à igualdade, especialmente no contexto de discriminações algorítmicas. Pesquisas empíricas analisadas por diversos autores - e reforçadas em artigos nacionais recentes - demonstram que sistemas de IA tendem a replicar vieses presentes nas bases de dados utilizadas para treinamento. No Brasil, onde desigualdades raciais, de gênero e socioeconômicas são historicamente profundas, o risco de discriminação algorítmica é ainda maior. A literatura jurídica alerta que a adoção de IA em políticas públicas, segurança pública, concessão de benefícios sociais e justiça criminal pode cristalizar desigualdades sistêmicas. O risco é que a tecnologia não apenas reproduza injustiças, mas as legitime sob uma aparência de neutralidade técnica.

A regulação adequada deve, portanto, reconhecer que algoritmos não operam de forma neutra. Eles aproximam-se mais de práticas interpretativas do que de operações matemáticas puras. A doutrina de Bruno Bioni enfatiza que a regulação deve incorporar mecanismos de prevenção, detecção e mitigação de vieses, além de auditorias independentes e transparência ampliada. Esse entendimento é reforçado por autores internacionais, como Frank Pasquale, que sustenta que sem instrumentos de accountability, algoritmos se tornam estruturas opacas de poder, capazes de afetar milhões de pessoas sem que haja mecanismos suficientes de contestação.

Outro ponto crítico amplamente discutido na literatura é o impacto da IA sobre o devido processo legal e o acesso à justiça. Sistemas de tomada de decisão automatizada utilizados por órgãos públicos - como classificação de demandas, triagem de benefícios ou análise preditiva - alteram a forma como o Estado exerce seu poder. O problema surge quando decisões produzidas por sistemas técnicos não oferecem caminhos claros de revisão humana ou quando sequer é possível compreender os critérios utilizados. O STF destaca que decisões automatizadas não podem suprimir as garantias constitucionais que asseguram a todos o direito de conhecer os fundamentos de decisões que os afetam. Assim, a IA desafia diretamente princípios do Estado Democrático de Direito.

A IA também impacta a esfera da liberdade de expressão, já que modelos utilizados por plataformas digitais modulam conteúdos, priorizam ou despriorizam publicações e estruturam o debate público. A literatura aponta que sistemas de recomendação baseados em otimização de

engajamento podem amplificar desinformação, reforçar polarização e criar bolhas informacionais. Lawrence Lessig já havia advertido que “código é lei”, e no ambiente digital isso se aplica com força ampliada: algoritmos privados exercem funções que, na prática, se assemelham à regulação do espaço público discursivo.

Além disso, a responsabilidade civil e administrativa por danos causados por IA tornou-se um dos debates centrais da doutrina. Os textos analisados destacam que o modelo tradicional de responsabilidade - baseado em culpa, conduta humana e nexos causal direto - não se ajusta a sistemas que aprendem, se adaptam e produzem resultados muitas vezes imprevisíveis. A literatura jurídica brasileira e internacional aponta que modelos de responsabilidade objetiva, baseados em risco, são mais adequados, sobretudo para sistemas de alto impacto. O PL 2.338/2023, influenciado pelo AI Act, incorpora essa lógica ao estabelecer responsabilidades proporcionais ao risco do sistema e ao impor deveres de governança, documentação, testes e supervisão humana.

No entanto, a doutrina ressalta que responsabilizar apenas o operador final, ou apenas o desenvolvedor, é insuficiente. A cadeia algorítmica envolve múltiplos agentes: desenvolvedores, fornecedores de dados, empresas que hospedam modelos, órgãos públicos que os utilizam e até entidades que reparam ou fazem fine-tuning. Por isso, autores defendem modelos de responsabilidade distribuída, nos quais cada agente responde pela parte da cadeia que controla. A literatura também enfatiza que o Estado deve exercer função regulatória ativa, fiscalizando práticas e impondo padrões técnicos mínimos para reduzir assimetrias informacionais.

A discussão internacional evidencia dois caminhos regulatórios distintos. A União Europeia adota abordagem preventiva, estruturada e de precaução, impondo obrigações rígidas a sistemas de alto risco. A China, por outro lado, adota modelo centralizador, com forte intervenção estatal voltada à segurança e ao controle social. Os Estados Unidos preferem modelo fragmentado, guiado por diretrizes técnicas e normas setoriais. O Brasil caminha para um modelo híbrido — com princípios constitucionais fortes, alinhamento à proteção de dados e inspiração na regulação baseada em risco. A literatura analisada mostra que esse modelo híbrido pode ser adequado ao contexto brasileiro, desde que dotado de mecanismos efetivos de fiscalização e de participação social.

De forma geral, esta seção demonstra que a IA não é apenas tecnologia; é fenômeno jurídico-político que altera a materialidade dos direitos fundamentais. Assim, sua regulação exige abordagem que combine constitucionalismo, teoria da responsabilidade, governança tecnológica e análise crítica das transformações sociais em curso.

A intensificação do uso da inteligência artificial também provoca uma redefinição da própria compreensão de poder no ambiente constitucional. A literatura sobre teoria crítica da tecnologia mostra que a IA produz uma forma nova de “poder infraestrutural”, conceito desenvolvido por autores como Nick Couldry e Ulises Mejías, que sustenta que sistemas digitais criam redes de influência que ultrapassam controles estatais tradicionais. Quando sistemas algorítmicos passam a organizar atividades econômicas, políticas e sociais, atuam como infraestruturas formadoras de realidade. Isso significa que a IA não apenas opera no mundo: ela organiza o mundo no qual opera.

Sob a ótica constitucional, isso altera a natureza dos direitos fundamentais porque desloca seu ponto de vulnerabilidade. Em vez de focar apenas no Estado ou em indivíduos isolados, os

direitos fundamentais passam a defender o cidadão contra “arquiteturas tecnológicas” inteiras, capazes de moldar comportamento coletivamente. Lawrence Lessig já havia antecipado isso ao afirmar que “o código regula com a mesma força que a lei”, mas Hildebrandt amplia o argumento ao dizer que a IA regula antes mesmo que percebamos.

Esse fenômeno afeta especialmente o direito à privacidade e proteção de dados que deixa de ser apenas a proteção de informações pessoais para se tornar a proteção contra “inferências automatizadas”.

Além disso, a literatura recente demonstra que a IA produz impactos sobre a própria epistemologia do Direito. A lógica estatística dos algoritmos não corresponde à lógica principiológica do constitucionalismo. Enquanto os direitos fundamentais operam por valores, contextos e circunstâncias concretas, os modelos de IA operam por correlações probabilísticas. Essa tensão normativa faz com que decisões automatizadas frequentemente ignorem nuances essenciais, produzindo riscos de injustiça estrutural.

Autores brasileiros já reconhecem esse dilema. Virgílio Afonso da Silva discute que a proporcionalidade exige ponderação contextual e argumentativa – algo que algoritmos não conseguem realizar. Da mesma forma, Ana Frazão argumenta que a IA pode gerar formas de abuso de poder econômico quando usada em mercados digitais altamente concentrados, prejudicando a concorrência e pluralismo econômico, que também são direitos constitucionais.

Do ponto de vista democrático, a IA interfere nos processos de formação da opinião pública. Modelos de recomendação usados por plataformas digitais filtram o conteúdo que cada indivíduo recebe, criando ambientes informacionais personalizados que podem reduzir o debate público comum, aumentar a polarização e favorecer a manipulação política. Isso coloca a IA no centro da discussão sobre a saúde democrática, tema tratado por Cass Sunstein e reforçado por relatórios internacionais sobre IA e democracia.

Por tudo isso, o impacto da IA sobre direitos fundamentais é transversal, estrutural e profundo, exigindo revisão do modo como o Direito compreende suas próprias categorias.

2.2 A Regulação Jurídica da Inteligência Artificial e os Caminhos para um Modelo Brasileiro

A consolidação de um modelo brasileiro de regulação da inteligência artificial exige análise cuidadosa das tradições normativas do país, das peculiaridades de seu constitucionalismo e das práticas institucionais que historicamente moldaram a proteção de direitos fundamentais. Diferentemente de sistemas jurídicos de tradição mais centralizadora ou mais liberal, o Brasil articula um constitucionalismo de caráter híbrido, no qual convivem responsabilidades estatais robustas, protagonismo judicial, diretrizes principiológicas abertas e forte ênfase na dignidade da pessoa humana. Esses elementos são essenciais para compreender como o país pode — e deve — estruturar sua regulação da IA.

O ponto de partida inevitável é a constatação de que a IA não pode ser regulada apenas por instrumentos tradicionais do Direito privado, como contratos ou responsabilidade civil isolada. A literatura demonstra, com crescente consenso, que a IA é infraestrutura crítica para a garantia de direitos e para o funcionamento do Estado. Assim, sua regulação deve ser tratada como política pública. Essa perspectiva aparece com clareza nas discussões do PL 2.338/2023, que busca

construir um marco legal abrangente para a IA no Brasil, articulando responsabilidades, princípios e obrigações técnicas.

O PL brasileiro está alinhado à lógica internacional baseada em risco — modelo que surgiu com o AI Act europeu e rapidamente ganhou adesão de organismos multilaterais. A classificação por risco não é mero exercício formal: trata-se de ferramenta que permite diferenciar sistemas que podem causar danos amplos e profundos (como IA usada em políticas públicas, decisões judiciais, policiamento preditivo, triagem de crédito e saúde) daqueles cujos efeitos são marginais. A literatura jurídica considera esse modelo adequado, pois permite calibrar obrigações regulatórias conforme o impacto potencial da tecnologia.

Contudo, embora o PL 2.338/2023 represente avanço importante, a discussão doutrinária destaca que o Brasil precisa de uma regulação não apenas inspirada em modelos externos, mas coerente com seus próprios desafios estruturais. Diferentemente da União Europeia, o país carrega desigualdades raciais, regionais, econômicas e educacionais profundas. Isso significa que a IA, se não regulada com rigor técnico e sensibilidade constitucional, pode cristalizar formas de exclusão muito mais severas do que aquelas identificadas em países centrais. A literatura brasileira alerta que sistemas de IA aplicados em políticas sociais, saúde pública ou segurança podem afetar grupos vulneráveis com intensidade desproporcional.

A ideia de risco algorítmico, analisada em artigos nacionais recentes, deve ser compreendida como fenômeno estrutural: não se limita ao risco de mau funcionamento técnico, mas inclui risco de discriminação histórica, assimetria informacional e ampliação de desigualdades. Nesse sentido, as obrigações de governança previstas no PL (documentação, testes, explicabilidade, supervisão humana, avaliação de impacto e mitigação contínua) assumem função constitucional de proteção da pessoa humana contra poderes automatizados. A literatura jurídica aponta que tais obrigações não constituem ônus regulatório exagerado, mas instrumentos essenciais de concretização dos direitos fundamentais.

Há também debate relevante sobre o papel das agências reguladoras. O PL 2.338/2023 prevê atuação coordenada de órgãos técnicos, mas ainda é incipiente na definição de competências específicas. A literatura comparada mostra que modelos eficazes de regulação da IA dependem de estruturas técnicas autônomas, capazes de fiscalizar, auditar, impor sanções e produzir padrões de conformidade. O Brasil dispõe de precedentes importantes nesse tema, como a atuação da Autoridade Nacional de Proteção de Dados (ANPD) na LGPD. Parte da doutrina sustenta que a ANPD deveria ter competências ampliadas para tratar da IA, enquanto outra parte defende a criação de autoridade específica. A solução ainda é objeto de intenso debate acadêmico e político.

Outro elemento crucial da regulação brasileira deve ser o fortalecimento da transparência algorítmica, especialmente em órgãos públicos. A Constituição impõe dever de publicidade, motivação e controle das decisões estatais. Quando decisões passam a ser automatizadas, esse dever não desaparece: migra para os algoritmos. Assim, sistemas utilizados pelo Estado devem ser auditáveis, explicáveis e contestáveis. A literatura jurídica enfatiza que o uso de caixas-pretas no setor público é incompatível com o Estado Democrático de Direito. Países como Canadá e Reino Unido já incorporaram obrigações explícitas de transparência em algoritmos governamentais; o Brasil ainda precisa avançar mais neste ponto.

A questão da responsabilidade jurídica também está no centro da construção do modelo brasileiro. A literatura analisada indica que a responsabilidade não deve recair exclusivamente sobre o usuário final da IA, pois este raramente detém controle efetivo sobre desenvolvimento, treinamento e atualização dos sistemas. Modelos modernos adotam o que se convencionou chamar de “responsabilidade escalonada” ou “responsabilidade distribuída”, na qual cada agente envolvido (desenvolvedor, provedor, integrador, operador e usuário) responde conforme sua capacidade de influenciar o risco. Essa lógica é compatível com a teoria civil brasileira, especialmente com a responsabilidade objetiva para atividades de risco e com a responsabilização solidária quando diversos agentes contribuem para o dano.

O debate sobre responsabilidade ganha especial relevância diante da autonomia crescente dos sistemas de IA. Como já apontado pela doutrina, a imprevisibilidade de modelos complexos desafia categorias tradicionais de culpa e nexos causal. A solução encontrada por diversas jurisdições, e defendida em artigos jurídicos nacionais recentes, é abandonar a busca ilusória de causalidade linear e adotar parâmetros de governança: se um agente falhou em cumprir obrigações de supervisão, testes, mitigação e monitoramento contínuo, responde juridicamente pelo dano. A responsabilidade, nesse modelo, deriva do descumprimento dos deveres de segurança e diligência, e não da tentativa inviável de rastrear cada conexão técnica realizada pelo algoritmo.

Por fim, a construção de um modelo brasileiro de regulação da IA deve ser orientada pelo que a doutrina chama de “constitucionalismo digital”. Isso significa compreender que a tecnologia altera estruturas sociais de poder e que a Constituição precisa fornecer diretrizes materiais para limitar esse poder. Ricardo Campos, ao tratar do tema, explica que a IA não apenas desafia direitos tradicionais, mas reconfigura práticas constitucionais fundamentais, exigindo interpretação renovada do devido processo, da igualdade substancial, da privacidade e do papel regulatório do Estado.

A regulação, portanto, não deve ser vista como obstáculo à inovação, mas como instrumento indispensável para garantir que o desenvolvimento tecnológico se dê em conformidade com os valores do Estado Democrático de Direito. A literatura jurídica é unânime em afirmar que inovação sem regulação adequada resulta em assimetria, opacidade e ampliação de vulnerabilidades sociais. Já a inovação regulada produz ambientes seguros, previsíveis e juridicamente estáveis, favorecendo inclusive o desenvolvimento econômico sustentado.

A discussão sobre a construção de um modelo regulatório brasileiro de inteligência artificial precisa levar em conta não apenas experiências estrangeiras, mas também a forma específica como o Brasil lida historicamente com desigualdades, com o papel do Estado e com a relação entre direitos fundamentais e inovação tecnológica. O país vive um paradoxo regulatório: ao mesmo tempo em que busca estimular inovação e competitividade, convive com estruturas sociais profundamente desiguais, que tornam indivíduos e grupos vulneráveis a decisões automatizadas injustas. Assim, a regulação brasileira não pode ser minimalista; precisa ser robusta, preventiva e orientada pelo constitucionalismo de proteção.

Um dos temas mais debatidos na doutrina é a necessidade de uma arquitetura institucional sólida. Países que implementaram regulações bem-sucedidas dispõem de órgãos técnicos altamente especializados, com autonomia e capacidade real de fiscalização. No Brasil, a experiência da ANPD demonstra que a proteção de dados só ganhou efetividade quando acompanhada de uma autoridade dedicada. A literatura majoritária defende que a regulação da IA

exige uma autoridade semelhante, com competências que incluem emissão de normas complementares técnicas, homologação de avaliações de impacto algorítmico, auditorias independentes, fiscalização contínua de sistemas de alto risco, aplicação de sanções administrativas e coordenação com outras agências reguladoras.

Mesmo com esse avanço, ainda falta ao Brasil uma integração institucional mais consistente. Por exemplo, setores como saúde, crédito, segurança e educação usam IA de maneiras específicas, exigindo diálogo entre o órgão regulador central e as agências setoriais (ANS, Bacen, ANS, Anatel etc.). A literatura internacional mostra que modelos “multicamadas” são os mais eficazes.

Outro aspecto central da construção de um modelo brasileiro é a elaboração de obrigações de transparência e explicabilidade adequadas ao contexto nacional. Países desenvolvidos possuem recursos técnicos e infraestrutura fortes para implementar explicabilidade algorítmica complexa. O Brasil, entretanto, precisa de um modelo que equilibre rigor regulatório com viabilidade operacional. A explicabilidade deve ser vista não como obrigação genérica, mas como dever jurídico calibrado ao risco, com diferentes níveis desde explicabilidade para o usuário afetado; explicabilidade para órgãos de fiscalização, explicabilidade para auditorias independentes à transparência pública para sistemas estatais.

A Constituição brasileira, ao exigir motivação e transparência de atos administrativos, praticamente impõe que sistemas de IA utilizados pelo Estado adotem o mais alto grau de explicação possível.

Outro ponto de expansão doutrinária essencial diz respeito ao uso de IA por órgãos públicos. A literatura crítica aponta que decisões estatais automatizadas sem supervisão humana adequada fere o devido processo legal e o princípio republicano. O Estado não pode delegar plenamente sua capacidade decisória a sistemas privados, opacos e incontrolláveis. Nesse sentido, parte dos juristas defende que: (i) sistemas de IA em políticas públicas devem ser de “uso seguro por padrão” com mecanismos explícitos de correção, reversão e contestação; (ii) IA não pode ser usada para decisões automáticas que afetem direitos civis e sociais sem revisão humana; (iii) a utilização de IA pelo Estado precisa obedecer ao princípio da proporcionalidade evitando automações que ampliem desigualdades.

Além disso, o modelo regulatório brasileiro deve lidar com o desafio da responsabilidade civil e administrativa em cadeias complexas de IA. A doutrina já reconhece que não haveria justiça nem segurança jurídica se a responsabilidade recaísse apenas sobre o usuário final. A responsabilidade precisa ser distribuída, com base nos “pontos de controle” exercidos por cada agente do ciclo da IA. Desenvolvedores que escolhem bases de dados inadequadas, fornecedores que liberam modelos com falhas conhecidas e operadores que não supervisionam adequadamente devem responder proporcionalmente ao risco que introduzem.

A teoria brasileira da responsabilidade objetiva por risco, somada ao art. 927 do Código Civil e às práticas consolidadas do CDC, oferece base sólida para adaptar o regime jurídico da IA: não importa se o dano foi causado por ação humana direta; importa se houve falha nos deveres de segurança, cuidado e monitoramento, especialmente em sistemas de alto risco.

Por fim, a doutrina aponta que a regulação da IA no Brasil precisa ser coerente com a lógica de um constitucionalismo digital. Não basta criar obrigações técnicas; é preciso garantir

que o uso da IA esteja alinhado aos valores e princípios constitucionais. Isso coloca o Brasil não apenas como receptor de tendências globais, mas como formulador de um modelo normativo original, capaz de responder às desigualdades e vulnerabilidades que caracterizam sua estrutura social.

3. A Responsabilidade Jurídica e os Modelos Regulatórios na Era da Inteligência Artificial

3.1 Responsabilidade Jurídica na Era Algorítmica

A responsabilização jurídica na era algorítmica representa um dos pontos de maior tensão na regulação contemporânea da inteligência artificial. À medida que sistemas automatizados assumem tarefas decisórias antes reservadas exclusivamente aos seres humanos, tornam-se evidentes lacunas e insuficiências nos modelos tradicionais de responsabilidade civil, administrativa e até penal. O Direito foi historicamente construído sobre categorias antropocêntricas - vontade, culpa, dolo, negligência, previsibilidade, causalidade linear. A IA, ao operar por meio de probabilidades e padrões autônomos aprendidos estatisticamente, rompe com parte dessas categorias, introduzindo fenômenos que não se ajustam às molduras clássicas.

A doutrina contemporânea tem chamado atenção para essa mudança paradigmática. Em diversos artigos jurídicos nacionais recentes, observa-se consenso crescente de que a responsabilidade por danos decorrentes de IA não pode se limitar a identificar o agente humano que acionou ou executou o sistema. As decisões não são meramente instrumentais: elas emergem de modelos treinados com enormes bases de dados, influenciados por escolhas técnicas, arquiteturas algorítmicas, parâmetros opacos e processos de treinamento que distribuem causalidade entre múltiplos atores.

Isso leva ao primeiro grande desafio: a fragmentação do nexo causal. Sistemas de IA não operam em linhas diretas de causa e efeito; operam em redes complexas. Quem alimenta o banco de dados exerce poder decisório. Quem cria o algoritmo exerce poder decisório. Quem faz manutenção, calibragem, fine-tuning, integração e atualização também exerce poder decisório. E, finalmente, quem utiliza o sistema exerce poder decisório. O problema é que todos esses fatores - dados, parâmetros, arquiteturas, ambientes de execução - interferem no resultado final. Determinar quem “causou” um dano muitas vezes se torna tarefa quase impossível.

Por esse motivo, a literatura jurídica sobre IA tem defendido que a responsabilidade não pode depender do modelo clássico de culpa. Ela deve se apoiar em fundamentos normativos novos, capazes de capturar o caráter coletivo e distribuído da produção algorítmica. Surge, então, o conceito de responsabilidade distribuída, que não significa diluição, mas atribuição diferenciada de deveres conforme o grau de controle exercido por cada agente. Desenvolvedores respondem pelos riscos introduzidos pela arquitetura do modelo. Fornecedores respondem pelos erros e limitações comunicadas inadequadamente. Integradores respondem por falhas de uso ou adaptação. Usuários respondem quando rompem parâmetros de operação.

Esse modelo distribui responsabilidades e evita a injustiça de concentrá-las em um único agente. Ele também reforça o caráter preventivo da regulação, pois impõe a cada parte da cadeia um dever de cuidado. O descumprimento desse dever — seja falha de mitigação, falta de supervisão, ausência de testes, negligência no manuseio de dados ou uso indevido — torna-se a

base da responsabilização.

Outro desafio estrutural diz respeito à responsabilidade objetiva. A doutrina mais atual defende a adoção desse modelo para sistemas de alto risco, o que é coerente com o art. 927 do Código Civil e com o regime do Código de Defesa do Consumidor. Isso porque sistemas de IA, especialmente em atividades sensíveis como saúde, crédito, segurança, transporte ou políticas públicas, geram riscos próprios, inerentes à sua complexidade técnica. Tais riscos, ainda que minimizados, não são elimináveis. Assim, quem introduz a tecnologia deve assumir o risco da atividade, independentemente de culpa. Esse raciocínio é amplamente aceito na literatura internacional e o AI Act adota solução semelhante.

Ainda assim, a responsabilidade objetiva não resolve tudo. Há situações em que danos resultam não de risco inerente, mas de negligência, falha na governança, ausência de supervisão humana ou má configuração. Para esses casos, a literatura recomenda sistemas híbridos, nos quais a responsabilidade objetiva se combine com elementos de responsabilidade subjetiva, sempre calibrada à natureza da operação algorítmica.

Outro ponto central é o problema da explicabilidade. Como responsabilizar alguém por uma decisão que nem mesmo especialistas conseguem compreender plenamente? A doutrina jurídica tem insistido que a responsabilidade deve se fundamentar não apenas no resultado final do algoritmo, mas no cumprimento de obrigações processuais de governança. Explicabilidade, assim, torna-se não apenas ferramenta técnica, mas requisito jurídico para imputação de responsabilidade.

Se um agente não consegue explicar minimamente como seu sistema chega a determinados resultados, isso configura violação de deveres de cuidado e, portanto, responsabilidade.

Esse raciocínio inspira o surgimento do que parte da literatura chama de responsabilidade por opacidade. Nesse modelo, não é necessário compreender tecnicamente o funcionamento interno do algoritmo, mas sim demonstrar que houve cumprimento dos deveres de documentação, supervisão, testes e monitoramento contínuo. A falha em cumprir tais deveres é suficiente para imputar responsabilidade, independentemente de se compreender exatamente como a falha técnica ocorreu.

Outra dimensão relevante da responsabilidade algorítmica é a necessidade de reconhecer que sistemas de IA têm impactos coletivos, o que exige respostas regulatórias que ultrapassem litígios individuais. A doutrina tem defendido mecanismos como ações civis públicas, obrigações de reparação coletiva, responsabilização administrativa e mecanismos de supervisão preventiva.

Além disso, a responsabilidade na era algorítmica se entrelaça com princípios constitucionais. O princípio da proporcionalidade pode exigir que tecnologias de alto impacto sejam submetidas a controles mais rigorosos, limitando ou até proibindo seu uso em determinadas situações. O princípio da dignidade humana impõe que decisões automatizadas não podem substituir, reduzir ou comprometer garantias mínimas da personalidade. O devido processo legal exige transparência, revisão humana e mecanismos de contestação. A igualdade exige mitigação ativa de vieses. A publicidade administrativa exige transparência algorítmica. A moralidade administrativa exige governança e supervisão humana efetiva.

Assim, a responsabilidade jurídica na era algorítmica não é apenas questão técnica; é questão constitucional. Ela se fundamenta no entendimento de que tecnologias que exercem poder precisam estar sujeitas às mesmas restrições que limitam o poder humano. No Estado Democrático de Direito, nenhum poder - nem técnico, nem político, nem econômico - pode operar sem controles.

O debate sobre responsabilidade algorítmica, portanto, revela que a regulação da IA não é apenas necessária, mas indispensável para garantir que avanços tecnológicos não comprometam direitos fundamentais, não ampliem desigualdades, não inviabilizem o devido processo e não produzam novas formas de poder obscuro. O Direito precisa responder com soluções normativas coerentes com a complexidade dos sistemas e com a fragilidade humana diante deles.

A complexidade do tema se adensa ainda mais quando se observa que a responsabilização jurídica da IA não envolve apenas danos materiais ou diretos, mas também danos difusos, coletivos e estruturais, que ultrapassam o indivíduo e afetam grupos inteiros, categorias sociais vulneráveis ou a própria ordem constitucional. Esse tipo de dano, amplamente discutido por Hildebrandt e por diversos juristas brasileiros que tratam de discriminação estrutural, exige ferramentas jurídicas capazes de lidar com violações que não se manifestam como eventos isolados, mas como padrões sistêmicos de injustiça algorítmica.

3.1.1 A responsabilidade civil, administrativa e penal na era algorítmica

A responsabilidade civil na era da IA exige ruptura com abordagens tradicionais centradas na culpa. Sistemas algorítmicos de alto risco, como aqueles usados em saúde, crédito, vigilância, educação, seleção de candidatos, policiamento ou políticas públicas, operam com riscos inerentes que independem da intenção ou negligência do operador final. Por esse motivo, a doutrina defende amplamente que o regime jurídico aplicável a tais sistemas deve ser o da responsabilidade objetiva, com base no risco tecnológico da atividade.

No contexto brasileiro, essa solução encontra apoio no art. 927 do Código Civil e na tradição consolidada do Código de Defesa do Consumidor. Assim, agentes que introduzem IA em ambientes sensíveis assumem responsabilidade por danos decorrentes do risco que voluntariamente colocaram em circulação. Isso não exclui a necessidade de responsabilização subjetiva em casos de dolo, negligência ou violação dos deveres de governança, mas complementa o sistema.

No setor público, a responsabilidade administrativa ganha contornos constitucionais. A administração pública é vinculada aos princípios da legalidade, publicidade, motivação e proporcionalidade. Sistemas algorítmicos utilizados pelo Estado, portanto, devem obedecer a padrões ainda mais elevados de transparência e explicabilidade. O STF já sinalizou que decisões automatizadas não podem esvaziar garantias processuais nem impedir o acesso à fundamentação das decisões.

A responsabilidade do Estado pode ocorrer por escolha inadequada de tecnologia, ausência de supervisão humana, uso discriminatório da IA, falha em garantir mecanismos de contestação, aquisição de sistema opaco e não auditável.

Há também a dimensão dos danos coletivos e estruturais, característicos da IA, que frequentemente afetam grupos inteiros como pessoas negras, mulheres, moradores de periferias, usuários de serviços públicos, sem um evento isolado. Nesses casos, a tutela coletiva e os instrumentos de defesa de direitos difusos são essenciais.

No âmbito penal, a responsabilização não recai sobre a IA - que não é sujeito de direito - mas sobre humanos que assumem riscos indevidos ao operar sistemas de alto impacto. A doutrina aponta que o uso negligente de tecnologias previsivelmente perigosas pode configurar culpa consciente ou até dolo eventual, dependendo da situação. A responsabilidade penal da pessoa jurídica também pode se aplicar quando a empresa deixa de implementar medidas mínimas de segurança e governança.

Assim, a responsabilidade na era algorítmica não desaparece: ela se transforma. Abandona a ilusão da culpa individual e reconhece a natureza distribuída, coletiva e frequentemente estrutural.

3.1.2 A responsabilidade constitucional como dever de proteção no constitucionalismo digital

No plano constitucional, a IA inaugura o que parte da doutrina denomina constitucionalismo digital: a necessidade de reinterpretar garantias fundamentais para enfrentar novos centros de poder algorítmico. A Constituição não perdeu validade — ao contrário, tornou-se ainda mais necessária. É ela que impede que sistemas automatizados operem como estruturas invisíveis de poder.

O Estado, como garante dos direitos fundamentais, possui deveres positivos de proteção: deve criar normas, mecanismos de controle e políticas públicas capazes de impedir que tecnologias violem direitos. Se o Estado se omite, incorre em responsabilidade constitucional. Canotilho e Sarlet tratam dos deveres de proteção como parte da eficácia horizontal e vertical dos direitos fundamentais - e na IA isso se torna ainda mais evidente.

Esse dever se expressa em diferentes dimensões:

- Proteção contra discriminação algorítmica: assegurar que bases de dados não reforcem desigualdades históricas.
- Proteção contra opacidade decisória: garantir que decisões automatizadas sejam compreensíveis e contestáveis.
- Proteção contra uso estatal abusivo: impedir reconhecimento facial indiscriminado, vigilância em massa, decisões automatizadas em políticas sociais sem revisão humana.
- Proteção da privacidade e da autodeterminação informativa: assegurar que dados e perfis inferidos não sejam utilizados de forma abusiva.
- Proteção da democracia: garantir pluralidade informacional diante de sistemas de recomendação que distorcem o debate público.

Assim, a responsabilidade constitucional não é retórica: ela exige que o Estado, empresas e desenvolvedores atuem dentro de limites estruturais capazes de proteger o cidadão de danos algorítmicos.

No Brasil, esse entendimento tem guiado a formulação do PL 2338/2023, que adota a lógica da regulação baseada em risco e impõe deveres explícitos de governança, supervisão humana, mitigação de vieses e transparência.

A responsabilidade constitucional é, portanto, o eixo que dá sentido às demais: civil, administrativa e penal. Sem ela, o Direito não seria capaz de proteger sujeitos em um ambiente no qual decisões se deslocam progressivamente para sistemas não humanos.

4. Experiências Internacionais, Casos Práticos e Lições Estruturais para o Modelo Brasileiro de Regulação da Inteligência Artificial

4.1 Casos concretos e violações de direitos fundamentais envolvendo sistemas de IA

Os casos internacionais envolvendo falhas ou abusos no uso de tecnologias de inteligência artificial não configuram episódios isolados ou circunstanciais. Pelo contrário: evidenciam padrões estruturais de risco, demonstrando que sistemas algorítmicos, quando utilizados sem governança adequada, podem produzir violações maciças de direitos fundamentais. Tais precedentes são importantes não apenas como ilustrações, mas como autênticos laboratórios jurídicos, a partir dos quais o Brasil pode extrair lições essenciais para a construção de sua própria regulação.

O caso COMPAS, nos Estados Unidos, tornou-se símbolo mundial de discriminação algorítmica. Investigações da ProPublica mostraram que o sistema não apenas apresentava viés racial, mas que esse viés era reforçado com o uso continuado, criando ciclos de retroalimentação que prejudicavam acusados negros. A controvérsia jurídica não se limitou ao viés, mas ao fato de que as partes não tinham acesso ao funcionamento interno do algoritmo, por ser software proprietário. Assim, decisões judiciais eram tomadas com base em um sistema que não podia ser auditado ou contestado, violando frontalmente o devido processo legal, a publicidade dos atos jurisdicionais e o direito à motivação das decisões.

Outro caso ilustrativo, frequentemente citado por pesquisas acadêmicas, envolve o algoritmo de alocação de recursos em saúde nos Estados Unidos, analisado em estudo da revista Science. O sistema classificava pacientes negros como menos necessitados de cuidado médico em comparação com pacientes brancos em condições semelhantes. A razão não estava em intenção discriminatória, mas no uso de “custo passado com saúde” como proxy de necessidade futura. Como indivíduos negros, historicamente, tiveram menor acesso ao sistema de saúde, o algoritmo interpretava esse menor gasto como menor demanda. Esse caso demonstra como escolhas técnicas aparentemente neutras podem reproduzir desigualdades históricas — alertando reguladores brasileiros para o cuidado na definição de bases de dados e critérios de classificação.

Na Europa, o caso do algoritmo de benefícios infantis da Holanda (Toeslagenaffaire) tornou-se uma verdadeira tragédia institucional. Mais de vinte mil famílias, muitas compostas por imigrantes, foram falsamente acusadas de fraude. O sistema definia perfis de suspeição com base em critérios opacos, classificando famílias inteiras como fraudulentas por pequenos erros

burocráticos. O escândalo revelou falhas: ausência de supervisão humana, ausência de auditorias, discriminação indireta e total opacidade. O governo foi pressionado ao ponto de renunciar. Trata-se de exemplo extremo dos riscos de delegar o poder sancionatório do Estado a sistemas opacos.

Outro caso relevante ocorreu na Austrália, com o sistema RoboDebt. Criado para identificar fraudes em benefícios sociais, o algoritmo gerou milhões de cobranças indevidas. Só anos depois reconheceu-se que o sistema era ilegal e baseado em premissas matemáticas incorretas. A tragédia social resultante - suicídios, endividamento, destruição de famílias - ilustra que o uso indiscriminado de IA em políticas públicas pode causar danos irreversíveis. O governo australiano pagou bilhões em indenizações. A lição para o Brasil é clara: sistemas automatizados não são neutros e podem transformar políticas sociais em aparatos de punição.

No campo do reconhecimento facial, há inúmeros casos de prisões injustas nos Estados Unidos. O mais conhecido é o de Robert Williams, detido em Michigan após ser erroneamente identificado por um sistema policial. Ele ficou preso horas até que a polícia percebesse que o algoritmo havia “confundido” seu rosto com o de um criminoso. Esse caso é emblemático porque evidencia risco duplo: o erro algorítmico e a confiança cega de agentes estatais na precisão da tecnologia. Para países como o Brasil, onde a seletividade penal é profunda, tal tecnologia pode amplificar injustiças já existentes.

Casos envolvendo plataformas digitais também ilustram riscos que o Brasil enfrenta. O escândalo da Cambridge Analytica mostrou que algoritmos podem manipular fluxos de informação e influenciar eleições democráticas. Investigações revelaram que modelos preditivos criaram perfis psicológicos de milhões de pessoas para direcionar propaganda política altamente segmentada. Esse episódio escancarou que a IA não afeta apenas indivíduos, mas estruturas democráticas, alterando o cenário político sem transparência ou controle.

Outros casos relevantes incluem:

- Amazon (EUA): algoritmo de contratação descartava automaticamente candidatas mulheres.
- Twitter e Instagram: estudos revelaram reforço algorítmico de corpos brancos e padrões europeizados.
- França: sistema de detecção fiscal automatizada gerou penalizações incorretas e massivas.
- Índia: sistemas de reconhecimento facial usados em protestos identificavam erroneamente ativistas, criando riscos democráticos.

Esses precedentes formam um quadro robusto: a IA, quando mal regulada, produz danos amplos, profundos e sistematicamente distribuídos.

Além dos casos já analisados, é importante destacar que a literatura jurídica contemporânea tem se dedicado a examinar fenômenos mais recentes e complexos envolvendo IA generativa, sistemas preditivos aplicados a políticas públicas e modelos automatizados usados por plataformas digitais para moderar conteúdo e estruturar interações online. Esses casos

ampliam o panorama de riscos e demonstram que os desafios não se limitam mais a decisões classificatórias ou preditivas, mas alcançam também a própria construção da realidade informacional.

Um exemplo emblemático aparece com o uso de IA generativa por plataformas tecnológicas que passaram a produzir textos, imagens e vídeos com alto grau de verossimilhança. Diversos relatórios internacionais documentam episódios em que deepfakes foram utilizados para manipular percepções sociais, influenciar processos eleitorais e atacar pessoas públicas e privadas. Em países europeus, houve episódios de vídeos gerados por IA utilizados para criar falsas acusações contra candidatos, influenciar debates públicos e disseminar boatos em larga escala. A capacidade da IA generativa de criar conteúdo convincente desafia diretamente garantias democráticas, exigindo que Estados adotem mecanismos normativos contra manipulação digital — tema que impacta o Brasil em ano eleitoral.

Outro caso relevante envolve o uso de sistemas de pontuação social privada por empresas de seguros, bancos e plataformas de mobilidade. Relatórios recentes mostraram que certas empresas utilizaram IA para prever comportamento de clientes e, a partir disso, estabelecer preços diferenciados, limitar acesso a determinados serviços ou até bloquear contas. Tais práticas, quando desprovidas de transparência, constituem formas privadas de scoring social, vedadas por legislações modernas como o AI Act europeu. No Brasil, onde a desigualdade estrutural e o histórico de exclusão financeira são marcantes, esse tipo de uso pode reforçar barreiras de acesso ao crédito, saúde e mobilidade urbana.

Há também estudos sobre a utilização de IA para controle migratório, especialmente em fronteiras dos Estados Unidos e da União Europeia. Sistemas de análise comportamental automatizada foram usados para identificar pessoas supostamente mentirosas em entrevistas de imigração, sem qualquer validação científica. Auditorias independentes classificaram essas tecnologias como pseudocientíficas, além de discriminatórias. Esse tipo de caso demonstra como tecnologias opacas podem afetar o direito humano à migração segura e ao devido processo, aspectos cruciais também para o Brasil, especialmente em regiões fronteiriças e no contexto de fluxos migratórios crescentes vindos da América Latina e do Caribe.

Outro campo relevante envolve sistemas de IA em educação, além do caso britânico já citado. Nos Estados Unidos, escolas públicas adotaram softwares de vigilância comportamental que monitoravam falas de estudantes em plataformas digitais. Alegando fins de “prevenção de violência”, os algoritmos monitoravam mensagens privadas, capturas de tela e pesquisas realizadas, o que gerou inúmeras violações de privacidade, especialmente entre adolescentes. Estudos mostraram que estudantes negros e latinos eram mais frequentemente classificados como “ameaças” pelo sistema. A extrapolação de vigilância algorítmica sobre populações juvenis configura violação grave da privacidade informacional, algo que preocupa o contexto brasileiro em que tecnologias similares começam a ser ofertadas para escolas públicas com pouca regulação.

Também é relevante mencionar experiências negativas envolvendo sistemas de IA aplicados à habitação. Em algumas cidades europeias, algoritmos usados para ranquear candidatos a moradias sociais privilegiaram perfis de famílias cujos dados indicavam “estabilidade financeira”, excluindo imigrantes, mães solo e famílias de baixa renda. Esses casos revelam que a IA pode transformar políticas públicas em mecanismos de exclusão automatizada, distorcendo o sentido original de programas sociais destinados à proteção de vulneráveis.

No setor privado, episódios envolvendo plataformas de entrega e transporte revelam formas de vigilância laboral automatizada. Empresas passaram a utilizar IA para monitorar produtividade, avaliar desempenho e até demitir trabalhadores, muitas vezes sem explicação clara sobre critérios de avaliação. Relatórios de organizações trabalhistas internacionais apontam que motoristas e entregadores foram penalizados por “inconsistências algorítmicas” geradas por falhas de GPS ou por erros de classificação da plataforma. O impacto sobre a dignidade laboral é direto — e especialmente sensível no Brasil, onde grande parte do trabalho digitalizado está alocada em plataformas.

Por fim, há casos envolvendo sistemas de análise de crédito automatizado que, ao utilizarem dados de redes sociais para prever “comportamento econômico”, acabaram penalizando usuários com base em seu círculo social, suas interações e até seu padrão linguístico. Isso constitui violação da autodeterminação informativa e cria uma espécie de “determinismo algorítmico”, onde o indivíduo passa a ser avaliado não por suas próprias ações, mas pelo comportamento do grupo ao qual o algoritmo o associa. No Brasil, onde crédito e consumo são marcadores centrais da vida social, esse tipo de prática pode aprofundar desigualdades já existentes.

Esses casos demonstram que a IA, quando usada sem governança adequada, produz danos que vão muito além do indivíduo: ela afeta grupos vulneráveis, distorce políticas públicas, interfere em processos democráticos e compromete garantias constitucionais. O Brasil, ao estruturar sua regulação, precisa considerar não apenas os riscos técnicos, mas os danos sociais, políticos e institucionais já documentados no mundo.

4.2 Lições normativas para o Brasil a partir de modelos estrangeiros e debates contemporâneos

Os modelos regulatórios internacionais oferecem ao Brasil uma base rica para reflexão, mas não um manual de transplante automático. O AI Act da União Europeia, por exemplo, tornou-se o marco regulatório mais detalhado e abrangente do mundo. Sua estrutura baseada em risco - proibido, alto risco, risco limitado e risco mínimo - oferece clareza e segurança jurídica. Porém, sua implementação depende de órgãos reguladores altamente capacitados, infraestrutura regulatória avançada e forte cultura de conformidade. O Brasil pode se inspirar nessa lógica, mas precisa adaptá-la às suas capacidades institucionais.

Outro ponto crucial da experiência europeia é a exigência de avaliações de impacto algorítmico (AI impact assessments), que se tornaram obrigatórias para sistemas de alto risco. Esse tipo de avaliação obriga empresas e órgãos públicos a identificarem, previamente, riscos de discriminação, falhas técnicas, efeitos sociais e impactos individuais. Para o Brasil, esse mecanismo é particularmente importante, dada a força normativa da igualdade substancial no ordenamento nacional e a magnitude das desigualdades estruturais existentes.

A OCDE também fornece princípios essenciais, como transparência, segurança, robustez e accountability, amplamente aceitos por jurisdições ao redor do mundo. Tais princípios podem ser incorporados diretamente no marco regulatório brasileiro, pois encontram amparo na Constituição brasileira e especialmente nos princípios da publicidade, moralidade, eficiência e dignidade humana.

A China oferece modelo distinto, de controle centralizado, com ênfase em segurança de Estado, moderação de conteúdo e forte regulação de plataformas digitais. Embora esse modelo não seja compatível com a lógica democrática brasileira, ele traz lições operacionais relevantes: exigência de documentação obrigatória, supervisão contínua e responsabilidade clara de plataformas por conteúdos gerados por sistemas automatizados.

Nos Estados Unidos, a ausência de marco regulatório federal unificado foi parcialmente compensada por diretrizes técnicas, como o NIST AI Risk Management Framework, que estabeleceu referências objetivas para avaliar risco, mitigar falhas e organizar governança. O Brasil pode se beneficiar dessa abordagem ao implementar normas complementares técnicas voltadas à operacionalização do PL 2.338/2023.

Além disso, a experiência internacional revela a necessidade de autoridades reguladoras especializadas, com autonomia e capacidade técnica real, auditorias independentes obrigatórias para sistemas de alto impacto, governança contínua, e não apenas conformidade prévia, transparência perante órgãos públicos e perante o usuário afetado, restrições explícitas ao uso estatal de tecnologias de vigilância, mecanismos de contestação humana obrigatória para decisões automatizadas e proibição de determinados sistemas considerados incompatíveis com direitos fundamentais.

Essas lições são fundamentais para o Brasil porque evidenciam que a regulação da IA não pode ser apenas principiológica. Ela exige mecanismos operacionais claros, instrumentos de fiscalização robustos e padrões de governança com força normativa. O risco brasileiro — amplamente apontado por juristas nacionais — é adotar regulação apenas retórica, sem os instrumentos necessários para dar-lhe efetividade.

O contexto brasileiro exige adaptação, não mera importação. O país precisa criar sistema regulatório que enfrente desigualdades raciais, socioeconômicas e territoriais, características próprias da realidade nacional. Além disso, o Brasil tem tradição constitucional que valoriza deveres positivos do Estado — o que facilita a adoção de políticas públicas de fiscalização, auditoria e mitigação de riscos.

Assim, o estudo comparado não serve para copiar, mas para mostrar caminhos possíveis, riscos reais e oportunidades. A regulação brasileira deve aprender com erros de outros países para evitar catástrofes jurídicas, sociais e institucionais que a IA já provocou em diversas partes do mundo.

A partir dos casos analisados, é possível extrair lições essenciais para qualquer país que deseje regular IA de maneira consistente. O primeiro ponto é a constatação de que transparência é pré-requisito básico, não recomendação voluntária. Sistemas decisórios que não podem ser auditados, compreendidos ou contestados produzem zonas de irresponsabilidade, blindando tanto desenvolvedores quanto órgãos públicos de escrutínio democrático. O Brasil, cuja Constituição exige publicidade e motivação dos atos administrativos, tem base jurídica sólida para exigir níveis elevados de transparência algorítmica.

Outro aprendizado é a importância da explicabilidade adaptada ao contexto. A União Europeia tem insistido que modelos de IA de alto risco devem ser acompanhados de documentação compreensível não apenas para especialistas, mas para usuários afetados e

autoridades reguladoras. A experiência internacional indica que sem explicabilidade suficiente, as garantias processuais são esvaziadas. O Brasil pode (e deve) incorporar explicabilidade como direito fundamental procedimental.

Os modelos estrangeiros também mostram que a mitigação de vieses é obrigação jurídica e não mera boa prática técnica. O caso holandês é particularmente importante porque evidencia que sistemas estatais, quando discriminam, violam direitos humanos em escala ampliada. O Brasil, com seu histórico de discriminação estrutural, precisa implementar avaliações de impacto algorítmico obrigatórias para sistemas usados em políticas públicas.

Outro ponto fundamental é o reconhecimento de que nem toda tecnologia pode ser regulada; algumas precisam ser proibidas. O AI Act europeu proíbe manipulação subliminar, scoring social e vigilância biométrica em tempo real. Vários municípios norte-americanos proibiram reconhecimento facial estatal. O Brasil precisa iniciar esse debate, sobretudo nas áreas de segurança pública e vigilância urbana.

A experiência internacional também demonstra que autoridades especializadas são indispensáveis. O Brasil, por meio da ANPD, já deu o primeiro passo ao criar um ente regulador para proteção de dados. Contudo, a IA exige competências adicionais: certificação de sistemas, análise técnica aprofundada, inspeção de modelos, auditorias estruturais e supervisão permanente. A criação de uma Autoridade Nacional de Inteligência Artificial é recomendada pela literatura.

Modelos de governança internacionais também demonstram que a responsabilidade distribuída é a solução mais adequada. A Europa já discute regimes de responsabilidade escalonada; a OCDE propõe abordagens de governança compartilhada; os Estados Unidos, embora sem marco federal, reconhecem que desenvolvedores e operadores não podem se eximir de responsabilidade.

Por fim, há lições sobre democracia. A IA afeta o debate público, a circulação de informações e a formação da opinião. Países avançados têm discutido formas de limitar a manipulação algorítmica, exigir transparência na publicidade política automatizada e fortalecer a regulação de plataformas digitais. O Brasil, cuja democracia já foi abalada por ondas de desinformação e manipulação digital, precisa incorporar essas preocupações no debate regulatório da IA.

Assim, as experiências internacionais fornecem parâmetros claros do que não deve ser feito, do que deve ser evitado, do que pode ser adotado com adaptações e do que deve ser implementado imediatamente.

Além das lições preliminares já discutidas, o exame aprofundado dos modelos internacionais evidencia que a regulação da IA não é apenas um conjunto de obrigações técnicas, mas um projeto político-institucional, cuja eficácia depende de escolhas estruturantes sobre governança, fiscalização, responsabilidade e limites éticos. Cada jurisdição estudada apresenta virtudes e falhas que o Brasil deve observar com cautela.

Um dos aspectos mais relevantes extraídos da experiência europeia é a centralidade das avaliações de impacto algorítmico. No AI Act e em legislações complementares, esse instrumento configura verdadeiro “ pilar procedimental ” da regulação. Avaliações de impacto

exigem que desenvolvedores e operadores antecipem riscos de discriminação, explique os objetivos do sistema, documentem métodos, examinem impactos sociais e identifiquem potenciais violações de direitos fundamentais antes da implementação. No contexto brasileiro, esse instrumento se mostra indispensável, considerando que grande parte dos danos potenciais da IA decorre de desigualdades pré-existentes, que tendem a ser reproduzidas matematicamente caso não sejam previamente mitigadas. Essas avaliações funcionam como mecanismos de prudência constitucional, incorporando a lógica dos deveres estatais de proteção.

Outra lição importante diz respeito ao papel das autoridades reguladoras independentes. Países que implementaram regulações eficazes, como Canadá, Reino Unido e membros da União Europeia, possuem órgãos técnicos com alto grau de autonomia, capazes de produzir diretrizes regulatórias, fiscalizar sistemas e impor sanções. Sem autoridade especializada, o marco normativo torna-se letra morta. Isso é particularmente relevante para o Brasil, que possui tradição de agências reguladoras, mas cuja efetividade depende de condições políticas, técnicas e orçamentárias. A ANPD demonstra como um órgão especializado pode influenciar positivamente a proteção de dados. Uma Autoridade de IA, caso criada ou acoplada à ANPD, precisaria ter prerrogativas ampliadas, corpo técnico altamente qualificado e capacidade de atuação transversal.

Também é fundamental observar como a experiência internacional tem lidado com o tema da supervisão humana obrigatória. O AI Act estabelece que sistemas de alto risco devem ser acompanhados por supervisores humanos treinados e capazes de intervir, interromper, revisar ou corrigir decisões automatizadas. Essa exigência decorre do reconhecimento de que a autonomia técnica dos sistemas não elimina a necessidade de agentes humanos continuarem responsáveis por resultados. No Brasil, isso se conecta diretamente ao princípio da motivação administrativa e ao devido processo legal. Nenhuma decisão estatal substancial pode ser tomada sem possibilidade de revisão humana real e informada.

No campo da vigilância e do reconhecimento facial, a tendência internacional é de restrição, não de expansão. A experiência de cidades norte-americanas e de autoridades europeias indica que tais tecnologias apresentam riscos inaceitáveis de discriminação e violação da privacidade, especialmente quando usadas em espaços públicos. Para o Brasil — país com histórico de violência policial e seletividade penal —, essa experiência recomenda prudência máxima. Qualquer uso de reconhecimento facial precisa ser considerado à luz de vieses raciais estruturais, evitando-se a adoção de sistemas que possam reforçar estereótipos criminais e produzir injustiças em massa.

Além disso, a experiência internacional mostra que a regulação deve ser dinâmica, acompanhando a evolução tecnológica. Normas estáticas tornam-se rapidamente obsoletas. Europa e EUA utilizam mecanismos de sandboxes regulatórios, consultorias públicas contínuas, revisões periódicas de normas e atualizações de padrões técnicos. Isso revela que a regulação da IA exige adaptação constante, diálogo social e revisão permanente. O Brasil, portanto, deve incorporar mecanismos de atualização regulatória automática, evitando congelar a legislação em modelos que podem se tornar inadequados em poucos anos.

Outro ponto central é a necessidade de regulação de plataformas digitais, especialmente no que diz respeito à moderação algorítmica de conteúdo e ao impacto no debate público. A União Europeia implementou o Digital Services Act (DSA), que obriga plataformas a fornecer transparência sobre algoritmos, explicar critérios de recomendação e oferecer possibilidade de desativação desses sistemas. Os Estados Unidos debatem modelos semelhantes. Para o Brasil,

onde as redes sociais desempenham papel crítico na formação da opinião pública e na circulação de informação, essas medidas são essenciais. A IA, quando usada para moderar conteúdo, deve obedecer a critérios de proporcionalidade, transparência e accountability, sob pena de afetar diretamente a liberdade de expressão e pluralismo político.

O contexto internacional também evidencia que dados não são apenas insumos técnicos, mas bens políticos, econômicos e sociais que afetam diretamente a autonomia dos indivíduos. Por isso, modelos regulatórios avançados colocam ênfase em minimização de dados; salvaguardas contra perfilamento excessivo; limites para inferências automatizadas; proteção de grupos vulneráveis e regras específicas para uso de dados sensíveis.

Como o Brasil possui desigualdades amplas e históricas envolvendo raça, renda, território e acesso a bens públicos, esses princípios adquirem especial relevância. Sistemas de IA que operam com dados enviesados podem perpetuar desigualdades sob aparência de neutralidade técnica.

Por fim, é importante reconhecer que as experiências internacionais demonstram que a IA gera desafios para soberania digital. Países que não desenvolvem regulações próprias tornam-se dependentes de modelos e produtos de empresas estrangeiras, o que pode resultar em submissão a interesses econômicos globais e perda de autonomia regulatória. O Brasil precisa garantir que sua legislação dialogue com modelos globais, mas preserve sua autonomia constitucional, social e cultural. Regulamentações importadas sem adaptação às realidades nacionais correm o risco de aprofundar desigualdades internas.

Assim, a análise comparada revela que o Brasil deve adotar modelo baseado em risco, instituir autoridade reguladora forte, garantir transparência e explicabilidade, assegurar supervisão humana, proibir sistemas incompatíveis com direitos fundamentais, exigir avaliações de impacto, criar mecanismos de atualização normativa, proteger dados e grupos vulneráveis, regular plataformas digitais e resguardar soberania tecnológica.

Esses elementos formam o núcleo das lições internacionais aplicáveis ao Brasil, compondo uma base sólida para construção de uma regulação constitucionalmente adequada da inteligência artificial.

5. Conclusão

O percurso analítico desenvolvido ao longo deste artigo demonstrou que a inteligência artificial se tornou um fenômeno jurídico-político de impacto estrutural, capaz de reorganizar práticas institucionais, reconfigurar direitos fundamentais e introduzir novas formas de poder que desafiam diretamente o constitucionalismo contemporâneo. A IA não é apenas tecnologia, mas um dispositivo normativo que produz efeitos concretos, frequentemente opacos, sobre indivíduos, coletividades e instituições. Por isso, sua regulação não pode ser compreendida como mero ajuste técnico, mas como exigência democrática, constitucional e civilizatória.

Os casos internacionais examinados revelam que os riscos não são abstratos ou hipotéticos: são reais, materiais e já produziram injustiças graves. Os episódios envolvendo discriminação algorítmica, erros em políticas públicas automatizadas, decisões judiciais baseadas em sistemas

opacos e vigilância tecnológica indevida mostram que a IA, quando desprovida de governança, amplia desigualdades, reforça estruturas discriminatórias e compromete direitos fundamentais essenciais — especialmente em sociedades marcadas por vulnerabilidades históricas, como é o caso brasileiro.

Da mesma forma, a análise comparada dos modelos regulatórios evidencia que não existe solução única ou universal. A União Europeia avança com uma abordagem baseada em risco, robusta e orientada à proteção de direitos; os Estados Unidos privilegiam diretrizes técnicas e abordagens setoriais; a China utiliza modelo centralizado e estratégico; outros países experimentam soluções híbridas, contextualizadas às suas realidades sociopolíticas. Para o Brasil, a lição central é que a regulação efetiva depende de adaptação, não de transplante. É indispensável considerar a realidade nacional: desigualdade estrutural profunda, seletividade penal, fragilidades institucionais e histórico de discriminação racial, territorial e socioeconômica.

O estudo também evidenciou que a responsabilidade jurídica na era algorítmica exige ruptura com modelos tradicionais baseados exclusivamente em culpa e causalidade linear. A fragmentação do nexos causal em sistemas complexos e a natureza distribuída dos danos exigem novos fundamentos normativos, como deveres de governança, avaliação de riscos, supervisão humana efetiva e responsabilidade distribuída ao longo da cadeia algorítmica. A responsabilidade civil, administrativa, penal e constitucional deve ser reinterpretada à luz do fenômeno técnico, sem perder de vista os valores fundantes da Constituição de 1988.

Por fim, sustenta-se que o Brasil precisa adotar um modelo regulatório que combine rigor ético e segurança jurídica com incentivo à inovação responsável. Devem ser implementados mecanismos de transparência, explicabilidade, prevenção de vieses, auditorias independentes, fiscalização robusta e limitação de usos incompatíveis com direitos fundamentais. A criação ou fortalecimento de autoridade reguladora especializada é peça indispensável para dar efetividade ao marco legal.

Em síntese, a inteligência artificial inaugura uma agenda normativa que não pode ser tratada como modismo ou simples atualização legislativa. Trata-se de desafio estrutural que demanda responsabilidade política, rigor acadêmico e compromisso constitucional. O Brasil tem a oportunidade — e a obrigação — de construir um modelo regulatório que proteja a dignidade humana, promova justiça algorítmica e fortaleça o Estado Democrático de Direito diante de um futuro cada vez mais mediado por tecnologias inteligentes.

A complexidade que gravita em torno da inteligência artificial reforça que o Direito brasileiro está diante de um ponto de inflexão histórico. Diferentemente de outras inovações tecnológicas, cujos impactos podiam ser assimilados de forma incremental, a IA exige do ordenamento jurídico uma postura reconstrutiva, capaz de revisitar conceitos, redefinir instituições e remodelar estratégias regulatórias. Não se trata apenas de disciplinar uma tecnologia emergente, mas de enfrentar um fenômeno que atravessa, simultaneamente, esferas econômicas, políticas, culturais e constitucionais.

A análise realizada também evidenciou que uma regulação adequada da IA não pode ser reduzida ao binômio “proteger versus inovar”. Essa perspectiva, frequentemente apresentada em debates superficiais, ignora que a proteção de direitos fundamentais e a inovação tecnológica não são agendas opostas, mas interdependentes. Sem proteção, a inovação se converte em fonte de insegurança jurídica, desigualdade e captura de poder. Sem inovação responsável, a proteção se

torna inócua e descolada da realidade digital que atravessa o cotidiano das pessoas. O equilíbrio entre esses dois polos não é estático, mas dinâmico, exigindo vigilância institucional contínua.

Outro ponto que emerge com força do estudo é que a regulação da IA precisa ser concebida como investimento democrático. Países que não regulam tornam-se dependentes tecnológica e juridicamente de empresas que operam algoritmos opacos e imposições técnicas que escapam ao controle público. Essa dependência fragiliza a soberania, compromete políticas públicas, enfraquece a transparência e permite que centros privados de decisão passem a exercer funções que, em regimes democráticos, pertencem ao Estado e à sociedade civil. O Brasil, ao formular seu marco regulatório, tem a oportunidade de se posicionar como protagonista regional na governança tecnológica, evitando que seu futuro digital seja ditado exclusivamente por corporações estrangeiras.

Além disso, a análise mostrou que sociedades marcadas por desigualdades profundas — como o Brasil — enfrentam riscos ampliados quando adotam sistemas de IA sem regulação robusta. A tecnologia não opera no vácuo: ela interage com estruturas históricas, reforçando padrões de exclusão racial, socioeconômica, territorial e de gênero. Assim, a regulação nacional precisa incorporar o princípio da igualdade material como núcleo da governança algorítmica, exigindo avaliações de impacto, mitigação de vieses e salvaguardas específicas para grupos vulnerabilizados. Sem essas medidas, a IA corre o risco de acelerar e aprofundar injustiças já consolidadas.

Também se reforça, ao final da análise, que a regulação da IA deve ser concebida como uma agenda interinstitucional. Nenhuma instituição isolada é capaz de enfrentar o desafio: legislativo, judiciário, executivo, agências reguladoras, universidades e sociedade civil organizada precisam atuar de maneira coordenada. A produção de normas deve dialogar com ciência de dados, engenharia, filosofia, sociologia e ciência política, refletindo a natureza multidisciplinar do fenômeno algorítmico. A ausência de diálogo entre saberes produz regulações frágeis, incapazes de compreender a complexidade técnica dos sistemas.

Do ponto de vista constitucional, este estudo reforça que a IA atua como verdadeiro teste de estresse para o Estado Democrático de Direito. Transparência, motivação, proporcionalidade, controle jurisdicional, igualdade e dignidade humana precisam ser reinterpretados para manter sua força normativa frente a um cenário em que decisões são tomadas por sistemas automatizados. Não se trata de criar direitos novos, mas de assegurar que os direitos existentes continuem válidos diante de um novo ambiente tecnológico.

Por fim, este trabalho indica que o Brasil se encontra diante de uma janela de oportunidade rara: construir um marco regulatório de IA antes da consolidação irreversível de práticas privadas e governamentais já estruturadas. Diferentemente do que ocorreu com a proteção de dados — em que a legislação veio após décadas de uso intensivo de tecnologias de vigilância e coleta de dados — a regulação da IA pode ser formulada preventivamente. Aproveitar essa oportunidade significa evitar danos massivos, promover justiça algorítmica, fortalecer instituições e orientar o uso da tecnologia para fins socialmente legítimos.

Diante de tudo isso, conclui-se que a regulação da inteligência artificial deve ser encarada como imperativo constitucional e democrático, indispensável para garantir que a tecnologia opere a serviço da sociedade, e não o contrário. Um marco regulatório robusto, contextualizado às realidades brasileiras e inspirado nas melhores práticas internacionais, permitirá que o país avance

rumo a um futuro digital justo, transparente e alinhado aos valores fundamentais da Constituição de 1988.

6. Referências

AMARAL, Gustavo. **Inteligência artificial e discriminação algorítmica**. São Paulo: Revista dos Tribunais, 2023.

BIONI, Bruno Ricardo. **Proteção de Dados Pessoais: A função e os limites do consentimento**. 2. ed. Rio de Janeiro: Forense, 2021.

BYGRAVE, Lee A. **Data Privacy Law: An International Perspective**. Oxford: Oxford University Press, 2014.

CAMPOS, Ricardo. **Constitucionalismo algorítmico: Direito, tecnologia e democracia**. São Paulo: Marcial Pons, 2022.

CANOTILHO, J. J. Gomes. **Direito Constitucional e Teoria da Constituição**. 7. ed. Coimbra: Almedina, 2016.

CARVALHO, Vinicius M.; ROSSI, Lucas. **Inteligência artificial, explicabilidade e devido processo**. Revista de Direito e Tecnologia, v. 12, n. 2, 2023.

CATH, Corien. **Governing artificial intelligence: Ethics, law, and policy for the 21st century**. Philosophical Transactions of the Royal Society A, v. 376, 2018.

COSTA, E. **Direito e algoritmos: a reconstrução da responsabilidade civil na era digital**. Revista Direito, Estado e Sociedade, n. 63, 2023.

DONEDA, Danilo. **Da privacidade à proteção de dados pessoais**. Rio de Janeiro: Renovar, 2006.

EUROPEAN UNION. AI Act – Artificial Intelligence Act. European Parliament and Council, 2024.

FRAZÃO, Ana. **Direito Concorrencial e Proteção de Dados**. São Paulo: Revista dos Tribunais, 2021.

HILDEBRANDT, Mireille. **Law for Computer Scientists and Other Folk**. Oxford: Oxford University Press, 2020.

LEONARDI, Marcel; BARBOSA, Matheus. **Responsabilidade civil por danos decorrentes de algoritmos**. Revista de Direito Civil Contemporâneo, v. 31, 2022.

LESSIG, Lawrence. **Code and Other Laws of Cyberspace**. New York: Basic Books, 1999.

MITTELSTADT, Brent et al. The ethics of algorithms: Mapping the debate. Big Data & Society, v. 3, n. 2, 2016.

NIST – National Institute of Standards and Technology. **AI Risk Management Framework**. U.S. Department of Commerce, 2023.

OCDE – Organização para a Cooperação e Desenvolvimento Econômico. **OECD Principles on Artificial Intelligence**. Paris, 2019.

PASQUALE, Frank. **The Black Box Society: The Secret Algorithms That Control Money and Information**. Cambridge: Harvard University Press, 2015.

PROPUBLICA. **Machine Bias: Investigating racial bias in algorithms**. ProPublica Report, 2016.

SARLET, Ingo Wolfgang; MARINONI, Luiz Guilherme; MITIDIERO, Daniel. **Curso de Direito Constitucional**. 7. ed. São Paulo: Saraiva, 2022.

SCHREIBER, Anderson. **Responsabilidade civil e tecnologia**. São Paulo: Atlas, 2020.

STF – Supremo Tribunal Federal. **Impactos da Inteligência Artificial no Constitucionalismo Contemporâneo**. Brasília: Supremo Tribunal Federal, 2024.

SUNSTEIN, Cass. **Republic.com 2.0**. Princeton: Princeton University Press, 2007.

UNESCO. **Recommendation on the Ethics of Artificial Intelligence**. Paris, 2021.

ZUBOFF, Shoshana. **The Age of Surveillance Capitalism**. New York: PublicAffairs, 2019.