



PONTIFÍCIA UNIVERSIDADE CATÓLICA DE GOIÁS
ESCOLA DE DIREITO, NEGÓCIOS E COMUNICAÇÃO
NÚCLEO DE PRÁTICA JURÍDICA
COORDENAÇÃO ADJUNTA DE TRABALHO DE CURSO
MONOGRAFIA JURÍDICA

**A RESPONSABILIDADE CIVIL POR ERROS COMETIDOS PELA INTELIGÊNCIA
ARTIFICIAL**

O DILEMA ENTRE CRIADOR, USUÁRIO E A PRÓPRIA MÁQUINA EM UM
CONTEXTO DE MODERNIZAÇÃO CRESCENTE

ORIENTANDO: LUCAS PALMEIRO INOHONA

ORIENTADOR: PROF.MS. ERNESTO MARTIM S. DUNCK

GOIÂNIA
2026

LUCAS PALMEIRO INOHONA

**A RESPONSABILIDADE CIVIL POR ERROS COMETIDOS PELA INTELIGÊNCIA
ARTIFICIAL**

O DILEMA ENTRE CRIADOR, USUÁRIO E A PRÓPRIA MÁQUINA EM UM
CONTEXTO DE MODERNIZAÇÃO CRESCENTE

Monografia Jurídica apresentada à disciplina
Trabalho de Curso II, da Escola de Direito e
Relações Internacionais, Curso de Direito, da
Pontifícia Universidade Católica de Goiás
(PUC-GOIÁS). Prof. Orientador: Ernesto
Martim S. Dunck.

GOIÂNIA
2026

LUCAS PALMEIRO INOHONA

**A RESPONSABILIDADE CIVIL POR ERROS COMETIDOS PELA INTELIGÊNCIA
ARTIFICIAL**

**O DILEMA ENTRE CRIADOR, USUÁRIO E A PRÓPRIA MÁQUINA EM UM
CONTEXTO DE MODERNIZAÇÃO CRESCENTE**

Data da Defesa: _____ de _____ de 2026

BANCA EXAMINADORA

Orientador: Professor Ernesto Martim S. Dunck

Nota

Prof.^a Professor José Aluisio e Araujo Junior

Nota

AGRADECIMENTOS

Em meio a toda pressão que um estudante passa durante a graduação, não poderia deixar de agradecer e celebrar a conquista de realizar esse trabalho.

Primeiramente, expresso minha gratidão a Deus, por me proteger, consolar nos momentos de estresse e sempre me guiar pelos caminhos.

A minha mãe Maria Valéria, por ser sinônimo de alicerce na minha vida.

Ao meu pai Sandro Inohona, que me ensinou sobre o peso das escolhas e das responsabilidades da vida, um exemplo de homem para mim.

A minha família, tanto materna quanto paterna, por demonstrar que família é o melhor presente que Deus pode dar a alguém.

A todos meus amigos, desde os mais antigos, como o Pedro Caixeta e o João Victor Brito que estiveram comigo desde o ensino fundamental, até meus amigos mais recentes.

Também agradeço ao meu professor e orientador Ernesto Martim S. Dunck, que com paciência e didática me orientou por este trajeto final rumo a conclusão da graduação.

RESUMO

A presente monografia tem como objetivo discutir sobre a responsabilidade civil oriunda por erros cometidos pela Inteligência Artificial e quem deveria ser responsabilizado perante a isso. Utilizando-se do método bibliográfico, juntamente com a análise crítica a legislação, doutrina, artigos e a dados de pesquisas. Diante disso, é possível discorrer sobre as nuances do direito civil referente à responsabilidade civil por erros da Inteligência Artificial, pontuando que mesmo não sendo previsto na legislação brasileira, ele é preservado pelo ordenamento jurídico atual, que visa priorizar a dignidade da pessoa humana.

Palavras-chave: Direito. Civil. Responsabilidade Civil. Inteligência Artificial.

ABSTRACT

This monograph aims to discuss civil liability arising from errors committed by Artificial Intelligence and who should be held responsible for them. Using the bibliographic method, along with a critical analysis of legislation, doctrine, articles, and research data, it is possible to discuss the nuances of civil law regarding civil liability for errors in Artificial Intelligence, noting that even though it is not explicitly provided for in Brazilian legislation, it is preserved by the current legal system, which prioritizes the dignity of the human person.

Keywords: Law. Civil. Civil Liability. Artificial Intelligence.

Sumário

RESUMO	5
ABSTRACT	5
INTRODUÇÃO.....	7
1. A ERA DA INTELIGÊNCIA ARTIFICIAL E O DILEMA JURÍDICO.....	11
1.1 CONCEITO DA INTELIGÊNCIA ARTIFICIAL.....	11
1.2 ORIGEM E FUNCIONAMENTO.....	14
1.3 O CENÁRIO DA MODERNIZAÇÃO E O DILEMA DA IA.....	16
2. CRÍTICA AO REGIME ATUAL E A BUSCA POR NOVOS PARADIGMAS DE RESPONSABILIDADE.....	19
2.1 O MARCO LEGAL DA INTELIGÊNCIA ARTIFICIAL NO CENÁRIO NACIONAL.....	20
2.2 ANÁLISE LEGISLATIVA E JURISPRUDENCIAL.....	21
2.3 PROBLEMAS DO USO IRREGULAR DA IA.....	24
2.4 O DESAFIO DA IMPUTAÇÃO E O ROMPIMENTO DO NEXO CAUSAL.....	26
3 ADAPTAÇÃO DE LEIS PARA UM NOVO TEMPO	28
3.1 REFLEXO NO PODER JUDICIÁRIO	31
3.2 A RESPONSABILIDADE CIVIL POR ATOS INDEPENDENTES DA INTELIGÊNCIA ARTIFICIAL	32
3.3 MEDIDAS ÉTICAS E REGULATÓRIAS PARA O CONTROLE E A SEGURANÇA DA IA...	33
3.4 REQUISITOS NECESSÁRIOS PARA ANÁLISE DA RESPONSABILIDADE CIVIL EM CASOS DE ERROS.....	34
CONCLUSÃO	36
REFERÊNCIAS	38

INTRODUÇÃO

Nas últimas décadas, a inteligência artificial passou de tema de ficção científica para elemento presente em serviços bancários, plataformas digitais, atendimento ao consumidor, diagnósticos em saúde e até mesmo em atividades ligadas à segurança pública e ao Poder Judiciário. Sistemas capazes de aprender com grandes volumes de dados passaram a apoiar decisões que afetam crédito, emprego, tratamento médico, concessão de benefícios e outras situações do cotidiano.

Quando esses sistemas funcionam dentro do esperado, trazem rapidez, redução de custos e novas oportunidades econômicas. Quando falham, porém, podem gerar prejuízos materiais, danos à imagem, discriminação e violação de direitos fundamentais, colocando em evidência a responsabilidade civil pelos erros cometidos por essas tecnologias.

O direito civil brasileiro foi estruturado com foco em condutas humanas, analisando culpa, dolo, nexo causal e dano a partir de ações ou omissões de pessoas físicas e jurídicas. A inteligência artificial, por sua vez, opera por meio de modelos estatísticos complexos, com algum grau de autonomia técnica e resultados que nem sempre são totalmente previsíveis ou transparentes para usuários e até mesmo para desenvolvedores.

Surge, então, o dilema central deste trabalho: quando um sistema de IA gera uma decisão injusta ou equivocada, quem deve responder perante a vítima? O programador que concebeu o modelo? A empresa que integrou o sistema à sua atividade econômica? O usuário que acionou a ferramenta em determinada situação concreta? Em certas discussões mais recentes, ainda aparece a hipótese teórica de responsabilizar a própria máquina, como se pudesse ocupar posição similar à de sujeito de direito.

Esse debate torna-se particularmente relevante em áreas de risco elevado, como saúde, transporte, finanças e segurança, nas quais decisões automatizadas podem impactar vidas, liberdade e patrimônio. O tema proposto responsabilidade civil por erros cometidos pela inteligência artificial: um dilema entre criador, usuário e a própria máquina busca examinar como o ordenamento jurídico brasileiro pode responder a essas novas

configurações tecnológicas sem abandonar princípios consolidados de proteção à pessoa e de reparação de danos.

A pergunta norteadora que orienta a pesquisa é: em que medida os institutos clássicos da responsabilidade civil, aliados à legislação já existente e aos projetos normativos em tramitação, são capazes de enfrentar adequadamente os danos decorrentes de erros de sistemas de inteligência artificial, distribuindo o risco entre desenvolvedores, fornecedores, usuários e demais agentes envolvidos?

O objetivo geral consiste em analisar, sob a ótica do direito civil brasileiro, os fundamentos e caminhos possíveis para a responsabilização de erros praticados por sistemas de inteligência artificial. Como objetivos específicos, pretende-se: a) apresentar um panorama da evolução da IA e de seu uso em atividades que envolvem risco jurídico relevante; b) examinar os critérios tradicionais de responsabilidade subjetiva e objetiva e sua aplicação em ambientes automatizados; c) discutir propostas doutrinárias e legislativas que tratam da responsabilidade pelos danos decorrentes de sistemas de IA; d) refletir sobre modelos que repartem o risco entre criador, usuário e demais participantes da cadeia tecnológica, afastando a ideia de atribuir responsabilidade direta à máquina.

A justificativa do estudo decorre da rápida expansão da inteligência artificial em serviços públicos e privados, muitas vezes sem que operadores do direito compreendam plenamente o funcionamento dos algoritmos ou os limites de sua utilização. A falta de compreensão pode levar à naturalização de erros técnicos como se fossem “fatalidades” ou à transferência indevida dos custos do dano para a vítima. Investigar a responsabilidade civil nesse contexto contribui para fortalecer a proteção de direitos, incentivar práticas de governança e transparência por parte de empresas e orientar a atuação de profissionais do direito diante de litígios envolvendo decisões automatizadas.

A relevância acadêmica e social da pesquisa está ligada à necessidade de atualizar a leitura dos institutos da responsabilidade civil frente à nova fase da digitalização das relações sociais. Discutir quem deve responder por erros cometidos por sistemas de inteligência artificial auxilia na construção de parâmetros mais claros para

juizadores, advogados, reguladores e desenvolvedores de tecnologia, favorecendo um ambiente em que inovação e tutela de direitos caminhem juntos, com repartição mais equilibrada dos riscos inerentes à automação de decisões.

Metodologicamente, trata-se de pesquisa bibliográfica, construída a partir da análise de livros, artigos científicos, teses, dissertações e documentos normativos que tratam de responsabilidade civil e inteligência artificial, com foco em produções recentes. A abordagem é qualitativa, voltada à interpretação de conceitos, argumentos e propostas regulatórias presentes na doutrina e nos textos legais, buscando articular diferentes contribuições em um quadro teórico que auxilie a compreensão do dilema entre criador, usuário e a própria máquina.

A primeira seção busca demonstrar que o conceito de inteligência artificial, explicar o seu funcionamento e sua origem. Diante disso, a inteligência artificial já é uma realidade consolidada no cotidiano social, presente em decisões que vão desde a concessão de crédito bancário até o diagnóstico de doenças complexas.

A segunda seção é a mais densa juridicamente e identifica três fragilidades estruturais do regime atual. A primeira delas é a ausência de um marco legal específico: o Brasil ainda não conta com uma lei geral sobre IA em vigor. O PL 2.338/2023, aprovado pelo Senado, segue em tramitação na Câmara dos Deputados, e enquanto isso juízes e advogados precisam adaptar o Código Civil de 2002, o CDC, a LGPD e o Marco Civil da Internet para situações que esses diplomas jamais previram diretamente.

A segunda fragilidade diz respeito ao desafio do nexo causal. Em atividades tradicionais, é possível identificar com clareza quem praticou o ato danoso. Em sistemas de IA, o resultado parece emergir de uma cadeia complexa e distribuída, de aparência impessoal, o que dificulta apontar o responsável. A doutrina, porém, rejeita que essa complexidade técnica sirva como escudo, sustentando que por trás de todo algoritmo há pessoas e organizações que tomaram decisões sobre dados, modelos e usos. A terceira fragilidade é ilustrada por casos concretos de uso irregular da IA.

A terceira seção aponta os instrumentos jurídicos já disponíveis no ordenamento brasileiro que podem ser aplicados aos danos causados pela IA. O artigo 927, parágrafo único, do Código Civil pode ser interpretado para incluir o uso de IA em larga escala como atividade de risco, sujeitando o agente à responsabilidade objetiva. O

artigo 14 do CDC protege os consumidores lesados por decisões automatizadas sem exigir que provem o defeito específico do algoritmo — basta demonstrar o dano e a falha do serviço. A LGPD, por sua vez, já exige revisão humana de decisões automatizadas que afetem direitos.

No campo legislativo, o PL 2.338/2023 representa o avanço mais estruturado. O projeto propõe que o fornecedor ou operador de sistema de IA que cause danos seja obrigado a repará-los integralmente, independentemente do grau de autonomia do sistema. Para sistemas de alto risco, a responsabilidade é objetiva; para os demais, há presunção de culpa com inversão do ônus da prova em favor da vítima. O texto ainda define excludentes específicas, como a comprovação de que o agente não colocou o sistema em circulação ou que o dano decorreu de fato exclusivo da vítima, de terceiro ou de caso fortuito externo.

Por fim, a seção aponta para a necessidade de uma governança algorítmica mais ampla, que combine documentação dos modelos, testes de impacto, auditorias independentes e canais de contestação de decisões automatizadas. O objetivo é criar incentivos econômicos para que empresas invistam em prevenção, e não apenas lidem com o dano depois que ele já ocorreu. A mensagem central é que inovação tecnológica e tutela de direitos precisam caminhar juntas.

1. A ERA DA INTELIGÊNCIA ARTIFICIAL E O DILEMA JURÍDICO

1.1 CONCEITO DA INTELIGÊNCIA ARTIFICIAL

A representação de máquinas que agem por conta própria, antigamente, apenas retratada em livros e filmes, vem cada vez mais tornando uma realidade. O cenário comumente representado de forma distópica, a título de exemplo filmes como: *Matrix* e *Blade Runner 2049*, pode gerar um medo constante do futuro, na medida que o desenvolvimento exponencial da Inteligência Artificial parece não apresentar limites. O cenário ulterior deverá revolucionar a forma como os indivíduos realizam as suas tarefas cotidianas, a forma que interagem no trabalho e na sociedade, gerando, assim, em novos fatos jurídicos.

Dessa forma, entender o conceito da Inteligência Artificial é indubitável antes de adentrar na temática da responsabilidade civil, chegar a um conceito definitivo é quase impossível, uma vez que apresenta diferentes entendimentos. Para John McCarthy, célebre cientista da computação, considerado o indivíduo que cunhou o termo *artificial intelligence* pela primeira vez, em 1956. Para ele a definição estaria relacionada à uma ciência e engenharia de criar máquinas inteligentes, estando associada com a tarefa de usar computadores para entender a inteligência humana, como segue a citação de McCarthy (2007, p. 01):

É a ciência e engenharia de criar máquinas inteligentes, especialmente programas de computador inteligentes. Está relacionado com a tarefa similar de usar computadores para entender a inteligência humana, mas a IA não precisa se restringir a métodos que são biologicamente observáveis.

Outra definição amplamente conhecida no meio da computação é a retratada no livro “Artificial intelligence: a modern approach”, no qual os autores Russel e Norvig, dividem a definição de inteligência artificial levando em consideração duas ideias fundamentais: a capacidade de aprendizagem e a manifestação de “comportamento inteligente”. Ademais, os autores afirmam que é possível dividir as definições de inteligência artificial em quatro categoria (Russel; Norvig, 2010.p. 02), conforme segue a tabela abaixo:

Thinking Humanly “The exciting new effort to make computers think . . . <i>machines with minds</i> , in the full and literal sense.” (Haugeland, 1985) “[The automation of] activities that we associate with human thinking, activities such as decision-making, problem solving, learning . . .” (Bellman, 1978)	Thinking Rationally “The study of mental faculties through the use of computational models.” (Charniak and McDermott, 1985) “The study of the computations that make it possible to perceive, reason, and act.” (Winston, 1992)
Acting Humanly “The art of creating machines that perform functions that require intelligence when performed by people.” (Kurzweil, 1990) “The study of how to make computers do things at which, at the moment, people are better.” (Rich and Knight, 1991)	Acting Rationally “Computational Intelligence is the study of the design of intelligent agents.” (Poole <i>et al.</i> , 1998) “AI . . . is concerned with intelligent behavior in artifacts.” (Nilsson, 1998)
Pensando Humanamente “O novo e excitante esforço para fazer os computadores pensarem... <i>máquinas com mentes</i> , no sentido pleno e literal.” (Haugeland, 1985) “[A automatização de] atividades que associamos ao pensamento humano, atividades como tomada de decisão, resolução de problemas, aprendizagem...” (Bellman, 1978)	Pensando Racionalmente “O estudo das faculdades mentais através do uso de modelos computacionais.” (Charniak e McDermott, 1985) “O estudo das computações que tornam possível perceber, raciocinar e agir.” (Winston, 1992)
Agindo Humanamente “A arte de criar máquinas que executam funções que exigem inteligência quando executadas por pessoas.” (Kurzweil, 1990) “O estudo de como fazer os computadores fazerem coisas nas quais, no momento, as pessoas são melhores.” (Rich e Knight, 1991)	Agindo Racionalmente “A Inteligência Computacional é o estudo do projeto de agentes inteligentes.” (Poole <i>et al.</i> , 1998) “IA... diz respeito ao comportamento inteligente em artefatos.” (Nilsson, 1998)

Diante do quadro supra descrito, as quatro categorias são divididas em sistemas que agem como seres humanos, sistemas que pensam como seres humanos, sistemas que pensam racionalmente e sistemas que agem racionalmente.

O sistema que agem como seres humanos, são caracterizados por aqueles sistemas que apresentam comportamento similar ao dos seres humanos, mas tudo isso depende do processamento de linguagem natural, argumentação automatizada, e aprendizado de máquina que permite adaptar-se a novas circunstâncias. No mesmo sentido, tem os sistemas que pensam como os seres humanos, sendo aqueles que tentam simular a capacidade de pensar dos seres humanos.

Dentro dessa perspectiva, concatena-se diretamente ao teste proposto por Alan Turing, conhecido como Teste de Turing, definido no artigo *Computing Machinery and Intelligence*, publicado em outubro de 1950. O teste consistia em torno de uma interação de 3 (três) participantes um humano examinador, um humano que serve como

"interlocutor" e uma máquina. A interação ocorre por meio de um sistema de mensagens, de forma que o examinador não tenha conhecimento da identidade dos participantes. O objetivo da máquina é enganar o avaliador, fazendo acreditar que praticamente qualquer um dos campos questionados é um humano respondendo.

Foi de grande relevância o estudo por provocar a discussão em torno da questão do que significaria ser um humano e se a inteligência poderia ser reproduzida artificialmente. O teste foi um dos principais marcos da computação, por analisar de forma criteriosa, e por muito tempo pendurou resultados em que facilmente seria possível distinguir uma resposta gerada por um humano em detrimento de uma máquina. No entanto, restou demonstrado que o Teste de Turing já se encontra ultrapassado, conforme noticiado na CNN BRASIL (2025), como apresentado:

Um estudo publicado em março de 2025 relatou que o ChatGPT, IA da empresa OpenAI, passou no experimento pela primeira vez ao ser testado com três participantes — a versão mais desafiadoras do Teste de Turing.

O GPT-4.5 conseguiu enganar as pessoas, fazendo-as pensar que se tratava de outro ser humano em 73% das vezes. Além dele, outras quatro linguagens de inteligência artificial conseguiram passar no experimento com taxas mais baixas de aprovação.

Alguns estudos que dialogam com o direito e a proteção de dados destacam que a IA deve ser entendida também como processo de tratamento intensivo de informações pessoais e sensíveis. Isso significa reconhecer que o “conceito” da tecnologia envolve não só algoritmos, mas também práticas de coleta, armazenamento e uso de dados, com potencial de gerar discriminações, exclusões e decisões opacas quando não há controles adequados de transparência e revisão humana.

Quando se desloca a análise para o campo da responsabilidade civil, a inteligência artificial passa a ser vista como elemento de uma atividade que gera risco para terceiros, especialmente quando empregada em contextos de alta complexidade ou em larga escala.

Nessa leitura, importa menos discutir se a máquina é “inteligente” em sentido filosófico e mais compreender que ela compõe uma engrenagem sociotécnica, dentro da qual se avaliam deveres de cuidado, previsibilidade de danos e formas de repartir o risco

entre desenvolvedores, fornecedores e usuários que se valem dessa tecnologia em proveito próprio.

1.2 ORIGEM E FUNCIONAMENTO

O estudo da Inteligência Artificial (IA), enquanto campo de pesquisa formal, iniciou-se por volta da década de 1950, mas a busca por dispositivos e métodos capazes de simular o raciocínio humano remonta a mais de dois mil anos. Desde a Antiguidade, filósofos têm explorado questões sobre a natureza da mente e da existência, onde já procuravam entender como são realizados os processos de visão, lembrança, aprendizagem e raciocínio.

Embora, o início do estudo da Inteligência Artificial, como campo de pesquisa formal tenha se dado por volta da década de 1950. Suas bases foram lançadas ainda na década de 1940, quando o biógrafo Warren McCulloch e o matemático Walter Pitts foram pioneiros ao criar o primeiro modelo computacional de redes neurais, combinando matemática e algoritmos para retratar o funcionamento do cérebro humano, funcionando como se fosse um neurônio artificial. Tal raciocínio foi apresentado no artigo “*A logical calculus of the ideas immanent in nervous activity*”, em 1943.

No caso em tela, convém mencionar que o uso da IA em setores do judiciário tem-se tornado foco de projetos, conforme o Painel da Pesquisa sobre Inteligência Artificial, publicado pelo CNJ (Conselho Nacional de Justiça) (CADIP, 2025, p. 10):

Entre 2022 e 2024, verificou-se o aumento de 26% no número de projetos de IA em desenvolvimento nos tribunais brasileiros, conforme levantamento promovido pelo CNJ. Esses projetos, parte do programa Justiça 4.0, incluíram a criação de plataformas em nuvem para integrar sistemas judiciais e compartilhar soluções tecnológicas entre os tribunais. Porém, dentre os 140 projetos existentes, apenas 37 encontram-se hospedados no Sinapses, plataforma do CNJ para impulsionar a IA no Judiciário.

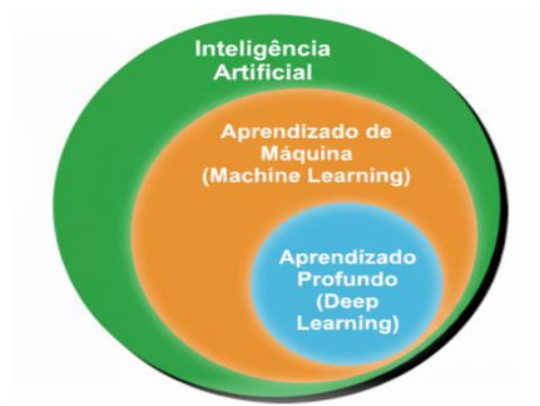
O dado surpreende, na medida em que houve um aumento considerável no número de projetos. O reflexo da IA para modernização do Judiciário, já se demonstra uma realidade, com o objetivo de alcançar uma acessibilidade, celeridade e eficiência na análise e desenvolvimento de processos.

No entanto, não é possível deixar de lado sobre o funcionamento da Inteligência Artificial e para isso, faz-se necessário comentar sobre os algoritmos. Para os educadores Fabricio Ferrari e Cristian Cechinel os algoritmos consistem (Ferrari, Cechinel, 2008, p. 14):

Como uma sequência finita de passos (instruções) para resolver um determinado problema. Sempre que desenvolvemos um algoritmo estamos estabelecendo um padrão de comportamento que deverá ser seguido (uma norma de execução de ações) para alcançar o resultado de um problema.

Nesse enfoque, a qualidade do algoritmo relaciona com a qualidade dos dados que serão, por ele, processados- em síntese, a qualidade da decisão automatizada está diretamente relacionada à qualidade dos dados que o algoritmo processa. O problema surge se o algoritmo se baseia em dados históricos repletos de preconceitos, ele reproduzirá, de forma automatizada, os mesmos padrões preconceituosos utilizados como base de seu processamento, conseqüentemente, ocasionando em situações injustas.

Ademais, o funcionamento da Inteligência Artificial divide em dois subconjuntos conhecidos como Machine Learning (aprendizado da máquina) e Deep Learning (aprendizado profundo), como representado no quadro ilustrativo a seguir:



(Figura: Representação da relação entre inteligência artificial, *Machine Learning* e *Deep Learning*. Adaptado de KODERA et al, 2021, p. 321)

No processo de Machine Learning, serve, basicamente como uma ferramenta que através dos algoritmos ensina o dispositivo, como um computador, a detectar

padrões e se adaptar aos problemas que são apresentados, a máquina aprende com a própria experiência e, por consequência, os dados inseridos no sistema são processados para gerar um novo algoritmo.

Já o Deep Learning, se apresenta como uma evolução ainda maior, mesmo sendo um subcampo do Machine Learning, são modelos baseados em redes neurais, que funcionam de forma semelhante ao funcionamento do cérebro humano. Em regra, o próprio dispositivo consegue realizar previsões de possíveis problemas, de modo que ela mesma encontra um meio de solucioná-los a partir da utilização de algoritmos, nesse caso atua com um grau elevado de autonomia. Como o computador obtém conhecimento da experiência, não há necessidade de um operador humano especificar formalmente todo o conhecimento que o computador precisa.

Por fim, tais fenômenos passaram a ser parte de todos os setores da vida individual e social, existindo inclusive uma grande dependência com seu uso, mesmo que o indivíduo desconheça do funcionamento da IA. A lista é extensa, mas alguns exemplos, nos mais diversos setores. No ramo do transporte, já existem uma grande quantidade de carros autônomos. No direito são múltiplas as suas aplicações, tanto no Poder Judiciário, como o surgimento da plataforma SINAPSE, na advocacia e nas demais atividades relacionadas ao Direito, como na área contratual.

1.3 O CENÁRIO DA MODERNIZAÇÃO E O DILEMA DA IA

Na atualidade, é impossível descartar a utilização da tecnológica nas mais diversas áreas, praticamente tudo que antes demandava trabalho manual, hoje pode ser significativamente facilitado com utilização da tecnologia. A partir disso, surgiram várias tecnologias com o fito de aumentar a eficiência, automatizar e tornar mais rápida a conclusão das mais diversas tarefas de nosso cotidiano. Nesse contexto a IA pode ser usada de maneira estratégica, alinhada de forma correta ao objetivo.

Um dos principais argumentos para seu uso consiste no aumento da precisão e na redução de custos. No entanto, observa-se que a divisão de trabalho entre o ser humano e máquinas- o dilema entre inteligência artificial e inteligência natural- sofre constante mudanças, surgindo um cenário, inclusive, mais favorável a inteligência

artificial.

Prova disso foi o relatório Inteligência Artificial e Empregos, divulgado pela Organização para a Cooperação e Desenvolvimento Econômico (OCDE), aponta que 27% dos empregos dos países da OCDE são de profissões com alto risco de automatização pela inteligência artificial. A organização argumenta sobre a necessidade da coleta de dados sobre o uso da tecnologia, para facilitar os futuros cenários prováveis.

No campo das relações de consumo, a inteligência artificial passou a atuar em análise de crédito, definição de limites, organização de filas virtuais e atendimento automatizado, afetando diretamente o acesso de milhões de pessoas a serviços financeiros e digitais. Estudos nessa área destacam que o fornecedor concentra o domínio da tecnologia e dos dados, enquanto o consumidor recebe apenas o resultado final, muitas vezes negativo, com pouca explicação, o que agrava a vulnerabilidade já reconhecida pelo direito do consumidor

Na saúde, hospitais e clínicas vêm incorporando sistemas de apoio diagnóstico, ferramentas de triagem e modelos de priorização de atendimento, sempre sob a marca da modernização e da inovação tecnológica. Pesquisas em bioética apontam que esses recursos podem ajudar equipes sobrecarregadas, mas também introduzem novos focos de erro quando são implementados sem transparência adequada, sem treinamento consistente das equipes e sem supervisão clínica real, o que dificulta a identificação de quem descumpriu deveres de cuidado quando um paciente sofre dano.

Especialistas no assunto, ressaltam sobre o perigo desse crescimento da IA. Para o pesquisador canadense Geoffrey Hinton, ex funcionário do Google, também conhecido como "padrinho" da IA, em uma entrevista para BBC evidencia que a inteligência que está sendo desenvolvida é diferente da que temos, conseqüentemente, ocasionando em um cenário repleto de incertezas (2023):

Neste momento, o que estamos vendo são coisas como o GPT-4 superar uma pessoa na quantidade de conhecimento geral que ela tem, e a supera de longe. Em termos de raciocínio, não é tão bom, mas já é capaz de raciocínios simples. Cheguei à conclusão de que o tipo de inteligência que estamos desenvolvendo é muito diferente da inteligência que temos.

Diante dos avanços, o jurista Carlos Affonso Pereira de Souza, na 1ª edição do livro *Inteligência Artificial e Direito: ética, regulação e responsabilidade* (2019) elucida a necessidade acerca da modernização jurisdicional em diversas áreas. Dentro desse contexto, Souza (2019, p. 100) expõe:

Acerca da responsabilidade jurídica em eventuais danos, para o Direito brasileiro serão exigidas soluções e atualizações de ordem administrativa, penal e cível com o objetivo de integrar carros e tecnologias autônomas. Em termos administrativos, a responsabilização exposta no Código de Trânsito Brasileiro ainda está atrelada à identificação de um motorista humano. No campo da responsabilidade penal, o sistema brasileiro é fundado na tipicidade e em elementos pessoais do causador do dano, como dolo ou culpa, o que traz patentes dificuldades em casos de falhas de funcionamento da IA. E, mesmo para o ramo da responsabilidade civil, em que a responsabilização objetiva encontra previsões casuísticas, discussões acerca do nexo de causalidade e de excludentes de ilicitude poderão afetar investimentos de empresas e reparações a vítimas.

O objeto de assunto já se mostrou de grande preocupação no cenário internacional, mais especificamente na União Europeia com a 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)), para eles o atual quadro jurídico não seria suficiente para abranger os danos provocados pela nova geração de robôs, na medida em que os robôs podem ser dotados de capacidades adaptativas e de aprendizagem que integram um certo grau de imprevisibilidade no seu comportamento. Logo, faz-se necessário um maior aprofundamento do ordenamento jurídico brasileiro nessa temática.

2. CRÍTICA AO REGIME ATUAL E A BUSCA POR NOVOS PARADIGMAS DE RESPONSABILIDADE

Segundo Silva (2023, p. 02), a entrada da inteligência artificial em serviços públicos e privados obrigou o direito civil a lidar com situações em que decisões com forte impacto na vida das pessoas são produzidas com apoio de modelos estatísticos complexos, treinados sobre grandes volumes de dados. O regime vigente, construído para lidar com condutas humanas individualizáveis, passa a ser testado em conflitos que envolvem código, bases de dados, escolhas empresariais e camadas de automação pouco transparentes para quem sofre o dano, o que tem motivado leituras críticas na doutrina.

Nesse contexto, explica Marques (2024, p. 60), ao mesmo tempo em que se exige reparação adequada para quem é prejudicado, mantém-se a preocupação de não travar totalmente a inovação tecnológica, o que coloca o debate em um ponto de tensão permanente. A questão deixa de ser apenas se há responsabilidade, passando a envolver quem deve responder, com qual intensidade e sob qual fundamento normativo, em um contexto em que o próprio desenho dos sistemas influencia a distribuição dos riscos na sociedade.

Nesse quadro, ganha corpo a ideia de que o arranjo atual é fragmentado, dependente de interpretação criativa de normas espalhadas em diferentes diplomas, o que torna o tratamento dos casos pouco previsível. A busca por novos paradigmas de responsabilidade procura justamente consolidar critérios mais claros para atividades baseadas em inteligência artificial, alinhando institutos tradicionais a exigências de governança algorítmica, transparência e prevenção de danos.

Por fim, é indubitável elucidar sobre a fragilidade no âmbito jurídico relacionada à problemática da IA. Dentro disso, a Advogada e também professora de direito civil da UNB, Ana Frazão (2025), elucida:

É fácil observar que, diante de uma realidade preocupante, o CNJ não se debruçou sobre vários dos riscos já mapeados nem se dedicou a, partindo da análise dos principais modelos de IA atualmente existentes, indicar providências concretas para a necessária adaptação destes à realidade do Judiciário. **Optou-se por uma abordagem excessivamente principiológica e procedimental, o que já seria complicado por si só, ainda mais quando esta se associa a uma**

classificação de riscos duvidosa, que considera como baixo risco praticamente todas as atividades para a elaboração de uma decisão judicial. O excesso de otimismo do CNJ também se evidencia com a expectativa de que soluções contratuais possam resolver problemas estruturais e – o que é mais grave – com a possibilidade de que, na ausência de uma solução institucional – o que seria imprescindível – magistrados individualmente decidam como utilizar a tecnologia. (Grifou-se)

2.1 O MARCO LEGAL DA INTELIGÊNCIA ARTIFICIAL NO CENÁRIO NACIONAL

No Brasil, ainda não há uma lei geral de inteligência artificial em vigor, mas o debate legislativo avançou com propostas que procuram estabelecer princípios, deveres e categorias próprias para quem desenvolve, fornece e utiliza sistemas automatizados, sendo a principal o PJ 2.338/23. Para Cardoso (2024, p. 150), estudos sobre regulação da IA destacam que projetos em tramitação buscam articular essa disciplina específica com o Código Civil, o Código de Defesa do Consumidor, a Lei Geral de Proteção de Dados e o Marco Civil da Internet, sem criar um universo normativo isolado do restante do ordenamento.

De acordo com autor Almeida (2023, p. 12), uma linha comum nesses textos legislativos é a preocupação em identificar atores da cadeia tecnológica, como desenvolvedor, fornecedor, operador e beneficiário, com o intuito de facilitar a atribuição de responsabilidade. A ideia é afastar a impressão de que o dano surge de uma entidade abstrata chamada IA reforçando que, por trás do sistema, há pessoas e organizações que tomam decisões sobre coleta de dados, parametrização de modelos e inserção da ferramenta em serviços de alto impacto social.

Ademais, convém mencionar que o uso da Inteligência Artificial em julgamentos não é de consenso geral. Conforme membro do centro de direito da sociedade da informação da universidade de Milão, Vinicius Almada Mozetic (2025, p. 1):

Ninguém deseja ser “julgado por um robô”, mas tampouco se pode ignorar que humanos também carregam vieses. A sinergia entre auditorias técnicas e supervisão judicial surge como estratégia para checar e balancear ambos os riscos, criando um ecossistema de governança responsável que impeça que preconceitos humanos ou falhas tecnológicas comprometam a imparcialidade das decisões. O debate sobre kill switches é igualmente central. A possibilidade de desligar ou suspender um sistema de IA em casos críticos, sem esvaziar garantias constitucionais, coloca em pauta a separação de poderes e os limites do controle estatal ou privado de plataformas. A jurisprudência do Supremo Tribunal Federal (STF), em casos de bloqueio de aplicativos e redes sociais, exemplifica a relevância da cautela para que medidas emergenciais não se

convertam em abuso ou censura. Em síntese, a inteligência artificial é uma realidade irreversível no Judiciário, mas não pode nem deve ser encarada como substituto de sua natureza humanística.

Como adverte o autor Sousa (2024, p. 140), o debate parlamentar também tem incorporado categorias ligadas a níveis de risco, distinguindo sistemas empregados em atividades com baixa exposição de direitos daqueles voltados a decisões que podem afetar saúde, liberdade ou subsistência econômica. A classificação por risco funciona como ponto de partida para graduar deveres de governança, exigindo mais de quem explora sistemas com potencial de dano grave ou irreversível, o que abre caminho para modelos de responsabilidade mais rígidos nesses contextos.

Ao lado dessas iniciativas, Nascimento (2022, p. 76) lembra que a doutrina do país já conta com bases normativas que incidem sobre o uso de IA, mesmo sem menção expressa à tecnologia. Regras sobre defeito do produto, responsabilidade pelo serviço e tratamento de dados pessoais criam um conjunto mínimo de deveres que se projetam automaticamente sobre quem adota soluções algorítmicas em suas atividades, ainda que o marco legal específico de inteligência artificial não esteja concluído.

Nesse cenário, Marques (2024, p.62) crítica o desenho atual do marco em construção aponta o risco de se produzir uma lei declaratória, com princípios amplos, mas pouca densidade operacional para casos concretos. Para evitar isso, parte da literatura sugere que a futura disciplina da IA incorpore parâmetros objetivos de diligência, como documentação de modelos, registros de testes, avaliação prévia de impacto e canais de contestação de decisões automatizadas, vinculando o descumprimento desses deveres à responsabilidade civil por danos decorrentes da atividade.

2.2 ANÁLISE LEGISLATIVA E JURISPRUDENCIAL

Como sintetiza Bessa (2023, p. 22), a análise do regime vigente mostra que grande parte dos conflitos envolvendo decisões automatizadas ainda é tratada com base em normas de responsabilidade por defeito do produto e falha na prestação do serviço. Em relações de consumo, tribunais tendem a enquadrar danos causados por sistemas de IA como vícios de segurança ou erros na execução de serviços, o que leva à aplicação da responsabilidade objetiva do fornecedor, independentemente de prova de culpa,

sempre que o resultado se afasta da legítima expectativa criada para o consumidor, sendo usado como alicerce o Código de Defesa ao Consumidor e o Marco Civil da Internet.

Nesse contexto, conforme decisão proferida pelo Tribunal de Justiça do Estado de Minas Gerais:

EMENTA: AGRAVO DE INSTRUMENTO - AÇÃO DE TUTELA ANTECIPADA EM CARÁTER ANTECEDENTE - PERDA DE OBJETO - NÃO OCORRÊNCIA - WHATSAPP E FACEBOOK - GRUPO ECONÔMICO - BANIMENTO - BLOQUEIO DE CONTA EM APLICATIVO - COMUNICAÇÃO PRÉVIA FUNDAMENTADA - NECESSIDADE - DIREITO DO CONSUMIDOR - ALTERNATIVIDADE DA CLÁUSULA RESOLUTIVA DO CDC - DEVER DE INFORMAÇÃO - CONTRADITÓRIO - AMPLA DEFESA - LIBERDADE DE EXPRESSÃO E COMUNICAÇÃO - AMBIENTE VIRTUAL - TROCA DE DADOS - LEI 12.965/2014 - MARCO CIVIL DA INTERNET - EFICÁCIA IMEDIATA E HORIZONTAL DOS DIREITOS HUMANOS - MÁQUINAS - ALGORITMOS - INTELIGÊNCIA ARTIFICIAL - RESPONSABILIDADE DA EMPRESA PELAS DECISÕES E CONSEQUÊNCIAS - TUTELA ANTECIPADA DE URGÊNCIA - REQUISITOS PREENCHIDOS. **Sendo frágeis as provas apresentadas em face da abrangência do objeto da demanda atinente à disponibilidade de conta em aplicativo, não ocorre perda de objeto. Constatados a probabilidade do direito e o perigo de dano (art. 300 do CPC), não se mostra razoável o banimento de conta em aplicativo, sendo assegurado ao usuário o direito de apresentar defesa à empresa que faz parte do mesmo grupo financeiro e tem ingerência sobre o aplicativo, sem representação institucional no Brasil, de modo a evitar, em sede de tutela antecipada, prejuízo ao consumidor.** Uma vez que o uso de recursos digitais de comunicação e de compartilhamento de dados tornou-se imprescindível no ambiente social, profissional e político, torna-se necessário imprimir coerência às relações jurídicas privadas modernas por meio da adoção de valores constitucionais, em conformidade com a eficácia imediata e horizontal dos direitos humanos, de modo a impedir que empresas de tecnologia, por meio de algoritmos, máquinas e inteligência artificial, violem princípios da Constituição da República e normas nacionais, em especial as dispostas no Marco Civil da Internet e no Código de Defesa do Consumidor. (grifou-se)

(TJ-MG - AI: XXXXX05976319001 MG, Relator: Marcos Henrique Caldeira Brant, Data de Julgamento: 23/06/2021, Câmaras Cíveis / 16ª CÂMARA CÍVEL, Data de Publicação: 24/06/2021)

Diante do supracitado, o Código de Defesa do Consumidor evidencia no seu artigo 14 a responsabilidade do fornecedor perante esses casos, no que segue a transcrição do artigo: O fornecedor de serviços responde, independentemente da existência de culpa, pela reparação dos danos causados aos consumidores por defeitos relativos à prestação dos serviços, bem como por informações insuficientes ou inadequadas sobre sua fruição e riscos.

No campo da responsabilidade civil geral, a jurisprudência vem utilizando o artigo 927 do Código Civil, sobretudo o parágrafo único, para fundamentar decisões em que atividades intensivamente tecnológicas são consideradas de risco. Como adverte Silva (2023, p. 3), a interpretação de que o uso de sistemas automatizados em larga escala, especialmente em setores como crédito e saúde, gera riscos superiores aos encontrados em atividades comuns tem sido empregada para justificar regimes mais rigorosos de imputação de danos, ainda que os julgados raramente usem expressões técnicas relacionadas à inteligência artificial.

A Lei Geral de Proteção de Dados (Lei nº 13.709/2018), também começa a aparecer como fundamento em casos que discutem decisões automatizadas e perfis construídos a partir de grandes bases de dados. Em harmonia com o apresentado por Schuck, (2025, p. 9), é necessário a exigência de transparência, revisão por pessoa natural e adoção de medidas de segurança passa a ser interpretada não apenas como garantia de privacidade, mas como parâmetro de diligência mínima, capaz de reforçar a responsabilidade de controladores que deixam de gerenciar adequadamente o ciclo de vida dos dados usados em sistemas de IA.

Em consonância Chaves Filho (2024, p. 3), apresentou que na esfera da saúde, resoluções profissionais e normas infralegais ainda ocupam espaço relevante, disciplinando o uso de tecnologias de apoio diagnóstico e de prontuário eletrônico. Essas regras, combinadas com princípios de bioética, servem como referência para verificar se houve descuido na avaliação de resultados gerados por sistemas automatizados, o que repercute na repartição da responsabilidade entre profissionais e instituições em casos de erro clínico associado à inteligência artificial.

Apesar desses avanços, a jurisprudência nacional ainda não consolidou padrões específicos para conflitos que envolvem IA, o que gera decisões oscilantes e baixa previsibilidade. Relatos de casos analisados em pesquisas acadêmicas mostram que, muitas vezes, o elemento tecnológico aparece apenas de forma lateral na fundamentação, sem discussão aprofundada sobre opacidade de modelos, origem dos dados ou razão técnica da falha, o que limita o potencial transformador da responsabilidade civil nesse campo.

2.3 PROBLEMAS DO USO IRREGULAR DA IA

Como apresentado por Schuck (2025, p. 9), o uso irregular de sistemas de inteligência artificial surge quando ferramentas são implementadas sem análise adequada de risco, sem compatibilização com normas de proteção de dados, consumo e responsabilidade civil ou sem preparo de equipes para interpretar resultados. Em contextos assim, decisões automatizadas podem negar crédito, rebaixar perfis de pacientes ou até mesmo proferir decisões de forma equivocada, sem que existam mecanismos claros de contestação ou revisão humana que impeçam a cristalização de injustiças.

Nesse contexto, a título de exemplo, no dia 11 de fevereiro de 2026, ficou marcado com o voto do desembargador Magid Nauef Láuar, este que foi relator do caso que havia absolvido um homem de 35 anos por estupro de vulnerável, no estado de Minas Gerais. Não obstante, com análise detida do acórdão, na página 45, há a frase "Agora melhore a exposição e fundamentação deste parágrafo". Desse modo, observa-se que foi usado a Inteligência Artificial em um caso considerado extremamente sensível.

De acordo com Chaves Filho (2024, p. 4), na saúde, o emprego de algoritmos de apoio diagnóstico sem validação local ou sem adaptação para a realidade epidemiológica de determinada região pode comprometer a segurança de pacientes. Estudos em bioética descrevem situações em que modelos treinados em populações específicas foram utilizados de forma generalizada, gerando risco de erro sistemático em grupos pouco representados na base original, o que fere o dever de cuidado de instituições que optam por soluções tecnológicas sem avaliação crítica robusta.

Além disso Bessa (2023, p. 25), opina que em serviços financeiros e de consumo, a utilização de modelos de scoragem ou de triagem automatizada sem transparência adequada favorece a reprodução de discriminações e a exclusão silenciosa de clientes considerados de alto risco por critérios não explicitados. Pesquisas nessa área apontam que, quando a empresa se refugia na explicação genérica de que "o sistema não aprovou", esvazia-se o direito de informação do consumidor e dificulta-se a produção de prova, o que agrava a posição de vulnerabilidade já reconhecida em lei.

Já existem julgados que exploram a fragilidade nos casos em que a Inteligência Artificial é utilizada com má-fé, referente a isso o Tribunal de Justiça de São Paulo, concluiu:

AGRAVO DE INSTRUMENTO. RECUPERAÇÃO JUDICIAL. Recurso interposto contra decisão que determinou a substituição da gestora judicial. Posterior convolação em falência. Perda superveniente do objeto recursal. Litigância de má-fé. Uso indevido de inteligência artificial, com citação de jurisprudência inexistente. Conduta que não se trata de mero equívoco ou desatenção técnica. Dever do profissional de conferência da veracidade das informações lançadas nas peças processuais. Conduta capaz de induzir o juízo a erro. DECISÃO MANTIDA. RECURSO PREJUDICADO, COM APLICAÇÃO DE MULTA.

(TJ-SP - Agravo de Instrumento: 2206069-59.2025.8.26.0000 Taubaté, Relator: AZUMA NISHI, Data de Julgamento: 18/12/2025, 1ª Câmara Reservada de Direito Empresarial, Data de Publicação: 18/12/2025)

Existe um problema recorrente no uso de IA fora do escopo para o qual a ferramenta foi concebida. Isso ocorre, por exemplo, quando um modelo treinado para apoio a decisões internas passa a ser empregado como critério exclusivo para conceder ou negar serviços a terceiros, sem revisão humana, ou quando algoritmos experimentais são integrados a processos decisórios sem testes suficientes. A literatura de responsabilidade civil considera que essas escolhas de implementação configuram violações de deveres de cuidado por parte de quem decide explorar economicamente a tecnologia.

Somado a isso para Cardoso (2024, p. 155), a ausência de documentação adequada sobre o desenvolvimento, o treinamento e a atualização de sistemas torna difícil reconstruir a cadeia de decisões técnicas que levou a determinado resultado danoso. Pesquisas que examinam a articulação entre regulação e responsabilidade civil sugerem que essa falta de registros, longe de ser mero detalhe administrativo, constitui fator de agravamento de responsabilidade, pois impede auditorias independentes e compromete a capacidade do próprio Judiciário de avaliar onde ocorreu a falha.

Conforme pesquisa realizada pelo escritório Tozzini Freire Advogados, uma das organizações de serviços jurídicos líderes na América Latina. O Brasil pode ser um dos melhores países para reivindicar indenização por falhas no uso de inteligência artificial, desde 2010, quando o termo foi mencionado pela primeira vez, a pesquisa constatou que 140 ações judiciais ajuizadas buscaram indenização por falhas de IA,

inclusive em reconhecimento facial e remoção de produtos nas plataformas digitais. Além disso, constatou que 64% das reivindicações foram mantidas em apelação. Outro dado importante apresentado foram os estados com o maior número de casos envolvendo IA, o qual São Paulo (61 casos) e Rio Grande do Sul (49 casos).

Ademais, Sofia Kilmar, sócia da área de contencioso da TozziniFreira, elucidou que a maior parte dos casos envolvia prevenção de fraudes financeiras, constou que antigamente a IA era mencionada apenas em casos relacionados à organização do conteúdo da internet, mas a partir de 2016 os temas começaram a se diversificar.

Segundo Kilmar, houve 50 casos que questionaram os processos de busca e os algoritmos utilizados pelas plataformas tecnológicas, afirmou também que os tribunais não fazem distinção entre algoritmo e a IA em processos judiciais, sendo uma estratégia intencional para ampliar a análise do nível de risco representado pela tecnologia, a título de exemplo são tratados da mesma maneira mesmo sendo diferentes: Reconhecimento facial, deepfakes, algoritmos, IA generativa e formas de automação.

2.4 O DESAFIO DA IMPUTAÇÃO E O ROMPIMENTO DO NEXO CAUSAL

Um dos pontos mais delicados na discussão sobre inteligência artificial e responsabilidade civil é o nexo causal entre a conduta imputada e o dano verificado. Em atividades tradicionais, costuma ser possível reconhecer com relativa clareza quem praticou a ação que desencadeou o prejuízo. Já em sistemas que combinam código, dados, decisões de desenho e uso em cadeia, surge a impressão de que o resultado decorre de uma sequência complexa e distribuída, muitas vezes apresentada como fruto de estatística impessoal, o que dificulta apontar o vínculo entre o comportamento de um agente específico e o dano.

Conforme Cestonaro (2023, p. 3), na saúde, essa dificuldade se manifesta de maneira muito concreta em casos de diagnóstico equivocado em que o profissional seguiu a sugestão de um sistema de apoio decisório. A pergunta central passa a ser em que medida o médico poderia ou deveria ter desconfiado do resultado, levando em conta a cultura institucional que incentiva a confiança na ferramenta, a apresentação da saída em linguagem técnica persuasiva e o próprio treinamento oferecido a quem utiliza a tecnologia.

Do ponto de vista estrutural, o emprego de modelos de aprendizado de máquina, cuja lógica interna nem sempre é explicável em linguagem acessível, alimenta a sensação de que a causalidade se dilui dentro da matemática, como se não houvesse ato humano identificável. Para Terranova (2024, p. 2), a doutrina que trata de responsabilidade médica e profissional em contexto de IA chama atenção para o risco de aceitar esse argumento como inevitável, pois isso poderia levar a um esvaziamento prático da tutela civil em nome da complexidade técnica.

Na seara regulatória, Cardoso (2024, p. 160) elucida que alguns autores defendem que, em vez de tentar reconstituir cada passo algorítmico, o direito deve concentrar a análise de causalidade em decisões humanas de nível mais alto, como a escolha de adotar determinado sistema, a forma de integrá-lo em fluxos de trabalho, o grau de supervisão mantido e a qualidade dos testes prévios. Nessa perspectiva, o rompimento do nexo causal seria afastado quando se verifica que o dano guarda conexão razoável com riscos inerentes à atividade escolhida pelo agente, ainda que o caminho técnico específico seja difícil de detalhar.

Existe, ainda, um artigo que trata sobre as circunstâncias que os agentes de inteligência artificial não serão responsabilizados. Conforme o artigo 28, do projeto de lei 2338/2023:

Art. 28. Os agentes de inteligência artificial não serão responsabilizados quando:

I – comprovarem que não colocaram em circulação, empregaram ou tiraram proveito do sistema de inteligência artificial; ou

II – comprovarem que o dano é decorrente de fato exclusivo da vítima ou de terceiro, assim como de caso fortuito externo.

A literatura que discute novos paradigmas aponta que regimes de responsabilidade objetiva e presunções relativas de culpa podem servir como resposta à alegada opacidade da inteligência artificial. Em conclusão ao apresentado Sousa levanta um questionamento (2024, p. 155), na medida que em vez de aceitar a quebra do nexo causal com base na complexidade técnica, essa abordagem transfere para desenvolvedores e operadores o encargo de demonstrar que adotaram práticas de governança suficientes para reduzir riscos, invertendo a lógica de prova e alinhando os incentivos econômicos à prevenção de danos em sistemas algorítmicos.

3 ADAPTAÇÃO DE LEIS PARA UM NOVO TEMPO

Conforme Silva (2023, p. 2), a inteligência artificial intensificou uma tensão que já existia entre códigos pensados para uma sociedade analógica e conflitos surgidos em ambientes digitais. O Código Civil, o Código de Defesa do Consumidor, a Lei Geral de Proteção de Dados e o Marco Civil da Internet foram escritos em momentos diferentes, com objetivos próprios, mas hoje são chamados a responder, em conjunto, por danos que envolvem algoritmos, dados massivos e decisões automatizadas. A doutrina lembra que a função da responsabilidade civil é justamente absorver mudanças sociais, ajustando interpretações para que a proteção a direitos da personalidade, ao patrimônio e à dignidade acompanhe a forma como as relações se organizam no tempo.

Uma linha importante de adaptação vem da releitura do Código Civil e do CDC à luz de atividades baseadas em inteligência artificial. O artigo 927, parágrafo único, do Código Civil, por exemplo, passou a ser discutido em conexão com empresas que utilizam sistemas automatizados em larga escala, sugerindo que o tratamento de dados em grande volume e a delegação de decisões a modelos opacos podem caracterizar atividade de risco, sujeita a responsabilidade objetiva. Ao mesmo tempo, as regras do CDC sobre defeito do produto e do serviço são revisitadas para alcançar falhas de ferramentas digitais que não entregam segurança compatível com a confiança despertada no usuário, mesmo quando o erro nasce de decisões algorítmicas invisíveis para o público.

É indubitável trazer em pauta, o capítulo V do Projeto de Lei 2338/2023, que trata da responsabilidade civil ligada aos sistemas de inteligência artificial, como segue:

Art. 27. O fornecedor ou operador de sistema de inteligência artificial que cause dano patrimonial, moral, individual ou coletivo é obrigado a repará-lo integralmente, independentemente do grau de autonomia do sistema.

§ 1º Quando se tratar de sistema de inteligência artificial de alto risco ou de risco excessivo, o fornecedor ou operador responde objetivamente pelos danos causados, na medida de sua participação no dano.

§ 2º Quando não se tratar de sistema de inteligência artificial de alto risco, a culpa do agente causador do dano será presumida, aplicando-se a inversão do ônus da prova em favor da vítima.

Art. 28. Os agentes de inteligência artificial não serão responsabilizados quando:

I – comprovarem que não colocaram em circulação, empregaram ou tiraram proveito do sistema de inteligência artificial; ou

II – comprovarem que o dano é decorrente de fato exclusivo da vítima ou de terceiro, assim como de caso fortuito externo.

Art. 29. As hipóteses de responsabilização civil decorrentes de danos causados por sistemas de inteligência artificial no âmbito das relações de consumo permanecem sujeitas às regras previstas na Lei nº 8.078, de 11 de setembro de 1990 (Código de Defesa do Consumidor), sem prejuízo da aplicação das demais normas desta Lei.

Para Cardoso (2024, p.10), o debate legislativo específico sobre inteligência artificial tende a reforçar esse movimento de atualização, sem substituir o núcleo dos códigos existentes. Projetos em tramitação, a exemplo do 2338/23, discutem categorias como “agente de IA”, classificação de sistemas por risco e deveres de governança, com o objetivo de criar uma moldura mais clara para a atuação de desenvolvedores, fornecedores e operadores. A literatura que acompanha esse processo ressalta que a nova disciplina deve dialogar com o que já está consolidado, servindo como camada adicional de proteção, e não como regime paralelo que fragilize conquistas anteriores em matéria de responsabilidade civil e de defesa do consumidor.

De outro modo, para Bessa (2023, p. 20), no campo do consumo, a adaptação das leis passa por reconhecer que muitos produtos e serviços hoje são essencialmente digitais, com funcionamento mediado por inteligência artificial. A análise de crédito em tempo real, a precificação dinâmica em plataformas, a triagem de reclamações por chatbots e as recomendações personalizadas de conteúdo deixam de ser detalhes tecnológicos e entram no radar dos deveres de informação, segurança e lealdade previstos no CDC. Estudos recentes apontam que, em casos de falha na decisão automatizada, o consumidor não deve ser compelido a decifrar o algoritmo para ter direito à reparação, pois a responsabilidade recai sobre quem organiza a atividade econômica e controla a tecnologia utilizada.

Reforçando a discussão, de acordo com Schuck (2025, p.9), a Lei Geral de Proteção de Dados também é reinterpretada em chave de adaptação a esse novo quadro. A previsão de revisão de decisões automatizadas por pessoa natural e o direito à explicação mínima sobre critérios utilizados ganham peso quando sistemas de IA começam a influenciar o acesso a crédito, emprego e serviços públicos. Autores que

estudam vieses algorítmicos sustentam que a LGPD funciona, nesse ponto, como ponte entre proteção de dados e responsabilidade civil, pois descumprimentos reiterados de deveres de transparência e segurança indicam falhas de diligência que justificam indenizações por danos materiais e morais ligados a decisões automatizadas injustas.

Em harmonia com Cestonaro (2023, p.3), na saúde, a adaptação tanto de normas legais quanto de regulamentos profissionais gira em torno da incorporação de sistemas de apoio diagnóstico e ferramentas de triagem baseadas em inteligência artificial. Em vez de enxergar essas tecnologias como substitutas do julgamento clínico, reformas regulatórias buscam reforçar o papel do profissional como responsável pela decisão final, exigindo validação local dos modelos, registro das versões utilizadas e protocolos claros de supervisão humana. Análises recentes ressaltam que, quando essas exigências não são observadas, instituições e profissionais assumem riscos adicionais, que devem ser considerados na apuração de culpa ou na aplicação de responsabilidade objetiva em contextos de dano a pacientes.

Outro eixo de adaptação ocorre na própria interpretação judicial. Decisões que enfrentam conflitos ligados a plataformas digitais e bancos de dados começam a reconhecer que a complexidade técnica não pode servir como escudo para afastar a responsabilidade de quem ganha com a exploração econômica da tecnologia. Desse modo, Marques (2024) analisou que em julgados comentados pela doutrina, tribunais aplicam regras tradicionais de inversão do ônus da prova, obrigação de guarda de registros e solidariedade entre integrantes da cadeia de fornecimento, ajustando o raciocínio jurídico para alcançar situações em que a prova do caminho algorítmico seria inacessível à parte lesada.

O tema da adaptação de leis também passa pela discussão sobre novos deveres de conduta para agentes que operam com inteligência artificial. Propostas que circulam em relatórios técnicos e estudos acadêmicos defendem a posituação de obrigações relativas à documentação de modelos, testes de impacto, auditorias independentes, canais de contestação e revisão de decisões automatizadas, associando o descumprimento desses deveres a presunções de culpa ou a regimes de responsabilidade mais rígidos. A ideia é aproximar responsabilidade civil e governança,

criando incentivos econômicos para que empresas invistam em controle e prevenção, e não apenas lidem com o dano quando ele já ocorreu.

Adaptações legislativas e interpretativas apontam para a necessidade de combinar diferentes instrumentos, e não apostar em uma norma única capaz de resolver todos os problemas ligados à inteligência artificial. O fortalecimento de regimes objetivos em atividades de alto risco, a articulação entre CDC, LGPD e Código Civil, a criação de um marco específico para IA e a valorização de práticas de governança algorítmica formam um conjunto em construção. O resultado esperado é um quadro no qual inovação tecnológica e tutela de direitos caminhem juntos, com uma distribuição mais justa do risco entre quem se beneficia economicamente da automação e quem fica exposto às falhas dos sistemas.

3.1 REFLEXO NO PODER JUDICIÁRIO

Como explica Marques (2024, p. 60), o Poder Judiciário brasileiro passou a lidar com demandas em que a inteligência artificial aparece como personagem central do conflito, seja em decisões de crédito, em bloqueios de contas digitais, em diagnósticos apoiados por software ou em filtragem de conteúdos em plataformas. Magistrados e servidores convivem com petições que mencionam algoritmos, aprendizado de máquina e decisões automatizadas, o que exige revisão de rotinas de prova, de valoração de laudos técnicos e de aplicação de categorias clássicas de responsabilidade civil a estruturas técnicas complexas que não foram previstas quando o Código Civil foi elaborado.

Na prática, muitos litígios chegam ao Judiciário com uma narrativa em que a parte lesada afirma ter sido prejudicada por uma decisão “do sistema”, enquanto a empresa ou instituição tenta deslocar a discussão para a esfera técnica, alegando falha pontual, erro do usuário ou mera aplicação de parâmetros neutros.

Para Silva (2023, p. 3) a experiência recente mostra que, se o juiz aceitar a ideia de que o algoritmo é uma espécie de caixa fechada, o nexos causal tende a se dissolver e a tutela reparatória enfraquece, motivo pelo qual decisões mais recentes

passaram a exigir que fornecedores expliquem, ao menos em linhas gerais, como se deu o processamento de dados em cada caso contestado.

Para terminar, Bessa (2023, p. 20) explica que Nas relações de consumo, essa discussão aparece de forma intensa em ações sobre scoragem de crédito, bloqueio de cadastros, negativa de serviços digitais e fraudes em ambientes automatizados. Julgados comentados pela doutrina indicam que tribunais vêm aplicando a responsabilidade objetiva do fornecedor quando o dano decorre de falha do serviço mediado por inteligência artificial, sem exigir que o consumidor prove o defeito específico do algoritmo, já que o controle técnico e organizacional da solução permanece nas mãos da empresa, que auferir lucro com sua utilização e deve suportar o risco correspondente.

3.2 A RESPONSABILIDADE CIVIL POR ATOS INDEPENDENTES DA INTELIGÊNCIA ARTIFICIAL

A expressão “atos independentes da inteligência artificial” costuma surgir em debates que descrevem o sistema como se ele agisse sozinho, quase como sujeito de vontade própria, afastado de escolhas humanas. A doutrina crítica alerta que essa narrativa mascara uma cadeia de decisões de projeto, de seleção de dados e de implementação que permanece inteiramente humana, além de interessar a atores econômicos que preferem diluir a responsabilidade em uma entidade tecnológica abstrata, em vez de assumir o risco de suas estratégias de digitalização.

Sob a ótica de Cardoso (2024, p. 18) na responsabilidade civil, erros atribuídos à inteligência artificial seguem vinculados a pessoas físicas e jurídicas que organizam a atividade, investem na solução, escolhem modelos, definem parâmetros de uso e colhem os frutos econômicos da automação. A análise desloca o foco do comportamento interno do algoritmo para decisões de mais alto nível, como a opção por utilizar determinada ferramenta em contexto sensível, a ausência de testes adequados ou a falta de supervisão humana em etapas em que o próprio fornecedor reconhece que há margem de falha, o que aproxima essas condutas do conceito de risco criado ou de defeito de serviço.

De acordo com Sousa (2024, p. 113), em setores de alto impacto, como saúde, transporte e crédito, ganha força a ideia de que o uso de inteligência artificial configura atividade que, por sua natureza, gera risco relevante para terceiros, o que justifica regimes de responsabilidade mais rígidos, inclusive objetivos, ainda que o dano decorra de combinações de dados e parâmetros difíceis de reconstruir caso a caso. Modelos em discussão na doutrina sugerem que, nesses contextos, quem decide explorar economicamente a IA responde pelos resultados típicos da atividade, cabendo direito de regresso interno contra desenvolvedores ou outros parceiros técnicos se ficar demonstrado que houve falha de execução contratual no desenho do sistema.

3.3 MEDIDAS ÉTICAS E REGULATÓRIAS PARA O CONTROLE E A SEGURANÇA DA IA

A incorporação de inteligência artificial em serviços públicos e privados traz consigo um conjunto de problemas éticos ligados à autonomia das pessoas, à justiça distributiva e à proteção da intimidade, que não podem ser reduzidos a questão técnica. Na saúde, por exemplo, comitês de ética e conselhos profissionais passaram a discutir o lugar dos algoritmos em decisões clínicas, reafirmando que a relação entre profissional e paciente não pode ser substituída por um cálculo automatizado, ainda que a ferramenta traga ganhos de agilidade em ambientes sobrecarregados e com recursos limitados

Outrossim, para Schuck (2025, p. 9), no campo regulatório, cresce a defesa de mecanismos de governança algorítmica que incluam documentação detalhada dos modelos, registros de versões utilizadas, rastreamento de decisões automatizadas e realização periódica de testes de impacto, sobretudo em organizações que usam inteligência artificial para conceder ou negar serviços essenciais. A literatura especializada mostra que essas medidas ampliam a possibilidade de auditoria e ajudam a identificar vieses discriminatórios ou padrões de erro antes que gerem danos em larga escala, funcionando também como parâmetro para verificação de culpa ou de descumprimento de deveres de cuidado em futuras ações de responsabilidade civil.

Especificamente na saúde, comenta Chaves Filho (2024, p. 4) que autores de bioética e de responsabilidade profissional vêm apontando caminhos que combinam

protocolos clínicos, validação científica de modelos e regras claras de supervisão humana. A recomendação recorrente é que hospitais e clínicas não tratem sistemas de apoio ao diagnóstico como oráculos infalíveis, mas como ferramentas auxiliares, integradas a rotinas em que o médico mantém a palavra final sobre a conduta, com obrigação de registrar quando decide seguir ou afastar a sugestão do algoritmo, o que facilita a reconstrução da cadeia decisória em caso de dano ao paciente.

3.4 REQUISITOS NECESSÁRIOS PARA ANÁLISE DA RESPONSABILIDADE CIVIL EM CASOS DE ERROS

Primeiramente, Silva (2023, p. 4) explica que em processos que envolvem erros associados à inteligência artificial, a primeira tarefa é identificar com clareza quem são os agentes relevantes na cadeia de decisão. Isso inclui desenvolvedores que criaram o modelo, empresas que comercializaram a solução, organizações que integraram a ferramenta em seus fluxos de trabalho e profissionais que, na ponta, acionaram o sistema na relação concreta com o usuário ou paciente. Sem essa reconstrução mínima da engrenagem humana que sustenta o algoritmo, o risco é tratar a inteligência artificial como ente isolado e esvaziar a própria lógica da responsabilidade civil, que se apoia na ideia de dever jurídico de cuidado assumido por pessoas e entidades determinadas.

Outro requisito central é a avaliação do dever de prevenção em cada etapa da cadeia, que envolve perguntas sobre previsibilidade do dano, possibilidade de mitigação e custo de medidas adicionais de segurança. Como elucidado por Nascimento (2022, p. 76), a doutrina que estuda regulação e responsabilidade em ambientes tecnológicos sugere que, em muitos casos, o dano ligado à IA não decorre de uma fatalidade imprevisível, mas da escolha de não testar adequadamente o sistema em condições reais, de não treinar equipes ou de não adaptar o modelo para contextos específicos, condutas que se enquadram em falhas de diligência analisáveis com as ferramentas já conhecidas do direito civil.

Ainda segundo Nascimento (2022, p. 3) também se mostra indispensável uma análise cuidadosa da natureza do dano, que nem sempre se limita a perdas materiais imediatas. A inteligência artificial pode produzir prejuízos ligados à reputação digital, à

exclusão silenciosa em processos seletivos, à negação injusta de crédito ou à classificação injusta de perfis de risco em saúde, efeitos que se acumulam no tempo e atingem direitos da personalidade e oportunidades futuras. estudos recentes em bioética e proteção de dados apontam que esses danos difusos exigem atenção especial do julgador, inclusive na fixação de indenização por dano moral e na definição de obrigações de fazer ligadas à correção de registros e perfis automatizados.

Na visão de Terranova (2024, p. 2), em alguns campos, como o da responsabilidade médica, a análise da responsabilidade por erros de inteligência artificial passa ainda por uma articulação fina entre perícia técnica e parâmetros jurídicos. Modelos voltados a diagnóstico ou estratificação de risco em saúde, por exemplo, exigem que o perito explique ao juiz se a ferramenta foi utilizada dentro das indicações aprovadas, se havia validação suficiente para aquele contexto e se o profissional teve acesso a alertas sobre limitações do sistema, elementos que influenciam a conclusão sobre culpa do profissional e da instituição, bem como sobre eventual necessidade de responsabilização solidária por parte de quem desenvolveu ou forneceu o software.

CONCLUSÃO

A pesquisa mostrou que a inteligência artificial não é um ente com plena autonomia pairando sobre a vida social, e sim uma tecnologia construída e operada por pessoas e organizações que tomam decisões sobre dados, modelos, usos e limites. Quando surgem erros em diagnósticos, análises de crédito, serviços digitais ou outras aplicações, o dano não nasce da “máquina em si”, mas de um conjunto de escolhas humanas distribuídas ao longo da cadeia técnica e econômica. A responsabilidade civil, nesse contexto, continua girando em torno de deveres de cuidado assumidos por desenvolvedores, fornecedores, instituições e profissionais que optam por incorporar a IA em atividades de risco, ainda que os caminhos internos do algoritmo sejam complexos e pouco transparentes.

O estudo também evidenciou que o ordenamento jurídico brasileiro já oferece ferramentas úteis para enfrentar esses conflitos, a partir da articulação entre Código Civil, Código de Defesa do Consumidor, Lei Geral de Proteção de Dados e normas setoriais. A discussão sobre um marco legal específico para a inteligência artificial tende a reforçar esse núcleo, desde que não crie brechas para diluir o dever de indenizar sob o pretexto de inovação. Regimes de responsabilidade objetiva em atividades de alto risco, deveres de governança algorítmica, documentação de modelos, testes de impacto e canais de contestação de decisões automatizadas surgem como caminhos para reequilibrar a relação entre quem lucra com a automação e quem suporta os prejuízos quando algo falha.

A análise aponta que o desafio não é apenas técnico ou normativo, mas também cultural e institucional. Tribunais, órgãos reguladores, empresas, profissionais de saúde e de outras áreas precisam aprender a dialogar com a linguagem da inteligência artificial sem naturalizar a opacidade como desculpa para a ausência de responsabilização. A construção de um ambiente em que IA e direitos caminhem juntos depende de escolhas políticas e jurídicas claras: identificar quem cria e explora a tecnologia, exigir prevenção de danos como parte do modelo de negócio e garantir que

vítimas tenham condições reais de ver reconhecida e reparada a lesão sofrida, inclusive quando ela se materializa por meio de decisões automatizadas.

REFERÊNCIAS

ALMEIDA, Guilherme Chabrour de. *Desafios da inteligência artificial em matéria de responsabilidade civil*. Anais do Congresso Internacional de Direitos Humanos de Coimbra, Coimbra, 2023. Disponível em: <https://www.cidp.pt/publicacoes/anais-direitos-humanos/desafios-inteligencia-artificial-responsabilidade-civil>. Acesso em: 8 dez. 2025.

AMADO, Tatiane Costa. *Bioética e inovações tecnológicas na saúde*. Revista Contemporânea (online), 2024. Disponível em: <https://ojs.revistacontemporanea.com/ojs/index.php/home/article/view/6358>. Acesso em: 8 dez. 2025.

BBC NEWS BRASIL. *O 'padrinho' da inteligência artificial que se demitiu do Google e adverte sobre perigos da tecnologia*. BBC News Brasil, 2 maio 2023. Disponível em: <https://www.bbc.com/portuguese/articles/cgr1qr06myzo>. Acesso em: 01 out. 2025.

BESSA, Leonardo Roscoe; NUNES, Ana Luisa Tarter. *Inteligência artificial (IA) e danos ao consumidor*. Revista de Direito do Consumidor, São Paulo, v. 32, n. 150, p. 15-48, nov./dez. 2023. Disponível em: <https://www.thomsonreuters.com.br/content/dam/ewp-m/documents/brazil/pt/pdf/other/rdc150.pdf>. Acesso em: 8 dez. 2025.

BITTAR, Eduardo Carlos Bianca. *Metodologia da pesquisa jurídica: teoria e prática da monografia para os cursos de direito*. 18. ed. São Paulo: Saraiva, 2024.

BRASIL. Senado Federal. Projeto de Lei nº 2338, de 2023. Dispõe sobre o uso da Inteligência Artificial. Brasília, DF: Senado Federal, 2023. Disponível em: <https://legis.senado.leg.br/sdleg-getter/documento?disposition=inline&dm=9347622>. Acesso em: 11 ago. 2025.

CARDOSO, Henrique Ribeiro; MELO, Bricio Luís da Anunciação. *Regulação e responsabilidade civil por danos causados por sistemas de inteligência artificial*. Revista de Direito Administrativo, Rio de Janeiro, v. 283, n. 3, 2024. Disponível em:

<https://bibliotecadigital.fgv.br/ojs/index.php/rda/article/view/91564>. Acesso em: 8 dez. 2025.

CENTRO DE APOIO AO DIREITO PÚBLICO (CADIP). *Inteligência Artificial no Poder Judiciário*. 2. ed., rev. e atual. São Paulo, 2025.

CESTONARO, Chiara; TODARO, Michele; SANTORO, Antonio; FERRARA, Santo Davide; OLIVA, Antonio. *Defining medical liability when artificial intelligence is applied on diagnostic algorithms: a scoping review*. *Frontiers in Medicine*, Lausanne, v. 10, 2023. Disponível em: <https://www.frontiersin.org/articles/10.3389/fmed.2023.1129874/full>. Acesso em: 8 dez. 2025.

CHAVES FILHO, Cláudio do Carmo; BYK, Jonas; FERREIRA, Luiz Carlos de Lima. *Desafios bioéticos do uso da inteligência artificial na oftalmologia*. *Revista Bioética*, Brasília, v. 32, e3727PT, 2024. Disponível em: https://revistabioetica.cfm.org.br/revista_bioetica/article/view/3727. Acesso em: 8 dez. 2025.

CNN BRASIL. *Teste de Turing: conheça experimento que identifica se IA superou humanos*. CNN Brasil, São Paulo, 2 dez. 2023. Disponível em: <https://www.cnnbrasil.com.br/tecnologia/teste-de-turing-conheca-experimento-que-identifica-se-ia-superou-humanos/>. Acesso em: 1 nov. 2025.

DONEDA, Danilo Cesar Maganhoto et al. *Considerações iniciais sobre inteligência artificial, ética e autonomia pessoal*. *Pensar*, Fortaleza, v. 23, n. 4, p. 1-17, out./dez. 2018. Disponível em: <https://www.techtudo.com.br/dicas-e-tutoriais/2020/11/onde-ficam-os-downloads-no-android-como-acessar-arquivos-salvos-no-celular.ghtml>. Acesso em: 7 dez. 2025.

EBAC. *Deep learning vs machine learning: o que é, qual é a diferença*. Blog da EBAC, 14 jun. 2024. Disponível em: <https://ebaonline.com.br/blog/machine-learning-e-deep-learning-seo>. Acesso em: 15 jan. 2026.

ELDAKAK, Ahmad; ALREMEITHI, Ahmed; DAHIYAT, Emran A. *Civil liability for the actions of autonomous AI in healthcare: an invitation to further contemplation*. *Humanities and Social Sciences Communications*, London, v. 11, n. 1, 2024. Disponível em: <https://www.nature.com/articles/s41599-024-01817-9>. Acesso em: 8 dez. 2025.

FERRARI, Fabricio; CECHINEL, Cristian. *Introdução a Algoritmos e Programação*. Bagé, RS: [s.n.], 2008. Versão 2.0.

FRAZÃO, Ana. *Inteligência artificial no Poder Judiciário: por que devemos continuar a nos preocupar com o assunto mesmo após a nova resolução do CNJ sobre o tema?*. JOTA, 26 fev. 2025. Disponível em: <https://www.jota.info/opiniao-e-analise/colunas/constituicao-empresa-e-mercado/inteligencia-artificial-no-poder-judiciario-2>.

FRAZÃO, Ana; MULHOLLAND, Caitlin (coords.). *Inteligência artificial e direito: ética, regulação e responsabilidade*. São Paulo: Thomson Reuters Brasil, 2019. p. 100 E-book.

INFOMONEY.OCDE diz que Inteligência Artificial pode substituir 27% dos empregos dos países do grupo. InfoMoney, São Paulo, 11 jul. 2023. Disponível em: <https://www.infomoney.com.br/carreira/ocde-diz-que-inteligencia-artificial-pode-substituir-27-dos-empregos-dos-paises-do-grupo/>. Acesso em: 1 out. 2025.

KODERA, S. et al. *Prospects for cardiovascular medicine using artificial intelligence*. *J Cardiol.*, [S.l.], v. 79, n. 3, p. 319-25, 2022. Disponível em: <https://www.sciencedirect.com/science/article/pii/S0914508721002835>. Acesso em: 10 nov. 2025.

LIANG, Lu-Hai. *Brazilians are winning court claims for AI failures, including on facial recognition*. *Biometric Update*, 29 abr. 2025. Disponível em: <https://www.biometricupdate.com/202504/brazilians-are-winning-court-claims-for-ai-failures-including-on-facial-recognition>. Acesso em: 15 jan. 2026

LIMA, Daniel Cândido. *Metodologia da pesquisa aplicada ao direito: ensinando em primeira pessoa*. Independently published, 2025.

MARQUES, Guilherme Raso. *Responsabilidade civil na era da inteligência artificial*. Revista da Advocacia Pública Federal, Brasília, v. 8, n. 1, p. 58-80, 2024. Disponível em: <https://seer.anafe.org.br/index.php/revista/article/view/202>. Acesso em: 8 dez. 2025.

MCCULLOCH, Warren S.; PITTS, Walter. *A logical calculus of the ideas immanent in nervous activity*. *Bulletin of Mathematical Biophysics*, [S.l.], v. 5, n. 4, p. 115-133, 1943. Disponível em: <https://www.cs.cmu.edu/~epxing/Class/10715/reading/McCulloch.and.Pitts.pdf>. Acesso em: 10 out. 2025.

MEDON, Filipe. *Inteligência Artificial e Responsabilidade Civil: Autonomia, Riscos e Solidariedade*. 2. ed. revista, atualizada e ampliada. Salvador: JusPODIVM, 2022.

MINAS GERAIS (Estado). Tribunal de Justiça (10ª Câmara Cível). Apelação Cível nº 1.0000.22.123456-7/001. Apelante: João da Silva. Apelado: Banco XYZ S/A. Relator: Desembargador Fulano de Tal. Jusbrasil, [S. l.], 20 out. 2025. Disponível em: <https://www.jusbrasil.com.br/jurisprudencia/tj-mg/1237538239>. Acesso em: 14 abr. 2026.

MONARD, Maria Carolina; BARANAUKAS, José Augusto. *Aplicações de Inteligência Artificial: Uma Visão Geral*. São Carlos: Instituto de Ciências Matemáticas e de Computação de São Carlos, 2000. p. 2.

MOZETIC, Vinicius Almada. *IA no Judiciário brasileiro: salvaguardas, riscos e novas fronteiras*. JOTA, 2 mar. 2025. Disponível em: <https://www.jota.info/opiniao-e-analise/artigos/ia-no-judiciario-brasileiro-salvaguardas-riscos-e-novas-fronteiras>. Acesso em: 11 jan. 2026.

MULHOLLAND, Caitlin. *O conceito de sujeito não-humano e a responsabilidade civil da Inteligência Artificial*. In: FRAZÃO, Ana; MULHOLLAND, Caitlin (coords.). *Inteligência*

artificial e direito: ética, regulação e responsabilidade. São Paulo: Thomson Reuters Brasil, 2019. p. 331-356.

NASCIMENTO, Hérica Cristina Paes; SOUZA, Maique Barbosa de; OLIVEIRA, Patrícia da Silveira. *A regulação da inteligência artificial e novos contornos para caracterização da responsabilidade civil*. *Revista de Direito, Governança e Novas Tecnologias*, Florianópolis, v. 7, n. 2, p. 73-90, 2022. Disponível em: <https://indexlaw.org/index.php/revistadgnt/article/view/8371>. Acesso em: 8 dez. 2025.

NASCIMENTO, Sabrina Maciel; PAIVA, Tábita Maika Gama de; KASUGA, Marcos Paulo Maciel; SILVA, Thiago de Assis Furtado; CROZARA, Conceição Maria Guedes; BYK, Jonas; FURTADO, Silvania da Conceição. *Inteligência artificial e suas implicações éticas e legais: revisão integrativa*. *Revista Bioética*, Brasília, v. 32, e3729PT, 2024. Disponível em: https://revistabioetica.cfm.org.br/revista_bioetica/article/view/3729. Acesso em: 8 dez. 2025.

NUNES, Heloá da Conceição; GUIMARÃES, Rita Miranda Coessens; DADALTO, Luciana. *Desafios bioéticos do uso da inteligência artificial em hospitais*. *Revista Bioética*, Brasília, v. 30, n. 1, p. 82-93, 2022. Disponível em: https://revistabioetica.cfm.org.br/revista_bioetica/article/view/2644. Acesso em: 8 dez. 2025.

OLIVEIRA, Maria Luísa Pinho Pereira de. *A responsabilidade civil e a inteligência artificial no Brasil: necessidade de atualização legislativa*. *Revista Brasileira de Direito Civil*, Belo Horizonte, v. 28, n. 1, p. 45-62, jan./mar. 2020.

RUSSELL, Stuart J.; NORVIG, Peter. *Artificial Intelligence: a modern approach*. 3ª ed. New Jersey: Pearson Education, 2010. p. 2.

SÃO PAULO (Estado). Tribunal de Justiça. Intimação. Agravo de Instrumento nº 2361481-80.2025.8.26.0000. Jusbrasil, [S. l.], 18 dez. 2025. Disponível em: <https://www.jusbrasil.com.br/diarios/documentos/5424927682/intimacao-agravo-de->

[instrumento-2361481-8020258260000-disponibilizado-em-18-12-2025-tjsp](#). Acesso em: 14 abr. 2026.

SCHUCK, Sabrina de Lima; ROSAL, Isabela Maria. *Insuficiência da LGPD na mitigação de vieses discriminatórios: desafios e responsabilidade civil dos sistemas de inteligência artificial*. Revista de Direito e as Novas Tecnologias, São Paulo, n. 27, 2025. Disponível em: <https://dspace.almg.gov.br/handle/11037/62822>. Acesso em: 8 dez. 2025.

SILVA, Francisco Alves da. *Responsabilidade civil e inteligência artificial: explorando soluções e desafios da era digital*. RECIMA21 – Revista Científica Multidisciplinar, São Paulo, v. 4, n. 11, e4114434, 2023. Disponível em: <https://recima21.com.br/index.php/recima21/article/view/2441>. Acesso em: 8 dez. 2025.

SOUSA, Andressa Evellyn de Oliveira. *A responsabilidade civil por danos causados por inteligência artificial: uma análise jurídica e ética*. Revista Brasileira de Direito Comercial, Empresarial, Concorrencial e do Consumidor, Brasília, v. 10, n. 57, 2024. Disponível em: <https://www.tjdft.jus.br/institucional/biblioteca/revistas-juridicas/revista-brasileira-de-direito-comercial-empresarial-concorrencial-e-do-consumidor/2024-v-10-n-57-fev-mar>. Acesso em: 8 dez. 2025.

STANFORD UNIVERSITY. John McCarthy. *What is Artificial Intelligence?* Stanford University, [s.d.]. Disponível em: <http://jmc.stanford.edu/articles/whatisai/whatisai>. Acesso em: 1 nov. 2025.

TEPEDINO, Gustavo; SILVA, Rodrigo da Guia. *Desafios da inteligência artificial em matéria de responsabilidade civil*. Revista Brasileira de Direito Civil - RBDCivil, Belo Horizonte, v. 21, p. 61-86, jul./set. 2019.

TERRANOVA, Claudio; BOUTONNET, Marion; FRACASSO, Tommaso; FERRARA, Santo Davide; OLIVA, Antonio. *AI and professional liability assessment in healthcare: a revolution in legal medicine?* Frontiers in Medicine, Lausanne, v. 10, 2024. Disponível em: <https://www.frontiersin.org/articles/10.3389/fmed.2024.1337335/full>. Acesso em: 8 dez. 2025.

TURING, Alan Mathison. Computing machinery and intelligence. *Mind*, Oxford, v. LIX, p. 460, out. 1950.

UNIÃO EUROPEIA. Parlamento Europeu. *Resolução do Parlamento Europeu, de 16 de fevereiro de 2017, com recomendações à Comissão sobre disposições de Direito Civil sobre Robótica (2015/2103(INL))*. Bruxelas: Jornal Oficial da União Europeia, 2017.

Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/PDF/?uri=CELEX:52017IP0048&from=EN>. Acesso em: 1 out. 2025.

ZENDESK. *Qual é a origem da Inteligência Artificial?* Zendesk Blog, [S.l.], [s.d.].

Disponível em: <https://www.zendesk.com.br/blog/qual-e-a-origem-da-inteligencia-artificial/>. Acesso em: 7 dez. 2025.