

RESPONSABILIDADE CIVIL POR DANOS CAUSADOS PELA INTELIGÊNCIA ARTIFICIAL: DESAFIOS E PERSPECTIVAS NO DIREITO BRASILEIRO

**CIVIL LIABILITY FOR DAMAGES CAUSED BY ARTIFICIAL INTELLIGENCE:
CHALLENGES AND PERSPECTIVES IN BRAZILIAN LAW**

**RESPONSABILIDAD CIVIL POR DAÑOS CAUSADOS POR INTELIGENCIA
ARTIFICIAL: DESAFÍOS Y PERSPECTIVAS EN EL DERECHO BRASILEÑO**

**Amanda Barros Batista Costa¹, Clara Kelliany Rodrigues de Brito², Danilo Souza Vale³, Glaucia
Fernanda Oliveira Martins Batalha⁴, Igor Carvalhal Frazão Corrêa⁵, Leandro Augusto dos
Remédios Costa⁶, Leticia Prazeres Falcão⁷, Maiane Cibele de Mesquita Serra⁸, Vitor Hugo Souza
Moraes⁹, Rossana Barros Pinheiro¹⁰**

DOI: 10.54899/dcs.v23i87.4401

Recibido: 30/01/2026 | Aceptado: 02/02/2026 | Publicación en línea: 09/02/2026.

RESUMO

A crescente integração de sistemas de Inteligência Artificial (IA) no cotidiano social e econômico impõe desafios sem precedentes ao ordenamento jurídico, especialmente no campo da responsabilidade civil. Caracterizada por sua autonomia e, em muitos casos, pela opacidade de seus processos decisórios (efeito “caixa-preta”), a IA desafia os pilares tradicionais da responsabilidade, como a culpa e o nexo de causalidade. O presente trabalho investiga o problema central de como atribuir a responsabilidade por danos gerados por sistemas autônomos no contexto brasileiro. Para isso, realiza-se uma análise crítica dos institutos clássicos do Código Civil e do Código de Defesa do Consumidor, avaliando sua aplicabilidade e suas limitações diante da complexidade algorítmica. Em seguida, aprofundam-se teorias contemporâneas, com ênfase

¹ Mestre em Ciências Sociais, Universidade Federal do Maranhão (UFMA), São Luís, Maranhão, Brasil.
E-mail: amandabarrosbatista@yahoo.com.br

² Doutora em Direito, Universidade de Marília (UNIMAR), São Paulo, São Paulo, Brasil.
E-mail: clarardebritoadv@gmail.com

³ Bacharel em Direito, Centro Universitário Santa Terezinha (CEST), São Luís, Maranhão, Brasil.
E-mail: danilo.vale98@icloud.com

⁴ Doutora em Ciências Sociais, Universidade Federal do Maranhão (UFMA), São Luís, Maranhão, Brasil.
E-mail: glauciamartinsbatalha@gmail.com

⁵ Especialista em Direito Internacional e Direitos Humanos, Pontifícia Universidade Católica de Minas Gerais (PUC - Minas), Belo Horizonte, Minas Gerais, Brasil. E-mail: igorcfc@hotmail.com

⁶ Doutor em Ciências Sociais, Universidade Federal do Maranhão (UFMA), São Luís, Maranhão, Brasil.
E-mail: leandroaugustocosta7@gmail.com

⁷ Mestre em Direito, Centro Universitário Christus (UNICHRISTUS), Fortaleza, Ceará, Brasil.
E-mail: leticiaprazeres@outlook.com.br

⁸ Doutora em Direito, Universidade Estadual Paulista (UNESP), Franca, São Paulo, Brasil.
E-mail: maianeserra@hotmail.com

⁹ Mestre em Direito e Instituições do Sistema de Justiça, Universidade Federal do Maranhão (UFMA), São Luís, Maranhão, Brasil. E-mail: vitorhugosmoraes@gmail.com

¹⁰ Mestre em Direito e Instituições do Sistema de Justiça, Universidade Federal do Maranhão (UFMA), São Luís, Maranhão, Brasil. E-mail: rossana.pinheiro@cest.edu.br

na teoria do risco da atividade (art. 927, parágrafo único, do CC), que emerge como o principal pilar para a responsabilização objetiva de sistemas de alto risco. A pesquisa possui natureza bibliográfica, dogmática e documental e conclui que a solução para a lacuna normativa atual tende à adoção da responsabilidade objetiva integral para sistemas de IA de alto risco, com fundamento na teoria do risco da atividade, buscando um equilíbrio entre a proteção das vítimas, a segurança jurídica e o incentivo à inovação tecnológica, com a indispensável garantia de transparência algorítmica.

Palavras-chave: Inteligência Artificial. Responsabilidade Civil. Dano. Nexos Causais. Risco da Atividade.

ABSTRACT

The growing integration of Artificial Intelligence (AI) systems into daily social and economic life poses unprecedented challenges to the legal system, particularly in the field of civil liability. Characterized by its autonomy and, in many cases, the opacity of its decision-making processes (the "black box" effect), AI challenges the traditional pillars of liability, such as fault and the causal link. This study investigates the central problem of how to attribute liability for damages caused by autonomous systems within the Brazilian context. To this end, a critical analysis of the classic institutes of the Civil Code and the Consumer Defense Code is conducted, evaluating their applicability and limitations in the face of algorithmic complexity. Subsequently, contemporary theories are deepened, with emphasis on the theory of the risk of activity (art. 927, sole paragraph, of the CC), which emerges as the main pillar for the objective liability of high-risk systems. This research, which is bibliographical, dogmatic, and documentary in nature, concludes that the solution to the current regulatory gap tends towards the adoption of full strict liability for high-risk AI systems, based on the theory of activity risk, seeking a balance between the protection of victims, legal certainty, and the encouragement of technological innovation, with the indispensable guarantee of algorithmic transparency.

Keywords: Artificial Intelligence. Civil Liability. Damage. Causal Link. Risk of the Activity.

RESUMEN

La creciente integración de los sistemas de Inteligencia Artificial (IA) en la vida social y económica cotidiana plantea desafíos sin precedentes al sistema jurídico, especialmente en el ámbito de la responsabilidad civil. Caracterizada por su autonomía y, en muchos casos, por la opacidad de sus procesos de toma de decisiones (el efecto de "caja negra"), la IA desafía los pilares tradicionales de la responsabilidad, como la culpa y la causalidad. Este artículo investiga el problema central de cómo asignar responsabilidad por daños causados por sistemas autónomos en el contexto brasileño. Para ello, se realiza un análisis crítico de las instituciones clásicas del Código Civil y del Código de Protección al Consumidor, evaluando su aplicabilidad y limitaciones ante la complejidad algorítmica. Posteriormente, se profundiza en las teorías contemporáneas, con énfasis en la teoría del riesgo de actividad (artículo 927, párrafo único, del Código Civil), que emerge como el pilar principal de la responsabilidad objetiva de los sistemas de alto riesgo. Esta investigación, de carácter bibliográfico, dogmático y documental, concluye que la solución al actual vacío regulatorio tiende a la adopción de una responsabilidad objetiva plena para los sistemas de IA de alto riesgo, basada en la teoría del riesgo de actividad, buscando un equilibrio entre la protección de las víctimas, la seguridad jurídica y el fomento de la

innovación tecnológica, con la indispensable garantía de la transparencia algorítmica.

Palabras clave: Inteligencia Artificial. Responsabilidad Civil. Daño. Nexos Causales. Riesgo de Actividad.



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

INTRODUÇÃO

A sociedade contemporânea testemunha uma transformação profunda, frequentemente denominada Quarta Revolução Industrial, impulsionada pela ubiquidade da Inteligência Artificial (IA). Seus sistemas permeiam desde as interações mais simples, como assistentes virtuais e recomendações de *streaming*, até operações de alta complexidade e potencial risco, como diagnósticos médicos, gestão de tráfego aéreo, análise de crédito e a condução de veículos autônomos. A rápida expansão dessa tecnologia, contudo, levanta questões cruciais para o campo do Direito, notadamente no que tange à proteção de direitos individuais e coletivos frente a possíveis danos gerados por decisões ou ações de sistemas autônomos.

Nesse contexto de transformação digital acelerada, a inteligência artificial (IA) passou a ser incorporada também como ferramenta estratégica para aumentar a eficiência do sistema de justiça, especialmente no desempenho de tarefas repetitivas e na elaboração de minutas de decisões. Desta forma, o desenvolvimento da inteligência artificial ganhou complexidade com o aprendizado de máquina (*machine learning*), que possibilita aos sistemas aprimorar seu desempenho a partir da análise de dados e experiências acumuladas. Esse avanço tornou-se ainda mais acelerado com as redes neurais artificiais, inspiradas no funcionamento do cérebro humano, que viabilizam a chamada aprendizagem profunda (*deep learning*).

Nesse sentido, a presente pesquisa justifica-se pela crescente lacuna normativa no Brasil. O Código Civil (Brasil, 2002), principal instrumento de regulação da responsabilidade civil, foi concebido em um contexto em que a tecnologia não possuía o grau de autonomia e imprevisibilidade da IA atual. Essa ausência de diretrizes claras gera insegurança jurídica, afetando não apenas as vítimas de danos, que enfrentam dificuldades probatórias para obter reparação (o que se agrava em casos de decisões algorítmicas discriminatórias ou falhas em

sistemas de alto risco), mas também desenvolvedores, operadores e usuários de IA, que carecem de parâmetros sobre suas obrigações e riscos.

Diante desse contexto, emerge o problema central desta pesquisa: quem deve ser responsabilizado pelos danos causados por sistemas de inteligência artificial no ordenamento jurídico brasileiro? A complexidade reside na autonomia dos sistemas, que podem operar sem intervenção humana direta, e no “efeito caixa-preta” (*black box*), que torna seus processos decisórios opacos e de difícil auditoria, pulverizando a cadeia de causalidade e dificultando a identificação do agente causador do dano.

A hipótese principal deste trabalho é que o modelo atual de responsabilidade civil brasileiro, baseado predominantemente na culpa e no nexo causal direto, mostra-se insuficiente para lidar com os riscos e desafios impostos pela Inteligência Artificial. A solução mais adequada e coerente com o desenvolvimento tecnológico reside na adoção de um regime de responsabilidade objetiva integral para sistemas de IA de alto risco, que se apoie fundamentalmente na teoria do risco da atividade (art. 927, parágrafo único, do Código Civil), complementado pela aplicação do Código de Defesa do Consumidor e pela garantia da transparência algorítmica prevista na Lei Geral de Proteção de Dados.

O objetivo geral deste trabalho é analisar a aplicabilidade e os desafios do instituto da responsabilidade civil no ordenamento jurídico brasileiro diante do uso da inteligência artificial, propondo critérios que garantam segurança jurídica e proteção aos envolvidos.

INTELIGÊNCIA ARTIFICIAL: DEFINIÇÕES, CARACTERÍSTICAS E O IMPACTO JURÍDICO

A Inteligência Artificial (IA) é um campo da ciência da computação que se dedica ao desenvolvimento de sistemas capazes de simular e executar funções cognitivas humanas, como raciocínio, aprendizado, percepção e tomada de decisão. Para o Direito, a relevância da IA reside em duas características centrais que desestabilizam o regime tradicional de responsabilidade: a autonomia e a opacidade algorítmica.

Conforme o Conselho Nacional de Justiça na plataforma Sinapses¹¹, o conceito de IA é

¹¹ Em agosto de 2020, foi aprovada a Resolução n. 332/2020 que instituiu o Sinapses como plataforma nacional de armazenamento, treinamento supervisionado, controle de versionamento, distribuição e auditoria dos modelos de Inteligência Artificial, além de estabelecer os parâmetros de sua implementação e funcionamento no âmbito do judiciário brasileiro.

aplicado em especial para soluções tecnológicas que se mostram capazes de realizar atividades de um modo considerado similar às capacidades cognitivas humanas. Assim, uma solução de IA envolve um agrupamento de várias tecnologias – redes neurais artificiais, algoritmos¹², sistemas de aprendizado, grande volume de dados (Big Data), entre outros – que fornecem os insumos e técnicas capazes de simular essas capacidades, como o raciocínio, a percepção de ambiente e a habilidade de análise para a tomada de decisão.

Cumprе esclarecer que a autonomia atribuída aos sistemas de inteligência artificial é de natureza estritamente técnica e operacional, não se confundindo com a autonomia jurídica própria dos sujeitos de direito. Trata-se de uma capacidade funcional de processamento e adaptação, incapaz de fundar imputação direta de deveres ou responsabilidades, mas suficiente para tensionar os modelos tradicionais de atribuição causal.

Como bem destacam Gustavo Tepedino e Rodrigo da Guia Silva (2019, p. 61), a complexidade e a autonomia desses sistemas “lançam influxos” diretos sobre cada um dos “elementos tradicionalmente exigidos para a deflagração do dever de indenizar”, forçando uma reavaliação da dogmática da responsabilidade civil.

Autonomia Decisória e Opacidade Algorítmica: o Desafio do Nexo Causal no Efeito “Caixa-preta”

O avanço da IA é impulsionado por subcampos como o *Machine Learning* (aprendizado de máquina) e o *Deep Learning* (aprendizado profundo). Nesses modelos, algoritmos são treinados com grandes volumes de dados para identificar padrões e tomar decisões sem serem explicitamente programados para cada cenário. A autonomia refere-se à capacidade de um sistema de IA operar e tomar decisões sem intervenção humana direta, aprendendo e se adaptando com base em novas informações.

Essa autonomia cria um vácuo, ou ao menos uma tensão profunda, na cadeia de causalidade tradicional. Em um sistema de responsabilidade subjetiva, busca-se a conduta humana (ação ou omissão) que, ligada por um nexo causal, deu causa ao dano. Quando um sistema autônomo – como um carro sem motorista ou um robô cirúrgico – toma uma decisão inesperada que resulta em dano, torna-se extremamente difícil, e por vezes impossível, reconectar

¹² Conforme o artigo 3º, I da Resolução 332/2020 do CNJ, considera-se algoritmo, a sequência finita de instruções executadas por um programa de computador, com o objetivo de processar informações para um fim específico.

essa decisão a uma falha específica e culposa do desenvolvedor, do operador ou do usuário.

A dificuldade reside em determinar se o dano decorreu de um erro de design ou programação, onde se poderia cogitar a culpa *in eligendo* ou *in vigilando* do desenvolvedor; de um erro de *input* de dados, caso os dados de treinamento se mostrem viciados ou incompletos, induzindo o sistema a padrões discriminatórios ou falhos; de um erro de operação, no qual o operador humano utilizou o sistema de forma inadequada ou negligenciou suas obrigações de supervisão; ou, finalmente, da própria decisão autônoma e imprevisível do sistema.

Nesta última hipótese, a mais desafiadora, o dano resulta de um aprendizado ou adaptação do sistema que estava objetivamente além da previsibilidade humana. Esta imprevisibilidade funcional, inerente aos sistemas de *deep learning*, pulveriza a responsabilidade na cadeia de fornecimento. Surge, assim, a necessidade de reflexão sobre novos paradigmas, como o que Ugo Ruffolo (2017 apud Tepedino; Silva 2019, p. 85) denomina “responsabilidade por algoritmo”, questionando se os regimes tradicionais são suficientes para lidar com danos causados por “produtos defeituosos inteligentes”.

Um dos maiores desafios impostos pela inteligência artificial ao direito é a opacidade de seus processos decisórios, fenômeno conhecido como “efeito caixa-preta” (*black box*). Diferentemente dos softwares tradicionais, baseados em uma lógica determinística de “se-então” (*if-then*), muitos modelos avançados de *deep learning* (aprendizado profundo) operam através de redes neurais multicamadas tão complexas que se tornam inescrutáveis. Isso significa que, muitas vezes, nem mesmo os desenvolvedores conseguem explicar, retroativamente, o raciocínio lógico exato que levou o sistema a uma decisão específica.

Essa preocupação não é exclusiva do ordenamento brasileiro. No âmbito europeu, a proposta do *Artificial Intelligence Act* adota uma abordagem baseada no risco, reconhecendo que determinados sistemas exigem regimes de responsabilidade reforçados justamente em razão de sua opacidade e autonomia decisória.

Desta forma, a falta de *accountability* algorítmica – entendida como a existência de instrumentos jurídicos e institucionais capazes de garantir responsabilização, transparência e possibilidade de auditoria dos sistemas automatizados, gera um cenário normativo permissivo. Conforme preleciona Frank Pasquale (2015, p. 8):

À medida que mais decisões importantes são confiadas a algoritmos opacos, desenvolvidos por entidades privadas ou públicas sem supervisão efetiva, o risco de arbitrariedades se intensifica. Os cidadãos são julgados por sistemas que não compreendem, baseados em dados que não podem contestar e por lógicas que não foram

submetidas ao debate democrático. Essa ‘sociedade da caixa-preta’ mina a possibilidade de responsabilização e abre caminho para um novo tipo de dominação informacional, na qual os detentores da tecnologia definem unilateralmente os critérios de justiça, eficiência e aceitabilidade social.

Ademais, o “efeito caixa-preta” entra em conflito direto com o direito à informação e à explicação, garantias basilares no ordenamento jurídico brasileiro. O art. 20 da Lei Geral de Proteção de Dados (LGPD) (Brasil, 2018) estabelece o direito do titular de solicitar a revisão de decisões tomadas unicamente com base em tratamento automatizado, exigindo informações claras sobre os critérios utilizados. Contudo, a transparência não é apenas uma exigência técnica, mas uma garantia de preservação da identidade do indivíduo frente à máquina. Nesse sentido, Danilo Doneda (2020, p. 81) alerta para o risco de perdermos o controle sobre como somos vistos e julgados pelos algoritmos:

[...] a privacidade é entendida não só em termos de um *ius excludendi alios* de uma esfera íntima, mas também como direito a não sofrer indevidas interferências externas relativamente à manifestação social de sua própria identidade. A partir do momento em que somos representados e avaliados por nossos dados pessoais, a privacidade acaba por ressoar em outros aspectos da personalidade do indivíduo.

Essa necessidade de “explicabilidade” (*accountability*) foi ecoada institucionalmente pelo Poder Judiciário. A Resolução nº 332 do Conselho Nacional de Justiça (CNJ), pioneira na regulação ética da IA nos tribunais, estabelece expressamente como princípios a “transparência” e a “possibilidade de auditoria” dos sistemas (Brasil, 2020, art. 3º, V e VI), reconhecendo que não pode haver justiça onde há segredo algorítmico.

OS FUNDAMENTOS CLÁSSICOS DA RESPONSABILIDADE CIVIL E SUA INSUFICIÊNCIA DIANTE DA INTELIGÊNCIA ARTIFICIAL

A responsabilidade civil, conforme os princípios tradicionais do Direito Civil brasileiro, baseia-se na obrigação de reparar um dano. Para que o dever de indenizar se configure, a doutrina clássica exige a presença de três elementos essenciais: a conduta (ação ou omissão), o dano (lesão a um bem jurídico) e o nexo de causalidade (o vínculo fático entre a conduta e o dano). Sobre este último elemento, Flávio Tartuce (2023, p. 515) leciona que o nexo de causalidade é o “fio condutor que liga o fato ao dano”, sendo o elemento mais complexo de ser determinado em situações que envolvem múltiplas concausas, como ocorre frequentemente em ambientes digitais.

O regime-regra do ordenamento brasileiro é o da responsabilidade subjetiva, fundamentada no art. 186 e no *caput* do art. 927. Nela, o dever de indenizar surge da comprovação da culpa *lato sensu* do agente, que abrange o dolo (a intenção de lesar) e a culpa *stricto sensu*. Esta última, como leciona a doutrina majoritária (Gonçalves, 2022; Tartuce, 2023), é a violação de um dever de cuidado, manifestando-se em três modalidades distintas: a) Negligência: caracteriza-se por uma omissão. É um *non facere*, a inércia ou passividade do agente que deixa de tomar as precauções esperadas (ex.: o programador que não atualiza um protocolo de segurança); b) Imprudência: configura-se por uma ação. É um *facere*, um comportamento precipitado, sem a cautela necessária (ex.: o desenvolvedor que lança um software sabidamente instável) e c) Imperícia: é a falta de aptidão técnica ou conhecimento específico para o exercício de uma profissão, arte ou ofício (ex.: um engenheiro de software que comete um erro crasso de codificação que um profissional médio não cometeria) (Cavaliere Filho, 2024).

Em contrapartida, como exceção, o ordenamento prevê a responsabilidade objetiva, disposta no art. 927, parágrafo único. Este regime, como ensina Flávio Tartuce (2023), prescinde da análise da culpa. O dever de indenizar nasce da simples ocorrência do dano e do nexo causal, bastando que a lei assim determine ou que a atividade normalmente desenvolvida pelo autor do dano implícito, por sua natureza, represente risco para os direitos de outrem (teoria do risco-proveito ou risco-criado).

A Inadequação da Responsabilidade Subjetiva em Sistemas de Inteligência Artificial

A responsabilidade civil subjetiva, regra geral no ordenamento jurídico brasileiro conforme os artigos 186 e 927, *caput*, do Código Civil, fundamenta-se na necessária comprovação de três elementos essenciais: a conduta humana, o nexo de causalidade e a culpa (em sentido amplo, abrangendo o dolo e a culpa estrita). Essa estrutura clássica foi desenhada para regular relações interpessoais, onde é possível identificar uma vontade humana na origem do ato ilícito.

Juristas clássicos, como Clóvis Beviláqua e Caio Mário da Silva Pereira (Pereira, 2023), construíram a dogmática da responsabilidade civil brasileira centrada na ideia de que não há reparação sem erro de conduta. No entanto, a aplicação desse modelo aos danos causados por Inteligência Artificial enfrenta barreiras intransponíveis. A principal delas reside na própria definição de culpa, que pressupõe um dever de cautela humano, conforme leciona Sérgio

Cavaliere Filho (2024, p. 32):

A ideia de culpa está visceralmente ligada à responsabilidade, por isso que, de regra, ninguém pode merecer censura ou juízo de reprovação sem que tenha faltado com o dever de cautela em seu agir. Daí ser a culpa, de acordo com a teoria clássica, o principal pressuposto da responsabilidade civil subjetiva.

Ocorre que, no contexto dos algoritmos de *Deep Learning* e *Machine Learning*, a “conduta” do sistema muitas vezes se dissocia da programação original. Quando uma IA toma uma decisão autônoma que causa danos, essa decisão pode não decorrer de um erro de código ou de negligência do desenvolvedor, mas sim do próprio processo de autoaprendizado da máquina, que identificou padrões imprevistos. Tentar rastrear uma “culpa humana” nesse processo equivale a buscar uma agulha em um palheiro algorítmico.

Essa opacidade do sistema, conhecida como “caixa-preta” (*black box*), gera uma assimetria probatória severa. Exigir que a vítima comprove a negligência, a imprudência ou a imperícia do programador é impor-lhe uma prova diabólica, ou seja, impossível de ser produzida. Sobre essa dificuldade probatória que torna a culpa um critério obsoleto, destacam Gustavo Tepedino e Rodrigo da Guia Silva (2019, p. 65):

[...] o comportamento imprevisível derivado de sistemas de IA torna difícil, e às vezes quase impossível, para a vítima comprovar se eventual dano adveio de uma conduta negligente do ser humano. Nesse sentido, entende-se que a aplicação do regime de responsabilidade subjetiva aos danos causados por sistemas de IA acabaria por deixar as vítimas sem reparação em grande parte dos casos.

Ademais, a inadequação da responsabilidade subjetiva agrava-se diante da fragmentação estrutural dos sistemas de inteligência artificial. Diferentemente dos atos humanos tradicionais, nos quais é possível identificar um agente claramente responsável pela tomada de decisão, os sistemas algorítmicos são resultado de uma cadeia complexa de atores: desenvolvedores, fornecedores de dados, empresas que treinam os modelos, operadores que os implementam e usuários finais que deles se valem.

Portanto, a insistência na aplicação da teoria subjetiva resultaria, na prática, na denegação de justiça. Se a vítima tiver que provar que houve falha humana para ser indenizada, a maioria dos danos causados por veículos autônomos, diagnósticos médicos automatizados ou discriminação algorítmica restará impune. O regime subjetivo, desenhado para a era analógica, revela-se, assim, inadequado para tutelar os riscos da sociedade da informação, exigindo a

transição para modelos baseados no risco ou na causalidade objetiva.

O Código de Defesa do Consumidor como Solução Parcial

O Código de Defesa do Consumidor (CDC) apresenta um regime jurídico mais avançado para a questão, ao prever a responsabilidade objetiva do fornecedor de produtos ou serviços, independentemente da existência de culpa (arts. 12 e 14). Quando o dano causado pela IA ocorre no âmbito de uma relação de consumo – o que abrange a vasta maioria das aplicações comerciais da tecnologia, desde aplicativos de crédito até softwares de diagnóstico médico, a aplicação do CDC é direta.

O regime do CDC rompe com a necessidade de provar a culpa e foca no “defeito” do serviço ou produto, que pode ser de concepção, informação ou segurança. No entanto, a contribuição mais valiosa do microsistema consumerista para o debate da IA vem da construção jurisprudencial do Superior Tribunal de Justiça (STJ) sobre o “fortuito”.

Tradicionalmente, a responsabilidade objetiva pode ser excluída por caso fortuito ou força maior. Contudo, a jurisprudência do STJ, consolidada na Súmula 479, estabeleceu uma distinção crucial entre fortuito externo e interno: “As instituições financeiras respondem objetivamente pelos danos gerados por fortuito interno relativo a fraudes e delitos praticados por terceiros no âmbito de operações bancárias”.

O fortuito externo é o evento imprevisível e inevitável que não guarda relação com a atividade empresarial (ex.: um desastre natural). O fortuito interno, por sua vez, é o evento que, embora também possa ser inesperado ou causado por terceiro, está intrinsecamente ligado aos riscos da atividade explorada. A fraude bancária é um exemplo clássico: embora seja um ato de terceiro (o fraudador), ela só ocorre porque o banco oferece serviços digitais, sendo um risco inerente e esperado daquele modelo de negócio.

A *ratio decidendi* do STJ é clara: quem auferir o bônus de uma atividade tecnológica e lucrativa (como o *internet banking*) deve arcar com o ônus dos riscos que ela gera, não podendo transferi-los ao consumidor.

É aqui que se constrói a analogia central para a IA. A análise de acórdãos recentes do STJ, como o REsp 1.995.458/SP, relatado pela Ministra Nancy Andrighi, reforça essa tese. No referido julgado, o Tribunal reitera que falhas de segurança em sistemas automatizados, mesmo que decorrentes de “ataques” ou falhas sofisticadas, são consideradas fortuito interno, pois estão

ligadas “à organização da atividade”. A Ministra destacou em seu voto que a sofisticação das fraudes deve ser acompanhada pela sofisticação da segurança do sistema, sendo este um ônus do fornecedor.

Ora, se uma falha de segurança causada por um terceiro (hacker) é considerada fortuito interno, com mais razão deve sê-lo uma falha causada pelo próprio sistema. Propõe-se, assim, a tese do “fortuito algorítmico”. Um dano causado por uma decisão autônoma, inesperada e imprevisível de um sistema de IA não pode ser tratado como “caso fortuito” externo, sendo, na verdade, a manifestação mais pura de um risco intrínseco à atividade de desenvolver e operar sistemas complexos, autônomos e opacos. O “defeito” aqui não é uma peça quebrada, mas um algoritmo que “decide” de forma danosa. Essa falha é o core do risco da atividade.

Apesar da inegável utilidade da responsabilidade objetiva do CDC e da construção jurisprudencial do fortuito interno, este regime possui uma limitação estrutural intransponível: seu âmbito de aplicação, uma vez que o CDC regula exclusivamente as relações de consumo e os danos causados por sistemas de IA não se restringem a eles. Sobre isso, Facchini Neto e Andrade (2023, p. 93) pontuam:

Do ponto de vista técnico, porém, há algumas dificuldades para se enquadrar a responsabilidade por danos causados por IA na moldura do CDC. Isto porque a IA é um sistema de autoaprendizagem, ou seja, o artefato aprende com sua própria experiência e pode vir a tomar decisões autônomas. Assim, para a vítima, seria difícil provar um defeito do produto dotado de IA e, especialmente, que o defeito já existia quando o artefato foi lançado no mercado. Nem sempre será fácil identificar a linha entre os danos resultantes da autodecisão da IA e os danos resultantes de defeito do produto, do ponto de vista técnico. A não ser que se entenda que uma decisão autônoma, por mecanismos de autoaprendizagem, que venha a se revelar danosa, de per si deva ser considerada defeituosa.

Nesses cenários, a vítima não estaria amparada pelo CDC. Isso exige que o ordenamento jurídico forneça uma solução para além do microssistema consumerista, buscando no próprio Código Civil o fundamento para a reparação, o que direciona a análise para a Teoria do Risco da Atividade.

A TEORIA DO RISCO DA ATIVIDADE COMO PARADIGMA CENTRAL

Diante da limitação do regime consumerista para cobrir todos os danos da IA, a solução pode ser encontrada no próprio Código Civil. A principal válvula de escape para a insuficiência do regime da culpa é a cláusula geral de responsabilidade objetiva, positivada no art. 927,

parágrafo único, do Código Civil: “Haverá obrigação de reparar o dano, independentemente de culpa, nos casos especificados em lei, ou quando a atividade normalmente desenvolvida pelo autor do dano implicar, por sua natureza, risco para os direitos de outrem”.

Esta cláusula é o pilar central para a responsabilização por danos causados por IA no Brasil. A aplicação de sistemas de IA de alto risco – especialmente aqueles dotados de autonomia e opacidade, enquadra-se perfeitamente no conceito de “atividade que implica, por sua natureza, risco”.

A definição de “risco” no art. 927, parágrafo único, foi objeto de ampla construção doutrinária e jurisprudencial. A interpretação dominante, endossada pelo Conselho da Justiça Federal (CJF) nas Jornadas de Direito Civil, afastou a ideia de que o risco deveria ser “excepcional” ou “anormal”.

O Enunciado 38 do CJF é taxativo ao afirmar que “a responsabilidade fundada no risco da atividade (art. 927, parágrafo único) configura-se quando a atividade normalmente desenvolvida pelo autor do dano, mesmo lícita, expõe outrem a uma probabilidade maior de dano, ínsita à sua própria natureza”. Ou seja, o risco não precisa ser extraordinário (como o de uma usina nuclear); basta que a atividade, pela sua própria natureza, aumente a probabilidade de danos para terceiros. O desenvolvimento e a operação comercial de sistemas algorítmicos complexos, autônomos e opacos são, por definição, atividades que se enquadram nesse preceito.

Complementarmente, o Enunciado 448 do CJF reforça que a atividade pode ser considerada de risco “quando os danos dela decorrentes, ainda que sem grande frequência, se apresentarem com particular gravidade ou desproporção”. No caso da IA, um dano (como um diagnóstico errado ou um acidente de carro autônomo) pode ter consequências de extrema gravidade.

A doutrina, liderada por juristas como Sérgio Cavalieri Filho, Gustavo Tepedino e Nelson Rosenvald, costuma classificar a teoria do risco em diferentes vertentes, todas convergindo para a responsabilização do agente que controla a atividade.

O Risco-Proveito (ou *Ubi emolumentum, ibi onus*) é a vertente mais clássica: quem auferir lucro (proveito) com uma atividade deve arcar com os danos (ônus) que ela causar. A exploração comercial da IA é, inegavelmente, uma atividade lucrativa que atrai essa lógica.

Contudo, a doutrina majoritária, incluindo Cavalieri Filho, posiciona a responsabilidade do art. 927, parágrafo único, na vertente do risco criado. O fundamento aqui não é o lucro, mas o simples fato de o agente ter criado e introduzido um novo risco na sociedade. Ao desenvolver e

comercializar um sistema de IA, o agente cria um risco de dano (o “fortuito algorítmico”) que antes não existia, e deve responder por ele, independentemente de ter tido lucro ou não.

A responsabilidade objetiva fundada no risco da atividade responde de maneira mais adequada à assimetria informacional estrutural existente entre os operadores de sistemas de IA e as vítimas dos danos. Os desenvolvedores, fornecedores e operadores detêm controle técnico, acesso aos dados, capacidade de monitoramento e poder de intervenção sobre os sistemas, encontrando-se em posição privilegiada para prevenir, mitigar ou absorver os efeitos danosos decorrentes do funcionamento algorítmico. A vítima, por sua vez, permanece alheia ao funcionamento interno da tecnologia, sendo irrazoável transferir-lhe o ônus dos riscos que não criou nem controla.

Sob essa perspectiva, a teoria do risco aproxima-se do critério da assunção do risco (*best risk bearer*), consagrado na análise funcional da responsabilidade civil, segundo o qual deve responder pelo dano aquele que se encontra em melhores condições de evitar sua ocorrência ou de socializar seus custos. No caso da inteligência artificial, o agente que introduz o sistema no mercado e dele retira benefícios econômicos ou estratégicos é quem melhor pode diluir os prejuízos por meio de seguros, repasse de custos, fundos de compensação ou ajustes no próprio design tecnológico.

Por fim, a adoção do risco da atividade como paradigma central evita soluções casuísticas ou artificiais de imputação, preservando a coerência sistêmica do direito civil. Em vez de forçar a identificação de uma culpa individual inexistente ou de recorrer a ficções jurídicas, reconhece-se que os danos algorítmicos decorrem de riscos estruturalmente associados a uma atividade lícita, porém potencialmente lesiva, cuja exploração deve ser acompanhada da correspondente assunção de responsabilidade.

Os Sujeitos da Responsabilidade e a Pulverização da Cadeia Causal

Definido o regime de responsabilidade objetiva aplicável aos danos decorrentes da inteligência artificial, impõe-se enfrentar a questão da legitimação passiva, isto é, a identificação dos sujeitos juridicamente responsáveis pela reparação. Diferentemente das relações tradicionais

de causalidade linear, os sistemas de IA inserem-se em cadeias produtivas complexas e descentralizadas, marcadas pela interdependência funcional entre diversos agentes econômicos.

Nesse ecossistema sociotécnico, concorrem para a produção e a operação do sistema, em graus variados de influência e controle: (i) o desenvolvedor do algoritmo, responsável pela arquitetura do modelo e por suas premissas técnicas; (ii) o fornecedor ou curador dos dados, cuja atuação impacta diretamente o treinamento, a acurácia e os potenciais vieses do sistema; (iii) o operador ou implementador, que integra a IA a um produto ou serviço específico, definindo seu contexto de uso; e (iv) o usuário final, que se vale das recomendações ou decisões automatizadas no exercício de sua atividade.

Essa multiplicidade de agentes rompe com o paradigma clássico da causalidade simples e dificulta a individualização de um nexos causal exclusivo entre a conduta de um único sujeito e o dano verificado. O resultado lesivo emerge, em regra, da interação sistêmica entre decisões humanas e processos automatizados, tornando artificial a tentativa de localizar um “autor” individual do dano. Conforme Gustavo Tepedino e Rodrigo Silva (2019, p. 85):

A novidade na matéria residiria, segundo parte da doutrina, na possibilidade de imputação do dever de indenizar também aos desenvolvedores de softwares ou algoritmos, e não apenas ao elo final da cadeia de fornecedores. Não se trata, contudo, ao menos no direito brasileiro, de conclusão propriamente inédita: conforme estabelecido pelos arts. 12 e 14 do CDC, a regra é a submissão dos variados fornecedores (incluindo o comerciante, guardadas as especificidades de sua responsabilização previstas pelo art. 13 do diploma) ao regime de responsabilidade objetiva pelos danos causados aos consumidores. Uma vez mais, o ineditismo parece estar não na solução jurídica, mas tão somente nas novas manifestações dos avanços tecnológicos sobre o cotidiano das pessoas.

Desta forma, a solução para a pulverização da responsabilidade na cadeia de fornecimento da IA reside na aplicação da responsabilidade solidária entre todos os agentes que participam da atividade de risco: desenvolvedores, fornecedores de dados, operadores e, em alguns casos, até mesmo o usuário final. Contudo, a fundamentação dogmática para tal solidariedade, fora do âmbito consumerista, é um dos pontos mais controversos do tema. A vedação à presunção de solidariedade, inscrita no art. 265 do Código Civil, impõe que ela decorra da vontade das partes ou de lei expressa. A tese que busca estender a aplicação do art. 942, *caput*, do CC que prevê a solidariedade entre coautores de ato ilícito para a responsabilidade objetiva por risco (art. 927, CC) carece de amparo doutrinário majoritário e jurisprudencial consolidado.

Trata-se de uma interpretação extensiva que, embora teleologicamente defensável para proteger a vítima, encontra barreiras dogmáticas. Portanto, a abordagem mais precisa e honesta é

reconhecer a existência de uma lacuna de *lege lata* (lei vigente). A solidariedade, neste contexto, é uma proposta de *lege ferenda* (lei a ser criada), sendo exatamente o que visa solucionar o Projeto de Lei nº 2.338/2023, ao prever expressamente a responsabilidade solidária entre o fornecedor e o operador do sistema de IA de alto risco.

Desta forma, uma vez estabelecida a solidariedade, a vítima pode acionar qualquer um dos agentes da cadeia para obter a reparação integral do dano. O agente que arcar com a indenização terá, então, o direito de regresso contra os demais corresponsáveis, na medida de sua participação no evento danoso.

Nesse contexto, a transparência algorítmica assume papel central não apenas como diretriz ética ou técnica, mas como verdadeiro pressuposto funcional da responsabilidade civil em ambientes automatizados. A possibilidade de auditoria dos sistemas de inteligência artificial, com acesso às informações relativas ao design do modelo, aos dados utilizados em seu treinamento, às métricas de desempenho e às condições concretas de operação é condição indispensável para a adequada identificação da origem do dano e, conseqüentemente, para a correta distribuição da responsabilidade entre os diversos agentes da cadeia produtiva.

A experiência europeia oferece um parâmetro relevante nesse debate. A Proposta de Diretiva de Responsabilidade por Inteligência Artificial da União Europeia (COM/2022/496 final) avança ao reconhecer expressamente as dificuldades probatórias inerentes aos sistemas opacos, instituindo uma presunção de nexo causal em favor da vítima. Nos termos da proposta, uma vez demonstrada a existência de uma falha no sistema de IA e a plausibilidade de que tal falha tenha causado o dano, transfere-se ao fornecedor ou operador do sistema o ônus de afastar a causalidade, mediante prova em contrário. Trata-se de um modelo que, sem afastar a responsabilidade objetiva, incorpora a transparência algorítmica como instrumento jurídico de equilíbrio entre proteção da vítima e adequada alocação dos riscos tecnológicos.

Dessa forma, o modelo europeu torna-se uma referência indispensável para o legislador brasileiro. Ao combinar a exigência de segurança prévia com mecanismos processuais que aliviam o fardo da vítima na hora de processar, a Europa demonstra que é possível criar um ecossistema de confiança, onde a inovação tecnológica não ocorre às custas da desproteção dos vulneráveis.

Nesse contexto, percebe-se que a pulverização da cadeia causal não pode servir como obstáculo à tutela da vítima. Ao contrário, ela reforça a necessidade de soluções jurídicas compatíveis com a complexidade tecnológica, privilegiando critérios como o controle da

atividade, a criação do risco e a capacidade de gestão dos danos, em detrimento da busca por uma imputação subjetiva individualizada. A responsabilidade civil, nesse cenário, deve funcionar como instrumento de recomposição e organização dos riscos sociais inerentes à exploração da inteligência artificial.

O Projeto de Lei nº 2.338/2023, atualmente em tramitação no Senado Federal, adota uma abordagem regulatória baseada em risco e prevê, em seu art. 33, a responsabilidade objetiva e solidária do fornecedor e do operador de sistemas de inteligência artificial de alto risco pelos danos causados. Embora represente um avanço ao positivizar expressamente a solidariedade na cadeia de fornecimento, suprimindo uma lacuna existente no Código Civil, o projeto é alvo de críticas quanto à vagueza do conceito de “alto risco” e à ausência de mecanismos probatórios específicos, como a presunção de causalidade em favor da vítima, nos moldes da proposta europeia.

Conforme observa Gustavo da Silva Melo (2024), tal omissão pode manter um ônus probatório excessivo sobre o lesado, esvaziando parcialmente a eficácia do regime objetivo. Assim, o PL nº 2.338/2023 deve ser compreendido menos como marco fundacional da responsabilidade civil por inteligência artificial e mais como instrumento de densificação normativa de princípios já consolidados no ordenamento jurídico brasileiro.

CONSIDERAÇÕES FINAIS

A presente pesquisa demonstrou que a Inteligência Artificial, por sua autonomia e opacidade, impõe desafios profundos ao regime tradicional de responsabilidade civil no Brasil. A análise crítica dos institutos clássicos revelou a insuficiência do modelo baseado na culpa para lidar com os danos causados por sistemas autônomos, tornando a busca por um novo paradigma uma necessidade premente.

Conclui-se que a solução mais adequada e coerente com o ordenamento jurídico brasileiro e com o avanço tecnológico é a adoção da responsabilidade objetiva integral para sistemas de IA de alto risco, com fundamento na Teoria do Risco da Atividade, prevista no art. 927, parágrafo único, do Código Civil. Este regime, que independe da prova de culpa, aloca o risco em quem o cria e dele auferir proveito, incentivando a inovação responsável e garantindo a reparação dos danos sofridos pelas vítimas.

A responsabilidade deve ser solidária entre todos os agentes da cadeia de fornecimento.

No entanto, reconhece-se a existência de uma lacuna *de lege lata* no Código Civil para fundamentar essa solidariedade de forma inequívoca. Por isso, conclui-se que a solidariedade entre os agentes da cadeia é uma proposta *de lege ferenda*, sendo o Art. 942, *caput*, do Código Civil o dispositivo mais próximo para uma interpretação extensiva, mas que deve ser complementado pela solução expressa trazida pelo Projeto de Lei nº 2.338/2023, garantindo à vítima o direito de buscar a reparação integral de qualquer um dos envolvidos, resguardado o direito de regresso entre eles.

O direito à explicação, previsto na LGPD, e o desenvolvimento de tecnologias de IA Explicável (XAI) são ferramentas indispensáveis para garantir a transparência algorítmica, que é um pressuposto para a efetivação da responsabilidade e para o exercício do direito de regresso.

As limitações desta pesquisa residem na natureza ainda incipiente do debate jurisprudencial sobre o tema e na constante evolução da tecnologia e da legislação. Como agenda para pesquisas futuras, sugere-se o aprofundamento, primeiramente, na Análise Econômica do Direito, para estudar o impacto da responsabilidade objetiva no mercado de IA, incluindo custos de seguro e incentivos à inovação. Em segundo lugar, propõe-se a investigação da viabilidade e dos modelos de Seguro de Responsabilidade Civil obrigatório para sistemas de IA de alto risco, como forma de garantir a solvência dos responsáveis. Por fim, recomenda-se o aprofundamento dos argumentos contrários e favoráveis à viabilidade da *E-Personality* (personalidade eletrônica), explorando se, em um futuro de IAs verdadeiramente autônomas, essa tese poderia ser revisitada.

Por fim, a responsabilidade civil, no contexto da inteligência artificial, revela-se não apenas como um mecanismo de reparação de danos, mas como um instrumento estruturante de governança tecnológica. Ao disciplinar a alocação dos riscos, incentivar padrões de segurança e orientar o comportamento dos agentes econômicos, o direito da responsabilidade assume papel normativo na conformação do desenvolvimento tecnológico. Nesse sentido, a adoção de um regime objetivo e funcionalmente orientado não se contrapõe à inovação, mas contribui para sua legitimação social, ao assegurar que o avanço tecnológico ocorra em consonância com a proteção dos direitos fundamentais.

Em suma, o Direito brasileiro possui os instrumentos necessários para responder aos desafios da IA, mas sua aplicação exige uma interpretação evolutiva e corajosa da doutrina e da jurisprudência, bem como a complementação legislativa para garantir a segurança jurídica e a proteção integral das vítimas na era da Inteligência Artificial.

REFERÊNCIAS

BRASIL. Conselho Nacional de Justiça. **Resolução nº 332, de 21 de agosto de 2020**. Dispõe sobre a ética, a transparência e a governança na produção e no uso de Inteligência Artificial no Poder Judiciário e dá outras providências. Brasília, DF: CNJ, 2020.

BRASIL. **Lei nº 10.406, de 10 de janeiro de 2002**. Institui o Código Civil. Diário Oficial da União, Brasília, DF, 11 jan. 2002.

BRASIL. **Lei nº 13.709, de 14 de agosto de 2018**. Lei Geral de Proteção de Dados Pessoais (LGPD). Diário Oficial da União, Brasília, DF, 15 ago. 2018.

BRASIL. **Lei nº 8.078, de 11 de setembro de 1990**. Dispõe sobre a proteção do consumidor e dá outras providências. Diário Oficial da União, Brasília, DF, 12 set. 1990.

BRASIL. **Projeto de Lei nº 2.338, de 2023**. Dispõe sobre o uso da Inteligência Artificial. Brasília, DF: Senado Federal, 2023.

BRASIL. Superior Tribunal de Justiça. **Recurso Especial nº 1.995.458/SP**. Relatora: Ministra Nancy Andrighi. Terceira Turma. Julgado em 09 ago. 2022. **Diário de Justiça Eletrônico**, Brasília, DF, 18 ago. 2022.

BRASIL. Superior Tribunal de Justiça. **Súmula nº 479**. As instituições financeiras respondem objetivamente pelos danos gerados por fortuito interno relativo a fraudes e delitos praticados por terceiros no âmbito de operações bancárias. Brasília, DF: Superior Tribunal de Justiça, 2012.

CAVALIERI FILHO, Sérgio. **Programa de Responsabilidade Civil**. 16. ed. São Paulo: Atlas, 2024.

DONEDA, Danilo. **Da Privacidade à Proteção de Dados Pessoais**. 3. ed. São Paulo: Thomson Reuters Brasil, 2019.

FACCHINI NETO, Eugênio; ANDRADE, Fábio Siebeneichler de. Reflexões sobre o modelo de responsabilidade civil para a Inteligência artificial: perspectivas para o direito privado brasileiro. In: SARLET, Gabrielle Bezerra Sales et al. (coord.); WESCHENFELDER, Lucas Reckziegel (coord. executivo). **Inteligência artificial e direito** [recurso eletrônico]. Porto Alegre: Editora Fundação Fênix, 2023. Disponível em: <https://fundarfenix.com.br/ebook/206iaedireito/>. Acesso em: 4 set. 2025.

GONÇALVES, Carlos Roberto. **Responsabilidade Civil**. 17. ed. São Paulo: Saraiva Educação, 2022.

MELO, Gustavo da Silva. **Inteligência Artificial e responsabilidade civil**: uma análise do anteprojeto do Marco Legal da Inteligência Artificial e do Projeto de Lei 2338/2023. Revista IBERC, v. 7, n. 1, p. 1-25, 2024.

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

PEREIRA, Caio Mário da Silva. **Responsabilidade Civil**. 13. ed. rev. e atual. por Gustavo Tepedino. Rio de Janeiro: Forense, 2023.

ROSENVALD, Nelson. **Direito Civil em Foco**. 2. ed. Salvador: JusPodivm, 2023.

TARTUCE, Flávio. **Manual de Direito Civil**: volume único. 13. ed. Rio de Janeiro: Forense; São Paulo: Método, 2023.

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.

TEPEDINO, Gustavo; TERRA, Aline de Miranda Valverde; GUEDES, Gisela Sampaio da Cruz. **Fundamentos do Direito Civil: Responsabilidade Civil**. 2. ed. Rio de Janeiro: Forense, 2020.

UNIÃO EUROPEIA. **Proposal for a Directive on adapting non-contractual civil liability rules to artificial intelligence (AI Liability Directive)**. COM(2022) 496 final. Bruxelas, 2022.