

**UNIVERSIDADE FEDERAL DE JUIZ DE FORA - CAMPUS GOVERNADOR
VALADARES**

**INSTITUTO DE CIÊNCIAS SOCIAIS E APLICADAS
GRADUAÇÃO EM DIREITO**

DANIEL ALVES LÍRIO RODRIGUES

**UMA ANÁLISE DA REJEIÇÃO ÀS DECISÕES JUDICIAIS POR INTELIGÊNCIA
ARTIFICIAL**

O que nos incomoda: a máquina ou resultado da decisão?

GOVERNADOR VALADARES

2025

DANIEL ALVES LÍRIO RODRIGUES

**UMA ANÁLISE DA REJEIÇÃO ÀS DECISÕES JUDICIAIS POR INTELIGÊNCIA
ARTIFICIAL**

O que nos incomoda: a máquina ou resultado da decisão?

Trabalho de Conclusão de Curso apresentado
ao Curso de Direito da Universidade Federal
de Juiz de Fora - Campus Governador
Valadares/MG, como requisito parcial para
obtenção do título de Bacharel em Direito

Orientador: Prof. Dr. Alisson Silva Martins

GOVERNADOR VALADARES

2025

AGRADECIMENTOS

Finda-se, neste momento, um importante ciclo de minha vida e inaceitável seria deixar de agradecer a pessoas que desempenharam tão grandes e importantes papéis em minha caminhada. Assim, agradeço:

Em primeiro lugar, a Deus, que me levou e sustentou a horizontes muito maiores do que eu jamais poderia imaginar ser capaz de alcançar.

Aos meus pais e minha irmã, que, sem exitar, sempre me incentivaram nos momentos difíceis e compreenderam a todo instante a minha ausência enquanto eu me dedicava à esta grande empreitada.

Ao Prof. Dr. Alisson Silva Martins, meu orientador, por ter prontamente aceito o convite de me guiar em um tema tão desafiador e, com dedicação, amizade e bom humor, me direcionado a resultados que, sozinho, eu não teria condições de alcançar.

Aos professores, que compartilharam não apenas conhecimento, mas também valiosos conselhos, ensinamentos e inspiração, deixo meu reconhecimento por todo o aprendizado que guardarei para a vida.

Aos verdadeiros amigos e amigas que estiveram ao meu lado com lealdade, sinceridade e generosidade, prestando apoio nos momentos de dificuldade e incerteza. Vocês tornaram essa trajetória mais leve e significativa, e cada memória será guardada com eterno carinho.

Aos colegas de curso, pela convivência e troca de experiências durante os últimos anos que me permitiram crescer não só como acadêmico, mas também como pessoa.

A todos que, de alguma forma, contribuíram para que esta conquista fosse possível, deixo registrado os meus sinceros agradecimentos.

Aos meus pais, que transformaram limites em caminhos e ausências em oportunidades. A vocês, que sempre se desdobraram para me entregar o que não tiveram, sem jamais pedir nada em troca, devo não apenas a conclusão deste curso, mas a própria possibilidade de sonhar mais alto.

Obrigado por terem me concedido as maiores riquezas que alguém pode carregar: um lar amoroso e acolhedor, valores sólidos que me sustentam e o exemplo diário de coragem, dignidade e generosidade.

E, por fim, pelo cuidado diário, pela paciência e pelo apoio em cada passo desta jornada acadêmica. O título alcançado com a aprovação deste trabalho e a superação desta fase é, antes de tudo, também de vocês! E, por isso, é a vocês que o dedico.

RESUMO

O presente trabalho tem por objetivo analisar os fatores que fundamentam a rejeição às decisões judiciais proferidas por sistemas de inteligência artificial, investigando se a resistência decorre da natureza não humana do decisor ou do resultado do julgamento. Para tanto, adota-se abordagem qualitativa e interdisciplinar, combinando pesquisa bibliográfica em doutrina jurídica, filosofia, psicologia cognitiva e documentos normativos. Inicialmente, são apresentados o histórico e os principais conceitos de inteligência artificial, com destaque para as aplicações no Direito e no processo judicial. Em seguida, examina-se a rejeição às decisões automatizadas a partir de dois eixos principais: o fenômeno do Vale da Estranheza, que evidencia respostas emocionais negativas diante de agentes artificiais que simulam o humano, e os vieses cognitivos, tais como a aversão à perda, o viés de confirmação e o viés do *status quo*, que influenciam a percepção de legitimidade. O estudo também contextualiza a utilização da IA no Judiciário brasileiro, destacando sistemas já implementados e a evolução normativa representada pelo Projeto de Lei nº 2338/2023, que busca regulamentar o uso responsável dessas tecnologias. Conclui-se que a rejeição não pode ser explicada por um único fator, mas pela interação entre reações emocionais e vieses cognitivos, o que reforça a necessidade de estratégias de implementação gradual, *design* centrado no humano e reforço da transparência e supervisão. Assim, a pesquisa contribui para o debate sobre a legitimação social do uso da inteligência artificial no processo judicial brasileiro.

Palavras-chave: Inteligência Artificial; Decisões Judiciais; Vale da Estranheza; Vieses Cognitivos; Legitimidade; Direito Processual Civil.

SUMÁRIO

INTRODUÇÃO.....	7
1. INTELIGÊNCIA ARTIFICIAL NO DIREITO: HISTÓRICO E CONCEITUAÇÃO...	10
1.1. Histórico e Conceituação Da Inteligência Artificial.....	12
1.2. Principais Conceitos e Funcionamento Das IA's Aplicadas Ao Direito.....	15
2. A REJEIÇÃO ÀS DECISÕES JUDICIAIS POR IA.....	17
2.1. O Vale Da Estranheza: Origem e Implicações.....	17
2.2. Implicação para a Legitimidade das Decisões Judiciais.....	19
2.3. Vieses Cognitivos E A Rejeição Às Decisões Judiciais Por Ia.....	20
2.3.1. Viés da Aversão à Perda e o "Jogo" Judicial.....	21
2.3.2. Viés de Confirmação e Resistência Institucional.....	22
2.3.3. Viés do Status Quo e Resistência à Mudança.....	22
2.4. A Interconexão Entre Os Fenômenos.....	23
3. A APLICAÇÃO DA IA NO JUDICIÁRIO.....	24
3.1. O Cenário Brasileiro.....	25
3.2. A Regulação Da Inteligência Artificial No Brasil.....	27
4. CONCLUSÃO.....	29
REFERÊNCIAS.....	33

INTRODUÇÃO

O Direito busca garantir a organização da vida em sociedade, e, a partir deste intuito, surge todo o ordenamento jurídico existente. Segundo Monnerat, é nesse sentido que deve ser entendida a máxima “onde está o homem está a sociedade e onde está a sociedade está o Direito” (Monnerat, 2020, p. 34). Porém, conforme expresso por Bobbio (2015), tão maior e mais complexo do que o dito ordenamento, é justamente a vida social que este busca garantir. Com isso, tem-se reconhecida a árdua missão do Direito de alcançar todos os aspectos da vivência social, principalmente ao se considerar que ela está em constante mudança e evolução. De diferente modo não poderia se dar ao considerar os incessantes avanços tecnológicos ocorridos nos últimos anos. O crescente desenvolvimento da Inteligência Artificial tem suscitado diversas discussões tanto no campo social como na academia, em projetos de leis, no mercado de trabalho, e em diversas outras áreas.

A utilização de Inteligências Artificiais Generativas (IAG's) no processo civil brasileiro tem crescido como alternativa voltada para o aprimoramento e garantia da eficiência e celeridade na resolução de conflitos em um sistema judiciário que contava, em 2024, com mais de oitenta milhões de processos em tramitação (CNJ, 2024). No entanto, sua implementação enfrenta desafios que transcendem questões estritamente técnicas e jurídicas, adentrando o campo dos fenômenos psicológicos e sociais que moldam a aceitação de novas tecnologias.

Um fenômeno particularmente relevante neste contexto é o que Sunstein e Gaffe (2024) denominam como "aversão a algoritmos", sendo este a tendência humana de preferir julgamentos e decisões proferidas por outros seres humanos em detrimento daquelas geradas por sistemas algorítmicos, mesmo quando estes demonstram desempenho superior em termos de precisão e eficiência. Esta aversão se manifesta de forma ainda mais acentuada em domínios tradicionalmente associados ao julgamento humano, como o Direito, em que as decisões impactam diretamente a vida, a liberdade e o patrimônio dos indivíduos.

A aversão a algoritmos, nos mais variados contextos, não se limita a uma simples resistência à inovação tecnológica. Segundo os autores, ela é produto de mecanismos psicológicos complexos e interconectados, incluindo : (1) o desejo de agência sobre decisões importantes; (2) reações morais ou emocionais negativas ao julgamento por máquinas; (3) a crença de que especialistas humanos possuem conhecimentos únicos e intuitivos, inacessíveis e incompreensíveis aos algoritmos; (4) a ignorância sobre o funcionamento e eficácia dos sistemas algorítmicos; e (5) o perdão assimétrico, manifestado pela maior intolerância a erros

cometidos por algoritmos em comparação com erros humanos similares (Sunstein; Gaffet, 2024, p. 01). Neste íterim, o presente trabalho busca compreender tais mecanismos de aversão algorítmica partindo da perspectiva propositiva do fenômeno do Vale da Estranheza e dos Vieses Cognitivos como sendo dois dos principais responsáveis pelo seu acontecimento.

A este respeito, temos que o fenômeno do Vale da Estranheza (*Uncanny Valley*), conforme formulado por Mori (2012), descreve a tendência de aceitação crescente de entes artificiais conforme se tornam mais semelhantes aos seres humanos, até o ponto em que essa semelhança provoca uma resposta negativa, gerando desconforto e repulsa. Embora tenha sido originalmente descrito para o campo da robótica, o conceito pode ser expandido para outras formas de tecnologia, incluindo, no presente contexto, as inteligências artificiais (IA's) que eventualmente sejam utilizadas no âmbito jurídico.

Mori observou que, à medida que um agente artificial se torna mais “humano”, sua aceitação tende a aumentar. No entanto, ao atingir um estágio de semelhança que ainda não é perfeito, ocorre um abrupto declínio na afinidade e aceitação, formando o chamado "vale" na resposta emocional humana (Mori, 2012, p.2). Ainda que sua análise tenha sido aplicada a robôs e próteses humanas, o mesmo efeito pode ser observado em sistemas de inteligência artificial que simulam capacidades cognitivas e decisórias humanas.

Em consonância a isso, é possível que a rejeição aos agentes artificiais não se limite à sua similitude imperfeita frente ao humano, mas possa persistir mesmo diante de uma semelhança imperceptível. Em tese, a mera consciência do fato de que um ser não é humano pode ser suficiente para gerar reações adversas, independentemente de sua aparência, comportamento ou mesmo capacidade técnica. Esse fenômeno pode ser observado, por exemplo, na ficção científica, como no romance ‘As Cavernas de Aço’, de Isaac Asimov (2013), no qual R. Daneel Olivaw, um detetive Android fisicamente indistinguível de um humano, ainda assim, enfrenta desconfiança e repulsa de seus colegas e da sociedade quanto ao exercício de suas atribuições.

Pensando nisso, podemos então tratar também do segundo fenômeno utilizado como base ao presente trabalho: os vieses cognitivos. Conforme expõe Kahneman (2012), estes consistem em padrões sistemáticos de erro no julgamento humano, resultantes de processos mentais automáticos que simplificam a tomada de decisões, mas que, em contrapartida, podem conduzir a percepções distorcidas da realidade. Esses mecanismos não apenas influenciam escolhas cotidianas, mas também impactam a atuação profissional em esferas altamente técnicas, como o Direito.

No contexto processual, tais vieses podem se manifestar de diferentes maneiras. A aversão à perda, por exemplo, tende a intensificar a percepção negativa da parte derrotada, sobretudo, no escopo do presente trabalho, quando a decisão é proferida por uma inteligência artificial, acentuando o desconforto de “perder para uma máquina”. Da mesma forma, a heurística da familiaridade leva à valorização das decisões proferidas por agentes humanos, em detrimento das oriundas de sistemas algorítmicos, cuja legitimidade ainda se mostra em disputa no imaginário jurídico. Por fim, o viés de confirmação pode reforçar resistências prévias quanto à confiabilidade de tais tecnologias, perpetuando uma postura de desconfiança estrutural.

Assim, a rejeição às decisões automatizadas pode ser compreendida não apenas pela ótica da estranheza diante do não humano, mas também pela influência de fatores psicológicos enraizados no modo como indivíduos, inclusive magistrados e operadores do Direito, processam informações e constroem julgamentos.

No que se refere às inteligências artificiais aplicadas ao Direito, mesmo que um sistema seja altamente sofisticado ao ponto de fornecer decisões tecnicamente perfeitas e indistinguíveis das humanas, sua aceitação ainda pode ser impossibilitada pela simples consciência de que não se trata de um operador do Direito real. Isso porque, conforme explica Max Weber (2015), a legitimidade do poder e, consequentemente, das decisões judiciais, depende da efetiva crença dos cidadãos na legalidade dos procedimentos que originam as normas. Nas sociedades modernas, essa legitimidade se fundamenta principalmente no princípio legal-racional, que pressupõe que as normas (e as decisões delas resultantes) devem advir de processos formalmente estabelecidos e reconhecidos socialmente. Assim, se um sistema artificial não estiver vinculado a essa tradição de autoridade, ou seja, se não for verdadeiramente percebido como parte integrante e socialmente aceita de um ordenamento jurídico, ele pode ser tido como destituído da legitimidade necessária, mesmo que suas conclusões sejam tecnicamente corretas.

Conclui-se, então, que a autoridade de uma decisão judicial não se baseia apenas na estruturação lógica de argumentos, mas também na legitimidade da figura que a emite. Com isso, um sistema artificial, por mais avançado que seja, pode ser percebido como ilegítimo simplesmente por não ser humano. Dessa forma, resta evidenciada a necessidade de proposição de uma perspectiva inovadora sobre a relação entre tecnologia e percepção da justiça, abordando novos aspectos na literatura jurídica em uma tentativa de oferecer subsídios

para uma regulamentação mais eficiente e humanizada do uso da IA no Processo Civil Brasileiro.

Para alcançar esse propósito, o presente trabalho tem como objetivo central analisar os fatores que fundamentam a rejeição às decisões judiciais proferidas por sistemas de inteligência artificial, indagando se essa resistência decorre da condição não humana do decisor ou da experiência do resultado processual. Para tanto, o estudo está estruturado em quatro capítulos, para além da introdução. O primeiro capítulo apresenta o histórico e os principais conceitos relativos à inteligência artificial, delimitando suas categorias e aplicações jurídicas para estabelecer as diretrizes centrais sobre o tema. O segundo capítulo examina a rejeição às decisões automatizadas a partir de dois eixos teóricos centrais, sendo estes o fenômeno do Vale da Estranheza e os vieses cognitivos, demonstrando de que forma ambos influenciam a percepção de legitimidade. O terceiro capítulo analisa a aplicação prática da inteligência artificial no Judiciário brasileiro, destacando experiências já em curso e a evolução normativa, com atenção especial ao Projeto de Lei nº 2338/2023. Por fim, o quarto capítulo apresenta as conclusões e proposições do trabalho, apontando que a rejeição às decisões automatizadas resulta da interação entre fatores emocionais e cognitivos, e sugerindo medidas de implementação gradual, transparência e supervisão humana como caminhos para sua aceitação social e institucional.

1. INTELIGÊNCIA ARTIFICIAL NO DIREITO: HISTÓRICO E CONCEITUAÇÃO

Neste momento, faz-se necessário e relevante o esclarecimento ao fato de, tratando-se de Inteligências Artificiais, termos uma extensa e profunda série de classificações e categorizações possíveis, as quais possuem, cada uma delas, uma abrangência conteudista suficiente para a escrita de inúmeras outras produções científicas. A título exemplificativo, podem ser mencionadas as inteligências artificiais generativas, voltadas à criação autônoma de novos conteúdos, como textos, imagens ou sons, a partir de padrões previamente aprendidos; os ‘Large Language Models’ (LLM’s), sistemas baseados em redes neurais treinadas em grandes volumes textuais para assimilar e gerar linguagem natural de maneira sofisticada (Sandhu, 2024); os modelos de aprendizado supervisionado e não supervisionado, cuja lógica reside, respectivamente, na identificação de padrões a partir de dados rotulados ou na descoberta de correlações em bases de dados sem categorização prévia; bem como o aprendizado por reforço, no qual o algoritmo aprimora seu desempenho pela lógica de

tentativa e erro, recebendo “recompensas” ou “punições” em função das escolhas realizadas (IBM, 2024).

Considerando este fato, optou-se no presente trabalho por uma abordagem abrangente do tema sem a aspiração de definir categoricamente qual das diversas construções algorítmicas possuiria mais ou menos impacto dos fenômenos descritos.

Dito isso, é sabido por todos que nos encontramos atualmente em uma era marcada pela dissolução de fronteiras tradicionais: o tempo, o espaço, o saber, o controle e até mesmo o próprio sujeito são constantemente ressignificados. Essa nova tessitura social é impulsionada pelo que Lacerda descreve como sendo uma sociedade na qual não se conhece mais o conceito de fronteiras, transmutando-se a noção de liberdade, poder, comunicação e democracia. Assim se caracteriza a sociedade da informação, impulsionada pela notável revolução tecnodigital operada nas últimas décadas (LACERDA, 2024). Esse cenário, ainda que repleto de promessas de inovação, também carrega consigo um sentimento difuso de descompasso e estranhamento. A partir da descrição precisa e provocadora do autor, temos que na sociedade da informação a velocidade de transformação é uma constante e, assim, os seus integrantes são invariavelmente tomados por uma certa estranheza sempre que sentem os impactos das mudanças promovidas, especialmente ao tentar entender o estado da arte em determinada área. Não há na contemporaneidade um sujeito sequer que não se sinta surpreendido ou ultrapassado rotineiramente, pois é impossível participar e se inteirar de todas as transformações operadas.

Isto significa que mesmo o mais cético dos indivíduos perceberá, ao fim e ao cabo, que os avanços tecnológicos não podem ser contidos por barreiras geográficas, culturais ou econômicas. A cultura tecnológica atual afeta a todos, pois mesmo que o ambiente em que cada pessoa esteja inserida possa interferir na intensidade com que ela percebe e se relaciona com os avanços tecnológicos isso não a torna imune ou inalcançável às transformações que marcam este momento social. O grande exemplo para tal situação tem sido o crescente alcance de redes de internet em locais que outrora seriam considerados como dificultosas ou mesmo inimagináveis, como o interior da floresta Amazônica¹, e regiões remotas do nordeste brasileiro².

¹ BBC News Brasil. *Starlink, de Elon Musk, domina internet por satélite na Amazônia com antenas em 90% das cidades*. BBC News, 21 maio 2025. Disponível em: <https://www.bbc.com/portuguese/articles/cv2edkw84zmo>. Acesso em: 21 maio 2025.

² GALILEU. *Como a internet chega em lugares remotos*. Revista Galileu, 30 setembro 2015. Disponível em:

Ademais, fica claro que a relevância da análise das inteligências artificiais no campo jurídico não pode ser mais postergada, pois trata-se de uma realidade que já nos alcança, e que seguirá, com o passar do tempo, nos desafiando em intensidades cada vez maiores. Assim, para garantir a ideal compreensão da discussão, é necessário o conhecimento de alguns elementos fundamentais, como veremos a seguir.

1.1. Histórico e Conceituação Da Inteligência Artificial

Não há, historicamente, um consenso absoluto sobre o significado da expressão “inteligência artificial”, porém, na visão de Fenoll (2023, p. 27), é possível dizer que esta descreve a possibilidade de máquinas, em alguma medida, “pensarem”, ou, de alguma forma, imitarem o pensamento humano para aprender a utilizar as habilidades empregadas pelas pessoas na tomada de decisões frequentes. Visando garantir o seu funcionamento, a inteligência artificial seria capaz de processar a linguagem, “entendendo” o que se expressa, como ocorre com um celular ao se mencionar o nome de uma pessoa para ligação, ou mesmo para a tradução de um texto à outro idioma. Segundo o autor, a palavra-chave ao se falar em inteligência artificial seria “algoritmo”, compreendido como um esquema executivo da máquina, responsável por armazenar todas as opções de decisão em função dos dados que vão sendo conhecidos. Estes, normalmente, são representados pelos chamados “diagramas de fluxos”, que são a descrição básica do esquema.

Nessa mesma esteira, compreendida a inteligência artificial como a capacidade de máquinas imitarem certos aspectos do raciocínio humano por meio de algoritmos, é necessário esclarecer, de forma ordenada, como esses sistemas funcionam na prática. O ponto de partida é a forma de linguagem da máquina, que só compreende dados convertidos em *bits*, ou seja, informações codificadas digitalmente em zeros e uns. Esses dados chegam ao sistema por meio do que se chama de *input*, que representa a entrada de informações a serem processadas, como números, comandos, texto ou imagem. Uma vez recebidos, esses dados percorrem uma sequência de instruções previamente programadas, executadas por meio do algoritmo. Esse algoritmo atua como um roteiro fechado, não sendo ele responsável por formular interpretações nem realizar julgamentos subjetivos, apenas aplicar, com rigor, regras que são previamente definidas. Ao final do processamento, a máquina retorna uma resposta, o chamado *output*, que precisa estar diretamente relacionado ao *input* recebido. Como observa Valentini (2017, p. 43-44), o algoritmo deve ser estruturado de forma que cada passo do

<https://revistagalileu.globo.com/Caminhos-para-o-futuro/Desenvolvimento/noticia/2015/09/como-internet-chega-em-lugares-remotos.html>. Acesso em: 21 maio 2025

processo seja simples, objetivo e preciso, a fim de garantir que a operação termine em tempo razoável e produza um resultado verificável. Essa característica de finitude é essencial, pois impede que o sistema entre em ciclos infinitos sem entregar uma resposta. Assim, o que se identifica como “inteligência” nos sistemas artificiais é, na verdade, o encadeamento lógico de pequenas operações, repetidas em alta velocidade, capazes de resolver tarefas complexas sem qualquer grau de consciência ou intencionalidade.

Ademais, conforme relata Lacerda (2022), o conceito de inteligência artificial, em seu período embrionário, referia-se à uma tentativa de emular os processos cognitivos humanos através dessas estruturas computacionais. Na visão do autor, a expressão “inteligência artificial” teria sido formulada em 1955, por um grupo de pesquisadores do Dartmouth College, uma universidade estadunidense. Neste momento fora proposta a criação de um comitê científico que tinha como objetivo estudar os benefícios potenciais da simulação da inteligência humana através de máquinas; e sua hipótese central era a de que elementos como o aprendizado e a capacidade de raciocínio poderiam ser descritos de maneira formal, o que tornaria possível a sua replicação por dispositivos computacionais programáveis.

Apesar desse marco institucional, o momento exato em que a IA passou a ser reconhecida como um campo autônomo de investigação ainda é objeto de debate. Isso porque, com base em leitura doutrinária especializada o processo de consolidação da IA é tido como sendo dividido em três grandes fases: um período de entusiasmo entre 1953 e 1969, marcado por avanços teóricos e grandes expectativas; um segundo momento, mais cético, denominado de “realismo”, entre 1966 e 1973; e, por fim, a etapa industrial da IA, que se inicia em 1980 e perdura até os dias atuais. Tais autores também reconhecem que os fundamentos conceituais da disciplina remontam às décadas de 1940 e 1950, sendo o projeto do Dartmouth College o grande ponto de inflexão. Por sua vez, Jahanzaib Shabbir e Tarique Anwer, também referidos pelo autor, atribuem a gênese do campo aos esforços de guerra conduzidos por Alan Turing, o considerado pai da computação, que ao atuar na quebra de códigos alemães durante a Segunda Guerra Mundial, teria estabelecido as primeiras bases teóricas da inteligência artificial. Segundo essa vertente, o próprio Turing teria cunhado a expressão em artigo publicado na década de 1950.

Ainda segundo a obra de Lacerda (2022), os estudos de Turing antecedem os marcos cronológicos citados, remontando à década de 1930, quando o matemático propôs o célebre ‘*Entscheidungsproblem*’, dando origem ao famigerado Teste de Turing, que possuía como sua grande indagação saber se uma máquina poderia processar informações e fornecer respostas

análogas às de um ser humano. A resposta inicial ao problema baseava-se na suposição de que o raciocínio humano poderia ser reduzido a fórmulas matemáticas, hipótese que à época, logo se revelou insuficiente frente à complexidade da conduta humana. Em investigações posteriores, Turing reformulou sua proposta, sugerindo que, se uma máquina fosse capaz de convencer ao menos um terço de seus interlocutores humanos de que se tratava de outro ser humano, então seria plausível admitir que esta estivesse, de alguma forma, “pensando”.

Curiosamente, hoje, quase cem anos após a gênese do referido teste, presenciamos o curso da história ser alterado por mais uma vez ao sermos contemporâneos da primeira inteligência artificial capaz de superar tal barreira. Segundo estudos recentes, o modelo GPT-4.5 foi o primeiro sistema artificial capaz de enganar um painel de interlocutores em um teste de Turing, sendo identificado como humano por 73% dos participantes, superando, inclusive, a taxa alcançada por humanos reais³.

Contudo, ainda que o êxito de um modelo de inteligência artificial no Teste de Turing represente um marco simbólico no desenvolvimento da tecnologia contemporânea, o resultado não permite, por si só, concluir pela existência de uma inteligência artificial consciente. Isso porquê, no experimento mental conhecido como “sala chinesa”, proposto por John Searle (apud Valentini, 2017, P. 55), nos é apresentada uma analogia para ilustrar essa limitação: imaginemos uma pessoa trancada em uma sala, recebendo mensagens escritas em chinês, idioma que ela não compreende. Com o auxílio de um manual com instruções precisas, essa pessoa pode ser capaz de responder às mensagens apenas combinando os símbolos corretamente, de forma que, para um observador externo, parecerá fluente no idioma sem, no entanto, sequer compreender o que está dizendo.

A metáfora demonstra que a simulação de compreensão não equivale à compreensão genuína. Da mesma forma, uma inteligência artificial que responde com fluência e coerência apenas manipula símbolos com base em regras predefinidas, sem necessariamente atribuir sentido ao que produz. Como observa Valentini (2017, p. 56), o argumento de Searle não busca negar a possibilidade de inteligência artificial, mas apenas enfatiza que o Teste de Turing, por si só, não pode servir como prova conclusiva de sua existência. Portanto, superar o Teste de Turing é um feito tecnológico relevante, mas, para a tranquilidade dos mais céticos,

³ REVISTA FÓRUM. Veja primeiro modelo de IA a passar em um teste de Turing. *Revista Fórum*, 14 abr. 2025. Disponível em: <https://revistaforum.com.br/ciencia-e-tecnologia/2025/4/14/veja-primeiro-modelo-de-ia-passar-em-um-teste-de-turing-177403.html>. Acesso em: 21 jun. 2025.

insuficiente para justificar a equiparação entre a atuação da máquina e a experiência cognitiva humana.

1.2. Principais Conceitos e Funcionamento Das IA's Aplicadas Ao Direito

A partir do desenvolvimento histórico da inteligência artificial e de sua relação com o Teste de Turing, estabeleceu-se uma distinção teórica fundamental entre os conceitos de IA fraca e IA forte, sendo evidente que a compreensão do funcionamento e da classificação das inteligências artificiais é indispensável para a análise dos impactos de sua aplicação no processo judicial brasileiro. Dentre os principais critérios classificatórios, destaca-se a divisão em três grandes categorias, a saber: inteligência artificial restrita (ou fraca), inteligência artificial geral (ou forte) e superinteligência.

Conforme explica Valentini (2017, p. 54), a IA fraca corresponde à hipótese de que máquinas podem agir como se fossem inteligentes, simulando racionalidade e comportamento inteligente sem, contudo, desenvolverem consciência ou compreensão genuína. Já a noção de IA forte sustenta que, ao agir de forma inteligente, essas máquinas estariam efetivamente pensando, ou seja, realizando operações cognitivas comparáveis às humanas. Essa diferenciação é essencial não apenas sob o ponto de vista filosófico, mas também prático, pois os sistemas utilizados atualmente no âmbito jurídico, como algoritmos de triagem processual, classificação de peças e organização de dados, operam exclusivamente dentro da lógica da IA fraca. Estes, apesar de sua alta eficiência, permanecem vinculados a rotinas programadas e à ausência de qualquer forma de intencionalidade ou autoconsciência, o que reforça a necessidade de cautela ao se atribuírem a tais sistemas funções que exigem interpretação, valoração ou julgamento.

A inteligência artificial restrita compreende sistemas projetados para executar tarefas específicas com elevado grau de eficiência, muitas vezes superando o desempenho humano naquela função isolada. Trata-se, porém, de um tipo de inteligência incapaz de operar fora de seu escopo programado, sem qualquer autonomia adaptativa. Nesta modalidade se enquadram aqueles sistemas de inteligência capazes de resolver com eficiência problemas em uma área específica, mas que, no entanto, se mostram incapazes de solucionar problemas de quaisquer outras áreas com a mesma autonomia (Interaction Design Foundation, [2023?]).

Exemplos notórios incluem o chamado Deep Blue, sistema desenvolvido para vencer partidas de xadrez, ou os sistemas de reconhecimento facial hoje largamente empregados. Ainda dentro dessa categoria, é possível subdividir os sistemas em duas classes: as máquinas

reativas, que não armazenam memória nem aprendem com a experiência, e os sistemas com memória limitada, que, por outro lado, utilizam dados passados para orientar decisões futuras (Lacerda, 2022).

O segundo nível, correspondente à inteligência artificial geral (IAG), pressupõe um salto qualitativo em relação à IA restrita. Esta classificação refere-se a sistemas com habilidades múltiplas, capazes de reconhecer informações, contextualizá-las, tomar decisões e interagir de forma mais abrangente com o meio, simulando o raciocínio humano (Interaction Design Foundation, [2023?]). Há, neste ponto, uma incrível aproximação deste grau de IA com o consciente humano, o que sugere a existência de um potencial de simulação de consciência e intuição. Embora alguns sistemas atuais já demonstrem aspectos de criatividade computacional ou raciocínio automatizado, o nível de inteligência equiparável ao humano ainda não foi completamente alcançado. A IAG também pode ser subdividida entre máquinas cientes, que percebem objetos e sujeitos ao redor, e máquinas autoconscientes, capazes de refletir sobre si próprias e seus próprios estados internos.

A terceira e mais avançada classificação é a superinteligência, conceito que se refere a um tipo de inteligência artificial superior à humana em praticamente todos os domínios. Conforme destaca Lacerda (2022), essa categoria poderia envolver desde sistemas com criatividade científica e habilidades sociais até formas de inteligência milhares de vezes superiores ao intelecto humano. Embora tal estágio ainda não tenha sido concretizado, sua possibilidade tem gerado debates éticos e jurídicos relevantes.

Esses graus de inteligências artificiais se articulam com três pontos centrais da aplicação contemporânea da IA: a organização de dados, o auxílio à tomada de decisões e a automação da decisão. O primeiro desses aspectos, a coleta e o tratamento de dados, constitui o alicerce de todo funcionamento inteligente não humano. Segundo Lacerda (2022), os dados são o epicentro da inteligência artificial, de modo que sua ausência poderá, em breve, ser considerada uma falha estrutural grave em qualquer ambiente organizacional. A análise eficiente dessas informações, por meio de algoritmos matemáticos, viabiliza decisões mais objetivas, baseadas em métricas e indicadores, reduzindo a influência de fatores emocionais e subjetivos.

Em um segundo plano, a inteligência artificial atua como ferramenta de apoio à decisão humana. Nesse cenário, os algoritmos contribuem para aprimorar a qualidade das decisões, sem, no entanto, substituí-las.

Finalmente, no terceiro estágio, surge a automação da decisão, em que a inteligência artificial deixa de ser um instrumento auxiliar e passa a deliberar autonomamente. Nesse caso, o agente humano se ausenta do processo decisório, confiando à máquina a responsabilidade final pela resposta a determinada situação. Como destaca o autor, apoiando-se em Kaplan e Haenlein, essa automação já permite imaginar sistemas que, de forma independente, realizam negociações comerciais, solucionam demandas de consumidores ou até atuam como tutores educacionais (Lacerda, 2022), e, em nosso caso, realizem de maneira autônoma a condução e julgamento de casos.

A classificação em graus de inteligência, associada aos três eixos funcionais descritos, oferece um panorama teórico fundamental para refletir sobre o impacto da IA nas práticas jurídicas. A depender do nível de autonomia envolvido, a resistência à sua adoção por parte dos operadores do direito pode se intensificar, especialmente quando confrontada com a ideia de que máquinas tomem decisões com efeitos concretos sobre direitos fundamentais.

2. A REJEIÇÃO ÀS DECISÕES JUDICIAIS POR IA

2.1. O Vale Da Estranheza: Origem e Implicações

A relação entre humanos e máquinas sempre foi permeada por uma constante entre fascínio e desconforto e, desde os primeiros autômatos mecânicos até os sofisticados sistemas de inteligência artificial contemporâneos, observa-se uma tensão fundamental entre o desejo de criar entidades que emulem características humanas e a inquietação provocada pelo alcance do sucesso em tal empreitada. Essa ambivalência encontra sua expressão mais precisa no conceito do Vale da Estranheza (*Uncanny Valley*), fenômeno que, embora originalmente concebido no contexto da robótica, oferece fundamentos valiosos para compreender a dinâmica de uma eventual resistência às decisões judiciais proferidas de maneira autônoma por sistemas de inteligência artificial.

O mencionado conceito foi introduzido pelo roboticista japonês Masahiro Mori (2012), que observou um curioso padrão na resposta emocional humana a robôs: à medida que estes se tornavam mais semelhantes aos seres humanos, a afinidade das pessoas por eles aumentava gradualmente, até chegar em um ponto crítico no qual, ao alcançar determinado patamar e similitude, especificamente quando os robôs atingiam um grau de características muito próximas, mas ainda perceptivelmente distintas do humano, ocorria uma queda abrupta na resposta afetiva, gerando sensações de estranheza, desconforto e repulsa. Essa afinidade voltaria a crescer apenas após o robô atingir um nível de semelhança praticamente perfeito,

formando assim um “vale” no gráfico que relaciona a semelhança humana e a resposta afetiva, vejamos:

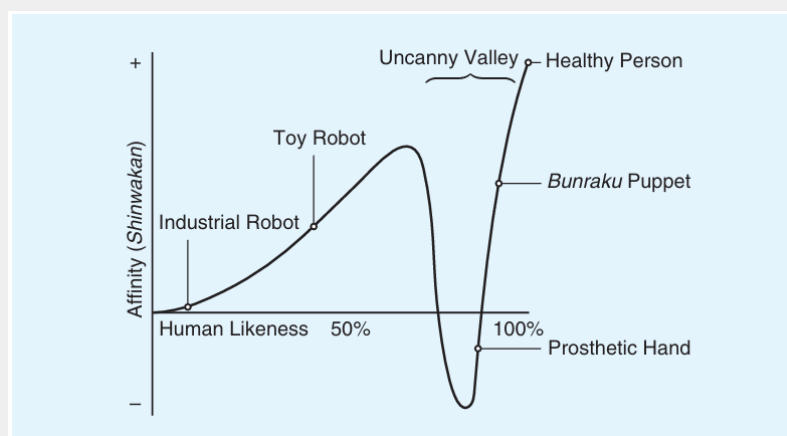


Figura 1: Gráfico demonstrando a relação entre similaridade humana e afinidade emocional.

Fonte: Mori; Masahiro (2012, p. 3)

Curiosamente, a ficção científica captou essa inquietação muito antes de ela se tornar uma realidade concreta. Duas décadas antes do surgimento acadêmico do referido conceito, em *Eu, Robô*⁴, de Isaac Asimov, robôs perfeitamente racionais convivem com humanos em um cenário de crescente desconfiança, ilustrando com precisão a sensação perturbadora que emerge quando uma criação artificial começa a se aproximar demais do que entendemos como humano.

Algumas correntes teóricas sustentam que essa resistência pode ter uma origem com base biológica. Moosa e Ud-Dean (2010) propõem uma explicação evolucionária para o Vale da Estranheza, argumentando que ele está relacionado a um mecanismo de autopreservação e evitação de perigo. Segundo o referido entendimento, a reação seria originada graças a uma herança evolucionária que nos levou a evitar seres ou situações que parecessem potencialmente ameaçadoras, como cadáveres ou indivíduos doentes. Além disso, é sabido que a espécie humana já teve que lidar, ao longo de sua evolução, com outras espécies semelhantes a ela como os neandertais (*Homo neanderthalensis*) e os *Denisovanos*, que coexistiram com os primeiros *Homo sapiens* por milhares de anos⁵. O contato com esses homínídeos pode ter influenciado um instinto de cautela ou desconfiança diante de seres que

⁴ ASIMOV, Isaac. *Eu, robô*. Tradução de Aline Storto Pereira. São Paulo: Aleph, 2014

⁵ GALILEU. Humanos e Neandertais conviveram por "muitos milhares de anos" na Europa. *Revista Galileu*, 1 fev. 2024. Disponível em: <https://revistagalileu.globo.com/ciencia/arqueologia/noticia/2024/02/humanos-e-neandertais-conviveram-por-muitos-milhares-de-anos-na-europa.ghtml>. Acesso em: 15 fev. 2025.

pareciam humanos mas que, ao fim e ao cabo, não eram. No caso das IAs jurídicas, a desconfiança surgiria não apenas da possibilidade de erros, mas, sim, da percepção de que um sistema não-humano está assumindo um papel essencial e tradicionalmente humano na aplicação do Direito.

Independentemente da explicação adotada, compreender este fenômeno no contexto do Direito é essencial para garantir a antecipação de desafios na implementação de sistemas de inteligência artificial na esfera jurídica. O que se busca esclarecer é o fato de que a aceitação dessas tecnologias pode depender não apenas de sua eficiência, mas também de como são introduzidas na prática jurídica, a fim de minimizar reações negativas e resistências culturais ou institucionais.

Particularmente relevante para o contexto jurídico é a constatação de que o Vale da Estranheza não se manifesta apenas em relação a entidades físicas, como robôs humanoides, mas também em interações com entidades virtuais que exibem comportamentos cognitivos semelhantes aos humanos. Conforme demonstrado por Saygin *et al.* (2012), a percepção de estranheza está intimamente ligada à violação de expectativas sobre comportamentos tipicamente humanos, especialmente quando estes envolvem capacidades cognitivas complexas como julgamento moral, empatia e tomada de decisões em situações ambíguas, que são precisamente as capacidades exigidas de forma constante nos contextos jurídicos.

O presente entendimento, então, vai no sentido de afirmar que, sendo um sistema de inteligência artificial limitado a funções mecânicas, como organização de processos ou análise de dados de forma transparente e objetiva, têm-se uma maior tendência à aceitação. No entanto, à medida que esses sistemas começam a assumir tarefas que tradicionalmente exigem julgamento humano (como a sugestão de sentenças ou a análise preditiva de decisões judiciais), temos o surgimento de reações de rejeição.

2.2. Implicação para a Legitimidade das Decisões Judiciais

O fenômeno do Vale da Estranheza tem implicações profundas para a legitimidade das decisões judiciais proferidas por sistemas de inteligência artificial. A legitimidade, como conceito sócio jurídico, não se reduz à correção técnica ou à conformidade normativa, mas envolve a aceitação social da autoridade decisória, dimensão esta que pode ser severamente comprometida pelos efeitos psicológicos aqui discutidos.

Weber, em sua análise clássica sobre as formas de legitimidade, as distingue entre legitimidade baseada na tradição, no carisma e na racionalidade legal-burocrática. As decisões

judiciais nas sociedades modernas derivam sua autoridade primariamente da terceira forma, fundamentando-se na aplicação racional de normas estabelecidas por procedimentos reconhecidos. Contudo, como observa Luhmann (1980), esta legitimidade formal é complementada por elementos simbólicos e rituais que conectam a decisão técnica a valores socialmente compartilhados e a uma tradição de autoridade. Os sistemas de IA judicial desafiam este equilíbrio delicado pois, por um lado, podem potencializar a dimensão racional-legal da legitimidade, oferecendo níveis superiores de consistência, previsibilidade e fundamentação técnica, e por outro, carecem dos elementos carismáticos e tradicionais que historicamente complementam a legitimidade formal das decisões judiciais.

Esta tensão manifesta-se concretamente na recepção social das decisões automatizadas. Estudos empíricos conduzidos por Tyler e Jackson (2013) demonstram que a percepção de legitimidade das decisões judiciais depende não apenas de seu resultado substantivo, mas do processo através do qual são produzidas e, particularmente, da percepção de que o decisor considerou as circunstâncias particulares do caso e as perspectivas das partes envolvidas. Esta "justiça procedimental" é frequentemente comprometida na percepção pública sobre sistemas automatizados, mesmo quando estes efetivamente incorporam tais considerações em seus algoritmos.

Esta dinâmica cria um paradoxo significativo para a implementação de sistemas de IA no judiciário, afinal, quanto mais estes sistemas se aproximam da capacidade decisória humana, que é, ao fim e ao cabo, o objetivo explícito de seu desenvolvimento, mais intensamente podem provocar o efeito do Vale da Estranheza, minando sua legitimidade percebida. Este paradoxo sugere que a aceitação social das decisões automatizadas pode não ser uma simples função de seu aperfeiçoamento técnico, mas depender de estratégias específicas para mitigar os efeitos psicológicos e culturais discutidos anteriormente. Contudo, dada a complexidade do tema, a análise da rejeição não se esgota no Vale da Estranheza; outros mecanismos psicológicos, os vieses cognitivos, também desempenham um papel crucial, como será explorado a seguir.

2.3. Vieses Cognitivos E A Rejeição Às Decisões Judiciais Por Ia

Para além da reação emocional e perceptual capturada pelo conceito do Vale da Estranheza, a resistência às decisões judiciais proferidas por sistemas de inteligência artificial encontra raízes em mecanismos psicológicos mais profundos conhecidos como vieses cognitivos. Estes atalhos mentais, embora frequentemente úteis para a tomada rápida de

decisões no cotidiano, podem levar a desvios significativos da racionalidade e da imparcialidade, especialmente em contextos complexos como o jurídico (Kahneman, 2012). A obra de Dierle Nunes, Natanael Lud e Flávio Quinaud Pedron tem destacado a presença desses vieses no processo judicial e a necessidade de mecanismos para mitigar seus efeitos (Nunes, Lud E Pedron, 2019), uma preocupação que se torna ainda mais crítica quando a interação envolve sistemas de IA. Isso porque, ao contrário do que se pode pensar, a introdução da IA no cenário judicial não elimina os vieses humanos; pode, na verdade, interagir com eles de maneiras complexas, muitas vezes amplificando a resistência ou gerando novas formas de erro.

2.3.1. Viés da Aversão à Perda e o "Jogo" Judicial

Conforme postulado por Kahneman e Tversky na Teoria da Perspectiva (Kahneman, 2012), a aversão à perda descreve a tendência humana de sentir o impacto de uma perda de forma mais intensa do que o prazer de um ganho equivalente. No contexto judicial, frequentemente percebido como um "jogo" com vencedores e perdedores (Goldschmidt, 2015), a perspectiva de uma decisão desfavorável proferida por uma IA pode ser sentida como uma perda particularmente agravada, afinal, ter sua derrota decretada por uma máquina pode carregar um estigma adicional, sendo percebido não como uma perda apenas do caso, mas também de parte da dignidade ou da validação que um julgamento humano, mesmo adverso, poderia conferir.

Dierle Nunes, ao analisar a imparcialidade dos sujeitos processuais, destaca que o processo decisório humano é permeado por uma complexa rede de emoções e valores que, embora possam eventualmente comprometer a objetividade, também conferem uma dimensão de reconhecimento intersubjetivo ao ato de julgar (Nunes *et al.*, 2019) e, assim sendo, a ausência de tal elemento humano nas decisões automatizadas pode intensificar a sensação de perda, pois elimina a possibilidade de identificação empática com o julgador.

Ademais, a falta de um interlocutor humano a quem se possa atribuir a derrota ou com quem se possa negociar simbolicamente o resultado intensifica a sensação de perda, alimentando a rejeição ao sistema automatizado. Kahneman observa que a aversão à perda frequentemente se manifesta como uma resistência desproporcional a abrir mão do *status quo*, mesmo quando alternativas objetivamente superiores estão disponíveis (Kahneman, 2012). No contexto judicial, isso pode se traduzir em uma preferência por decisões humanas potencialmente falíveis em detrimento de decisões algorítmicas estatisticamente mais

precisas. A imparcialidade fria da máquina, paradoxalmente, pode tornar a "derrota" mais difícil de aceitar do que a decisão de um juiz humano, cujas falhas e subjetividades são, de certa forma, esperadas e compreendidas dentro do sistema. Como observa Nunes, a percepção de justiça processual frequentemente depende não apenas do resultado, mas da sensação de ter sido ouvido e compreendido por um julgador capaz de empatia (Nunes *et al.*, 2019). A ausência dessa dimensão nas decisões automatizadas pode exacerbar a aversão à perda, transformando derrotas processuais em experiências que sejam psicologicamente mais custosas.

2.3.2. Viés de Confirmação e Resistência Institucional

O viés de confirmação leva os indivíduos a buscar, interpretar e recordar informações de maneira a confirmar suas crenças preexistentes, funcionando como uma forma de "pensamento motivado", no qual o desejo de manter determinadas crenças direciona seletivamente a atenção e o processamento cognitivo (Kahneman, 2012). Isto significa que operadores do Direito com visões céticas sobre a IA tenderão a focar nos erros e limitações dos sistemas, ignorando seus sucessos e potenciais benefícios, ao passo que entusiastas podem superestimar as capacidades da IA, negligenciando seus riscos, ou seja, ambas estas posturas dificultam uma avaliação equilibrada e informada. No mais, como observa Nunes, a formação jurídica tradicional, com sua ênfase na autoridade humana e na interpretação hermenêutica, pode predispor os profissionais a resistir a abordagens algorítmicas, percebendo-as como reducionistas ou incompatíveis com a complexidade do fenômeno jurídico (Nunes *et al.*, 2019).

Institucionalmente, o viés de confirmação pode se manifestar na resistência a dados que desafiem práticas estabelecidas, favorecendo narrativas que justifiquem a manutenção do *status quo* e dificultem a adoção de novas tecnologias, mesmo quando estas demonstram potencial para aprimorar a eficiência ou a consistência das decisões. Neste contexto, mesmo evidências objetivas de eficácia dos sistemas de IA podem ser filtradas através de lentes cognitivas que privilegiam a manutenção de paradigmas estabelecidos, dificultando a adoção de novas tecnologias.

2.3.3. Viés do *Status Quo* e Resistência à Mudança

O viés do *status quo*, intimamente relacionado à aversão à perda, reflete a preferência pela manutenção da situação atual, percebendo qualquer mudança como uma potencial perda. É desnecessário dizer que a introdução da IA no Judiciário representa uma mudança profunda

nas rotinas, na distribuição de poder e na própria concepção do trabalho jurídico e a resistência a essa mudança pode ser explicada, em partes, por uma preferência ao familiar e pelo receio dos custos e incertezas associados à adaptação com novos sistemas e procedimentos. Dierle Nunes observa que o campo jurídico é particularmente suscetível ao viés do *status quo*, dada sua tradição histórica, hierarquia estabelecida e mecanismos institucionais que valorizam a estabilidade e a previsibilidade (Nunes *et al.*, 2019).

Mesmo que a IA prometa ganhos de eficiência, redução de custos e maior consistência decisória, a inércia cognitiva e o conforto com os métodos tradicionais podem levar à rejeição ou à implementação hesitante. Kahneman observa que o viés do *status quo* é particularmente forte em contextos de alta complexidade ou incerteza, na qual os custos e benefícios da mudança são difíceis de quantificar precisamente (Kahneman, 2012), sendo o campo jurídico, com sua inerente complexidade normativa e social, a representação precisa desse tipo de ambiente, potencializando a resistência à inovação tecnológica.

2.4. A Interconexão Entre Os Fenômenos

A análise precedente nos torna possível a percepção de que tanto o Vale da Estranheza quanto os vieses cognitivos podem contribuir significativamente para a resistência às decisões judiciais proferidas por sistemas de inteligência artificial. Contudo, uma compreensão completa do fenômeno exige reconhecer que estes fatores não operam isoladamente, mas interagem de maneiras complexas, reforçando-se mutuamente e moldando diferentes facetas da rejeição.

O Vale da Estranheza, com sua ênfase na resposta emocional e perceptual pode atuar como um gatilho ou amplificador para certos vieses cognitivos, dado que o desconforto visceral provocado por uma IA que habita o "vale" pode intensificar a aversão à perda, tornando a derrota decretada por uma entidade ainda mais psicologicamente custosa. A seu tempo, de semelhante modo, a sensação de estranheza pode fortalecer o viés de confirmação, levando operadores do Direito a interpretar qualquer ambiguidade ou erro do sistema como sendo a prova cabal de sua inadequação fundamental, confirmando a crença preexistente de que o julgamento é uma prerrogativa exclusivamente humana. Por outro lado, os vieses cognitivos podem explicar a resistência à IA judicial mesmo em cenários onde o Vale da Estranheza não seja predominante. Sistemas de IA que não buscam emular a cognição humana, operando de forma puramente instrumental (como algoritmos de análise de risco baseados em dados estatísticos), podem ainda assim enfrentar resistência devido ao viés do

status quo (preferência pelo sistema atual) ou à aversão ao algoritmo (maior intolerância a erros de máquinas).

Neste ponto, é plausível supor que o Vale da Estranheza e os vieses cognitivos expliquem diferentes tipos ou níveis de resistência. O Vale da Estranheza pode ser mais relevante para explicar reações imediatas, emocionais e talvez mais prevalentes no público leigo ou em relação a sistemas que explicitamente tentam simular o comportamento humano (como pode ocorrer, em um futuro não muito distante, por meio de juízes robôs ou avatares judiciais); enquanto os vieses cognitivos, por sua vez, podem oferecer uma explicação mais robusta para a resistência institucional e profissional.

A hipótese central deste trabalho, ao questionar se a origem da rejeição decorre primariamente do "robô" (por meio do Vale da Estranheza) ou da "derrota" (através de vieses como aversão à perda), encontra aqui sua complexidade. A resposta mais provável é que ambos os fatores atuem de forma dinâmica. A percepção do "robô" pode modular a experiência da "derrota" (ou da simples interação com o sistema), enquanto a predisposição à "derrota" (ou à resistência à mudança, à perda de controle) molda a forma como o "robô" é percebido e avaliado.

Compreender esta interação é crucial para o desenvolvimento de meios eficazes de implementação da IA no Judiciário, haja vista que estratégias focadas apenas em superar o Vale da Estranheza (por exemplo, através de design de interface ou maior transparência) podem ser insuficientes se não abordarem simultaneamente os vieses cognitivos subjacentes ao passo que, inversamente, abordagens focadas apenas na demonstração da superioridade técnica da IA podem falhar se ignorarem as barreiras emocionais e perceptuais erguidas pelo Vale da Estranheza.

3. A APLICAÇÃO DA IA NO JUDICIÁRIO

Com a crescente utilização de inteligências artificiais pela sociedade como um todo e, principalmente, pelos operadores do Direito, podemos ser levados a acreditar que a aplicação de tais ferramentas têm sua origem datada de forma recente. Contudo, conforme ressalta Fenoll (2023, p. 25), estes instrumentos já vêm sendo utilizados há décadas, sendo restringidos, porém, a uma aplicação meramente rudimentar como processadores de textos ou buscadores de jurisprudência. Valle et al. (2023), a título de exemplo, mencionam o software TAXMAN, desenvolvido pela universidade de Rutger em Nova Jersey no ano de 1972, que atuava na fiscalização de sociedades por ações e tinha como objetivo a identificação de

mudanças no contrato que resultassem (ou não) em isenções fiscais; ou, mais recentemente, o sistema ROSS, criado em 2015 na Universidade de Toronto, que tinha como seu principal diferencial a capacidade de responder a perguntas formuladas em linguagem natural e, inclusive, sugerir a leitura de doutrinas e decisões relevantes para a resolução do caso concreto.

Ainda segundo os autores, a modernização e informatização do Poder Judiciário brasileiro data do ano de 2001, com a Lei nº 10.259/2001, que instituiu o Juizado Especial no âmbito da Justiça Federal. Isto devido ao fato de, pela primeira vez, perceber-se o cuidado do legislador na positivação da possibilidade de o Judiciário recorrer aos meios eletrônicos quando necessário. Neste mesmo ano fora também firmado convênio pelo Supremo Tribunal Federal junto ao Banco Central do Brasil (BACEN) para instaurar a ferramenta de busca de bens denominada BACENJUD (ou, SISBAJUD nos dias atuais). Contudo, destaca-se pelos autores que neste período não houve o uso da Inteligência Artificial da forma como concebemos hoje, podendo ser o início da década de 2000 considerado como um prólogo para a modernização do Poder Judiciário brasileiro que ocorreria entre 2010 e os dias atuais. Cita-se o ano de 2010 como o grande ponto de virada no âmbito da tecnologia da informação pois fora este o ano de instauração do Processo Judicial Eletrônico (PJe) na Vara de Natal da Justiça Federal do Rio Grande do Norte. Resta, então, evidente que, com o avanço da ciência, torna-se diuturno a introdução de instrumentos cada vez mais sofisticados ao contexto jurídico.

3.1. O Cenário Brasileiro

O Poder Judiciário brasileiro tem também se destacado pela adoção de soluções tecnológicas inovadoras, incluindo sistemas de IA. No âmbito da Administração Pública Nacional, os tribunais têm liderado as aplicações de inteligência artificial. De acordo com dados de pesquisa do CNJ⁶, 66% dos tribunais brasileiros têm projetos de IA em desenvolvimento e, no âmbito do Sinapses, já há registro de 147 sistemas de IA aplicados a diferentes tarefas nos tribunais. Entre os sistemas mais relevantes, destacam-se: O Sinapses⁷,

⁶ CONSELHO NACIONAL DE JUSTIÇA. Pesquisa uso de inteligência artificial (IA) no Poder Judiciário. 2023. Conselho Nacional de Justiça; Programa das Nações Unidas para o Desenvolvimento. Brasília: CNJ, 2024. Disponível em: https://bibliotecadigital.cnj.jus.br/jspui/bitstream/123456789/858/1/Pesquisa%20uso%20da%20inteligencia%20artificial%20IA%20no%20poder%20judici%C3%A1rio_2023.pdf. Acesso em: 15 jul. 2024. p. 27.

⁷ CONSELHO NACIONAL DE JUSTIÇA. **Plataforma Sinapses / Inteligência Artificial**. [S. l.], [s. d.]. Disponível em: <https://www.cnj.jus.br/sistemas/plataforma-sinapses/>. Acesso em: 30 de maio de 2025.

plataforma desenvolvida pelo CNJ que, diferentemente de um sistema específico, opera como um 'framework' possibilitando aos tribunais o desenvolvimento, treinamento e implementação de modelos de IA adaptados às suas necessidades particulares, minimizando a dependência de soluções externas; Victor⁸, desenvolvido em parceria entre o Supremo Tribunal Federal e a Universidade de Brasília, o VICTOR foi o primeiro projeto de IA do STF, focado inicialmente na identificação de temas de repercussão geral; VitorIA⁹, que busca agrupar processos através da similaridade de temas para possibilitar a identificação de novas controvérsias; e o STJ Logos¹⁰, desenvolvido inteiramente no próprio tribunal, com o objetivo de modernizar a análise e a elaboração de conteúdos judiciais, o mecanismo conta com duas funcionalidades principais: geração de relatório de decisão e análise de admissibilidade de agravos em recurso especial (AREsps). Os usuários podem fazer perguntas diversas sobre o processo na caixa de diálogo do STJ Logos e também solicitar a execução de tarefas, como a exibição de uma lista de argumentos apresentados na petição.

Estas implementações concretas demonstram que há, de fato, um avanço significativo da IA no judiciário brasileiro, mas também evidenciam os desafios relacionados à sua aceitação e legitimação pois a implementação bem-sucedida de sistemas de IA no Poder Judiciário não depende apenas de sua eficiência técnica, mas também de sua capacidade de inspirar confiança em magistrados, servidores, advogados e jurisdicionados; e essa confiança, por sua vez, está diretamente ligada à percepção de legitimidade, transparência e subordinação destes sistemas aos valores fundamentais do Estado Democrático de Direito.

Tal observação dialoga diretamente com a hipótese central deste trabalho, sugerindo que fatores psicológicos e sociais, como o Vale da Estranheza e os vieses cognitivos, podem desempenhar um papel tão ou mais importante que fatores técnicos na determinação do sucesso ou fracasso da implementação de sistemas de IA no contexto jurídico brasileiro.

⁸ SUPREMO TRIBUNAL FEDERAL. **Presidente do STF apresenta sistema de IA que auxilia na triagem de processos.** [S. l.], 15 de agosto de 2023. Disponível em: <https://portal.stf.jus.br/noticias/verNoticiaDetalhe.asp?idConteudo=471331&ori=1>. Acesso em: 30 de maio de 2025.

⁹ SUPREMO TRIBUNAL FEDERAL. **Ministra Rosa Weber lança robô VitorIA para agrupamento e classificação de processos.** [S. l.], 17 de maio de 2023. Disponível em: <https://noticias.stf.jus.br/postsnoticias/ministra-rosa-weber-lanca-rob-vitoria-para-agrupamento-e-classificacao-de-processos/>. Acesso em: 30 de maio de 2025.

¹⁰ SUPERIOR TRIBUNAL DE JUSTIÇA. **STJ lança novo motor de inteligência artificial generativa para aumentar eficiência na produção de decisões.** [S. l.], 12 de fevereiro de 2025. Disponível em: <https://www.stj.jus.br/sites/portalt/Paginas/Comunicacao/Noticias/2025/11022025-STJ-lanca-novo-motor-de-inteligencia-artificial-generativa-para-aumentar-eficiencia-na-producao-de-decisoes.aspx>. Acesso em: 30 de maio de 2025.

3.2. A Regulação Da Inteligência Artificial No Brasil

O Brasil tem também avançado significativamente na construção de um marco regulatório para a inteligência artificial, destacando-se o Projeto de Lei nº 2338/2023 (BRASIL, 2023), também conhecido como "PL da IA", que expressa um abrangente esforço no sentido de estabelecer normas gerais de caráter nacional para o desenvolvimento, implementação e uso responsável de sistemas de inteligência artificial no país, buscando alcançar o delicado ponto de equilíbrio entre a proteção de direitos fundamentais e o estímulo à inovação tecnológica. A proposta legislativa estabelece como fundamento central a primazia da pessoa humana, colocando-a como eixo norteador de toda a regulação e, para além deste princípio antropocêntrico, têm-se a complementação por outros fundamentos essenciais como o respeito aos direitos humanos, aos valores democráticos, à proteção ambiental, à igualdade, à não discriminação e ao desenvolvimento tecnológico sustentável (Brasil, 2023).

Distintivo aspecto do projeto brasileiro é a adoção de uma abordagem baseada em riscos, na qual se estabelece diferentes níveis de exigências regulatórias conforme o potencial impacto dos sistemas de IA, com o objetivo de evitar a imposição de ônus excessivos a sistemas de baixo risco, enquanto assegura um controle mais rigoroso para aplicações que possam afetar significativamente direitos fundamentais, como as utilizadas em decisões judiciais, segurança pública ou acesso a serviços essenciais.

No que tange especificamente ao uso da IA no contexto jurídico, o PL estabelece salvaguardas importantes que dialogam diretamente com as preocupações levantadas neste trabalho sobre a legitimidade e aceitação das decisões automatizadas. Entre os princípios norteadores que impactam diretamente o uso judicial da IA, destacam-se a transparência, explicabilidade e auditabilidade dos sistemas; a participação e supervisão humana efetiva; o devido processo legal, contestabilidade e contraditório; e a rastreabilidade das decisões como meio de prestação de contas (Brasil, 2023). O projeto reconhece expressamente o direito das pessoas afetadas por sistemas de IA à informação prévia sobre suas interações com tais sistemas, à explicação sobre decisões ou previsões algorítmicas, e à contestação de resultados que produzam efeitos jurídicos ou impactos significativos.

Particularmente relevante para o contexto judicial é o direito à determinação e participação humana em decisões de sistemas de IA, considerando o contexto e o estado da arte tecnológica, o que reforça a noção de que a automação completa de decisões judiciais, mesmo não sendo contemplada pelo marco regulatório em desenvolvimento, é vislumbrada como uma real possibilidade no horizonte de eventos. Ademais, a proposta legislativa também

aborda diretamente questões relacionadas à discriminação algorítmica, estabelecendo o direito à não-discriminação e à correção de vieses discriminatórios diretos, indiretos, ilegais ou abusivos. Esta preocupação é especialmente pertinente no contexto judicial, onde decisões enviesadas podem perpetuar ou amplificar desigualdades sociais existentes.

Um aspecto particularmente inovador do projeto é o tratamento dado à responsabilidade civil, que estabelece um regime de responsabilidade solidária entre fornecedores e operadores de sistemas de IA de alto risco, com possibilidade de regresso proporcional à participação no dano. Para os sistemas de baixo risco, adota-se um regime de responsabilidade subjetiva, baseado na verificação de culpa. Esta gradação busca equilibrar a proteção efetiva das vítimas com incentivos adequados à inovação responsável.

O PL também reconhece a importância da pesquisa e desenvolvimento em IA, prevendo medidas de fomento à inovação e à capacitação de recursos humanos nesta área. Estabelece, ainda, a necessidade de programas educacionais para conscientização pública sobre os impactos da IA na sociedade, elemento crucial para mitigar a aversão algorítmica discutida anteriormente.

No contexto específico do Judiciário, o projeto estabelece que sistemas de IA utilizados em decisões que afetem direitos fundamentais devem garantir a supervisão humana significativa, a transparência algorítmica e a possibilidade de contestação. Estas salvaguardas alinham-se às preocupações levantadas sobre o Vale da Estranheza e os vieses cognitivos que podem influenciar a aceitação de decisões judiciais automatizadas.

Por fim, é importante ressaltar que o PL 2338/2023 ainda está em tramitação no Congresso Nacional e pode sofrer alterações até sua eventual aprovação. No entanto, seus princípios fundamentais e diretrizes gerais já oferecem um panorama valioso sobre o caminho regulatório que o Brasil está construindo para a inteligência artificial, com implicações diretas para sua aplicação no sistema judicial brasileiro e, nesse sentido, revela uma abordagem equilibrada que busca proteger direitos fundamentais sem inibir a inovação tecnológica, reconhecendo tanto os benefícios quanto os riscos potenciais da IA no contexto jurídico. Esta regulação, quando implementada, poderá contribuir significativamente para mitigar os fenômenos de aversão algorítmica e rejeição às decisões judiciais por IA, ao estabelecer garantias de transparência, supervisão humana e contestabilidade que atendem a muitas das preocupações psicológicas e sociais aqui analisadas.

4. CONCLUSÃO

Ao longo deste trabalho, buscamos compreender os elementos psicológicos e sociais que fundamentam a rejeição às decisões judiciais proferidas por sistemas de inteligência artificial. Retomando os estudos de Sunstein e Gaffe (2024), que identifica distintas formas de aversão algorítmica, podemos, finalmente, enquadrar nossa tese central, correlacionando tais mecanismos com os fenômenos do Vale da Estranheza e dos vieses cognitivos analisados.

Os cinco mecanismos de aversão algorítmica identificados por Sunstein e Gaffe, quais sejam: (1) o desejo de agência sobre decisões importantes; (2) reações morais ou emocionais negativas ao julgamento por máquinas; (3) a crença de que especialistas humanos possuem conhecimentos únicos e intuitivos inacessíveis aos algoritmos; (4) a ignorância sobre o funcionamento e eficácia dos sistemas algorítmicos; e (5) o perdão assimétrico; não operam de forma isolada, mas se manifestam através de processos psicológicos mais profundos que podem ser compreendidos à luz do Vale da Estranheza e dos vieses cognitivos.

Em nosso sentir, o desejo de agência sobre decisões importantes (1) e as reações morais ou emocionais negativas ao julgamento por máquinas (2) podem ser perfeitamente compreendidos como manifestações diretas do fenômeno do Vale da Estranheza. O desconforto visceral provocado por entidades que simulam capacidades humanas fundamentais como o julgamento moral e a empatia, gera uma resistência instintiva à delegação de decisões importantes a sistemas artificiais. Esta resistência é particularmente pronunciada no contexto judicial, em que decisões afetam diretamente aos indivíduos e a dimensão simbólica e ritual do julgamento humano carrega um peso histórico e cultural significativo.

Por outro lado, a crença de que especialistas humanos possuem conhecimentos únicos e intuitivos inacessíveis aos algoritmos (3), bem como a ignorância sobre o funcionamento e eficácia dos sistemas algorítmicos (4), parecem derivar primariamente dos vieses cognitivos discutidos. O viés de confirmação leva operadores do Direito a valorizar excessivamente a intuição e o "feeling" jurídico, qualidades supostamente inimitáveis pela máquina, enquanto minimizam evidências de que algoritmos podem, em muitos casos, superar o desempenho humano em tarefas de previsão e classificação. Simultaneamente, o viés do *status quo* e a resistência à mudança alimentam a ignorância ativa sobre o funcionamento e a eficácia dos sistemas de IA, criando um ciclo de desinformação que perpetua a desconfiança.

Já o perdão assimétrico, quinto mecanismo identificado por Sunstein e Gaffe, e que resulta em uma maior intolerância a erros cometidos por algoritmos em comparação com

erros humanos similares, representa, ao fim e ao cabo, a consequência da influência exercida pelos mecanismos de aversão anteriores.

Esta correlação entre os mecanismos de aversão algorítmica e os fenômenos psicológicos mais amplos do Vale da Estranheza e dos vieses cognitivos oferece uma perspectiva mais nuançada sobre a questão central deste trabalho: "O que nos incomoda: a máquina ou o resultado da decisão?". A análise sugere, como dito, que ambos os fatores estão intrinsecamente entrelaçados, com o Vale da Estranheza alimentando uma resistência primária à própria ideia de julgamento por máquinas, enquanto os vieses cognitivos moldam a interpretação e avaliação dos resultados dessas decisões.

Diante dos conhecidos e complexos desafios psicológicos, técnicos e éticos que permeiam a implementação da inteligência artificial no sistema judicial, torna-se imperativo delinear estratégias que possam mitigar as resistências observadas e promover uma integração mais harmoniosa e legítima dessas tecnologias. A superação do Vale da Estranheza e a neutralização dos vieses cognitivos não são tarefas simples, e nem tão pouco se pretende esgotar a questão no presente trabalho, mas abordagens multifacetadas podem ajudar a pavimentar o caminho para um futuro onde a IA sirva como ferramenta eficaz e confiável a serviço da justiça e da sociedade.

Uma primeira linha de ação envolve o design centrado no humano e na explicabilidade pois, afinal, sistemas de IA devem ser projetados não apenas para otimizar a eficiência, mas também para inspirar confiança e facilitar a compreensão por parte dos usuários. Isso implica investir em interfaces intuitivas, fornecer justificativas claras e compreensíveis para as decisões ou sugestões algorítmicas (mesmo que simplificadas) e permitir níveis adequados de controle e supervisão humana (Nunes *et al.*, 2019). A opacidade dos algoritmos de aprendizado de máquina (também chamada de "caixa-preta") alimenta tanto o desconforto do Vale da Estranheza (por criar uma entidade que pensa de forma incompreensível) quanto vieses como o de confirmação (por dificultar a avaliação de sua confiabilidade), de forma que a transparência sobre as capacidades e limitações de atuação em cada sistema se torna fundamental para gerenciar expectativas e evitar tanto a complacência quanto a aversão algorítmica.

Em segundo lugar, é crucial promover a integração e o conhecimento técnico entre os operadores do Direito pois, como defendem Rocha e Lírío (2025), discutir IA com propriedade exige um entendimento técnico que ultrapasse o senso comum. Iniciativas de formação continuada, *workshops* e a inclusão de conteúdos sobre tecnologia e IA nos

currículos jurídicos são essenciais para capacitar juízes, advogados, promotores e defensores a interagir criticamente com essas ferramentas, compreendendo seus potenciais e riscos. Esse conhecimento é vital para desmistificar a IA, combater a "algoritmofobia" (Eigen, 2020) e permitir uma avaliação mais informada e menos enviesada.

Uma terceira estratégia reside na implementação gradual e contextualizada pois, como alternativa à introdução de sistemas de IA de forma abrupta em funções decisórias críticas, pode ser mais prudente começar (como já tem sido feito por grande parte dos tribunais ao redor do país), por tarefas de apoio, automação de rotinas e análise de dados em áreas menos sensíveis ou com maior grau de objetividade. A demonstração de valor em contextos de menor risco pode gradualmente construir confiança e familiaridade, facilitando a aceitação em etapas posteriores. Ademais, a ênfase na colaboração humano-máquina, e não na lógica de substituição, pode ser mais eficaz. Apresentar a IA como uma ferramenta de "inteligência aumentada", que potencializa as capacidades humanas em vez de suplantá-las ajuda a mitigar o ocasionamento de efeitos negativos.

No campo regulatório, é necessário evitar uma regulação baseada exclusivamente no medo, buscando um equilíbrio sensato e funcional entre a proteção de direitos fundamentais e o fomento à inovação (Rocha; Lírio, 2025), algo que, como vimos, têm sido realizado em nosso contexto. Normas excessivamente restritivas ou tecnicamente inviáveis podem sufocar o desenvolvimento tecnológico no país. A regulação deve ser adaptativa, baseada em evidências e desenvolvida através de um diálogo genuinamente multidisciplinar entre juristas, tecnólogos, cientistas sociais e a sociedade civil.

As perspectivas futuras apontam para sistemas de IA cada vez mais sofisticados e integrados ao ecossistema jurídico, e a capacidade de dialogar com outras áreas do saber e de compreender os fundamentos técnicos das ferramentas utilizadas será cada vez mais essencial para o exercício competente e ético da profissão jurídica. Nesse cenário, não se trata mais de discutir a viabilidade hipotética da atuação da IA no direito, mas sim de reconhecer que a sua aplicação prática, inclusive na atividade jurisdicional, configura uma tendência concreta e irreversível. Como destacam os autores (Nunes *et al.*, 2019, p. 146), embora a ideia de um "computador-juiz" ainda cause certo impacto ou mesmo incredulidade, ela não pode mais ser reduzida a uma mera profecia ou ficção científica. Sua emergência representa, antes, uma ruptura cognitiva no processo decisório tradicional, cuja assimilação demandará não apenas domínio técnico, mas também um novo olhar sobre os critérios de legitimidade e racionalidade no processo judicial.

O desafio reside, então, em garantir que essa integração ocorra de forma a fortalecer, e não enfraquecer, os pilares da justiça: a imparcialidade, a equidade, a transparência e a legitimidade. A superação da rejeição às decisões por IA dependerá menos da perfeição técnica dos algoritmos e mais da capacidade de construir pontes de confiança e compreensão entre a tecnologia, os operadores do Direito e a sociedade a quem servem. Compreender os mecanismos psicológicos subjacentes a essa rejeição, sejam eles derivados do Vale da Estranheza ou dos vieses cognitivos, constitui um primeiro passo essencial nessa jornada.

REFERÊNCIAS

- ASIMOV, Isaac. *As cavernas de aço*. Tradução de Aline Storto Pereira. São Paulo: Aleph, 2013. 302 p.
- BRASIL. Senado Federal. Projeto de Lei nº 2338, de 2023. Dispõe sobre o uso da Inteligência Artificial. Brasília, DF, 2023.
- CNJ. Com 84 milhões de processos em tramitação, 2024. Disponível em: <https://www.cnj.jus.br/com-84-milhoes-de-processos-em-tramitacao-judiciario-trabalha-com-produtividade-crescente/>. Acesso em: 25 ago. 2025.
- EIGEN, Zev J. Algoritmofobia: superando seu medo de algoritmos e inteligência artificial. Boletim Revista dos Tribunais Online, vol. 8, out. 2020.
- FENOLL, Jordi Nieva. Inteligência Artificial e o Processo Judicial. São Paulo: JusPodivm, 2023.
- GOLDSCHMIDT, James. Derecho, Derecho Penal y Proceso: El proceso como situación jurídica. Tradução de Jacobo López Barja de Quiroga, Ramón Ferrer Baquero e León García-Comendador Alonso. Madrid: Marcial Pons, 2015.
- IBM. O que é aprendizado supervisionado? IBM, 28 dez. 2024. Disponível em: <https://www.ibm.com/br-pt/think/topics/supervised-learning>. Acesso em: 23 ago. 2025.
- INTERACTION DESIGN FOUNDATION. What is Narrow AI? [2023?]. Disponível em: <https://www.interaction-design.org/literature/topics/narrow-ai>. Acesso em: 23 ago. 2025.
- KAHNEMAN, Daniel. Rápido e devagar: duas formas de pensar. Tradução de Cássio de Arantes Leite. Rio de Janeiro: Objetiva, 2012.
- LACERDA, Bruno Torquato Zampier. Bens digitais: cybercultura, redes sociais, e-mails, músicas, livros, milhas aéreas, moedas virtuais. 2. ed. São Paulo: Foco, 2021.
- LACERDA, Bruno Torquato Zampier. Estatuto jurídico da inteligência artificial: entre categorias e conceitos, a busca por marcos regulatórios. 1. ed. Indaiatuba, SP: Foco, 2022. E-book. Disponível em: <https://plataforma.bvirtual.com.br>. Acesso em: 19 maio 2025.

LASALVIA, Raquel; MAEJI, Vanessa; CIEGLINSKI, Thaís (Ed.). Com a plataforma Sinapses, Judiciário assume protagonismo no desenvolvimento de soluções de IA. Conselho Nacional de Justiça, 27 jun. 2023. Disponível em: <https://www.cnj.jus.br/com-a-plataforma-sinapses-judiciario-assume-protagonismo-no-desenvolvimento-de-solucoes-de-ia/>. Acesso em: 25 ago. 2025.

LUHMANN, Niklas. Legitimação pelo procedimento. Tradução de Maria da Conceição Côrte-Real. Brasília: Editora Universidade de Brasília, 1980.

MONNERAT, Fábio Victor da Fonte. Introdução ao estudo do Direito Processual Civil. 5. ed. São Paulo: Saraiva Jur, 2020.

MORI, Masahiro. The uncanny valley: the original essay by Masahiro Mori. Tradução de Karl F. MacDorman e Norri Kageki. IEEE Spectrum, 12 jun. 2012. Artigo original: Energy, v. 7, n. 4, p. 33-35, 1970. Disponível em: <https://spectrum.ieee.org/autoton/robotics/humanoids/the-uncanny-valley>. Acesso em: 28 set. 2024.

NUNES, Dierle; LUD, Natanael; PEDRON, Flávio Quinaud. Desconfiando da imparcialidade dos sujeitos processuais: contributo de uma crítica da racionalidade da prova para a superação das premissas idealistas do sistema processual. Salvador: Juspodivm, 2019.

SANDHU, Jamie A. What are LLMs and generative AI?: a beginner's guide to the technology turning heads. Toronto: Schwartz Reisman Institute for Technology and Society, 25 jan. 2024. Disponível em: <https://srinstitute.utoronto.ca/news/gen-ai-llms-explainer>. Acesso em: 23 ago. 2025.

SAYGIN, Ayse Pinar et al. The thing that should not be: predictive coding and the uncanny valley in perceiving human and humanoid robot actions. SCAN (Social Cognitive and Affective Neuroscience), v. 7, n. 4, p. 413-422, 2012. DOI: 10.1093/scan/nsr025. Disponível em: <https://academic.oup.com/scan/article/7/4/413/1738009>. Acesso em: 25 maio 2025.

SUNSTEIN, Cass R.; GAFFET, Jared H. An Anatomy of Algorithm Aversion. Columbia Science & Technology Law Review, v. 26, p. 290-317, 2024.

TYLER, Tom R.; JACKSON, Jonathan. Popular legitimacy and the exercise of legal authority: motivating compliance, cooperation and engagement. Psychology, Public Policy,

and Law, v. 19, n. 3, p. 275-291, 2013. DOI: 10.1037/a0034514. Disponível em: https://www.researchgate.net/publication/249643866_Popular_Legitimacy_and_the_Exercise_of_Legal_Authority_Motivating_Compliance_Cooperation_and_Engagement. Acesso em: 23 jun. 2025.

VALLE, Vivian Lima López; FUENTES i GASÓ, Josep Ramón; AJUS, Atílio Martins. Decisão judicial assistida por inteligência artificial e o Sistema Victor do Supremo Tribunal Federal. Revista de Investigações Constitucionais, Curitiba, vol. 10, n. 2, e252, maio/ago. 2023. DOI: 10.5380/rinc.v10i2.92598.

VALENTINI, Rômulo Soares. Julgamento por computadores?: as novas possibilidades da juscibernética no século XXI e suas implicações para o futuro do direito e do trabalho dos juristas. 2017. 302 f. Tese (Doutorado em Direito) - Faculdade de Direito, Universidade Federal de Minas Gerais, Belo Horizonte, 2017.

WEBER, Max. Economia e sociedade: fundamentos da sociologia compreensiva. 4ª reimpressão. Brasília: Editora Universidade de Brasília, 2015.