

IX ENCONTRO VIRTUAL DO CONPEDI

**DIREITO, GOVERNANÇA E NOVAS TECNOLOGIAS
VI**

Todos os direitos reservados e protegidos. Nenhuma parte destes anais poderá ser reproduzida ou transmitida sejam quais forem os meios empregados sem prévia autorização dos editores.

Diretoria - CONPEDI

Presidente - Profa. Dra. Samyra Haydêe Dal Farra Napolini - FMU - São Paulo

Diretor Executivo - Prof. Dr. Orides Mezzaroba - UFSC - Santa Catarina

Vice-presidente Norte - Prof. Dr. Jean Carlos Dias - Cesupa - Pará

Vice-presidente Centro-Oeste - Prof. Dr. José Querino Tavares Neto - UFG - Goiás

Vice-presidente Sul - Prof. Dr. Leonel Severo Rocha - Unisinos - Rio Grande do Sul

Vice-presidente Sudeste - Profa. Dra. Rosângela Lunardelli Cavallazzi - UFRJ/PUCRio - Rio de Janeiro

Vice-presidente Nordeste - Prof. Dr. Raymundo Juliano Feitosa - UNICAP - Pernambuco

Representante Discente: Prof. Dr. Abner da Silva Jaques - UPM/UNIGRAN - Mato Grosso do Sul

Conselho Fiscal:

Prof. Dr. José Filomeno de Moraes Filho - UFMA - Maranhão

Prof. Dr. Caio Augusto Souza Lara - SKEMA/ESDHC/UFMG - Minas Gerais

Prof. Dr. Valter Moura do Carmo - UFERSA - Rio Grande do Norte

Prof. Dr. Fernando Passos - UNIARA - São Paulo

Prof. Dr. Edinilson Donisete Machado - UNIVEM/UENP - São Paulo

Secretarias

Relações Institucionais:

Prof. Dra. Claudia Maria Barbosa - PUCPR - Paraná

Prof. Dr. Heron José de Santana Gordilho - UFBA - Bahia

Profa. Dra. Daniela Marques de Moraes - UNB - Distrito Federal

Comunicação:

Prof. Dr. Robison Tramontina - UNOESC - Santa Catarina

Prof. Dr. Liton Lanes Pilau Sobrinho - UPF/Univali - Rio Grande do Sul

Prof. Dr. Lucas Gonçalves da Silva - UFS - Sergipe

Relações Internacionais para o Continente Americano:

Prof. Dr. Jerônimo Siqueira Tybusch - UFSM - Rio Grande do Sul

Prof. Dr. Paulo Roberto Barbosa Ramos - UFMA - Maranhão

Prof. Dr. Felipe Chiarello de Souza Pinto - UPM - São Paulo

Relações Internacionais para os demais Continentes:

Profa. Dra. Gina Vidal Marcilio Pompeu - UNIFOR - Ceará

Profa. Dra. Sandra Regina Martini - UNIRITTER / UFRGS - Rio Grande do Sul

Profa. Dra. Maria Claudia da Silva Antunes de Souza - UNIVALI - Santa Catarina

Educação Jurídica

Profa. Dra. Viviane Coêlho de Séllos Knoerr - Unicuritiba - PR

Prof. Dr. Rubens Beçak - USP - SP

Profa. Dra. Livia Gaigher Bosio Campello - UFMS - MS

Eventos:

Prof. Dr. Yuri Nathan da Costa Lannes - FDF - São Paulo

Profa. Dra. Norma Sueli Padilha - UFSC - Santa Catarina

Prof. Dr. Juraci Mourão Lopes Filho - UNICHRISTUS - Ceará

Comissão Especial

Prof. Dr. João Marcelo de Lima Assafim - UFRJ - RJ

Profa. Dra. Maria Creusa De Araújo Borges - UFPB - PB

Prof. Dr. Antônio Carlos Diniz Murta - Fumec - MG

Prof. Dr. Rogério Borba - UNIFACVEST - SC

D598

Direito, governança e novas tecnologias VI [Recurso eletrônico on-line] organização CONPEDI

Coordenadores: Jonathan Barros Vita; Patricia Eliane da Rosa Sardeto; Thiago De Mello Azevedo Guilherme – Florianópolis: CONPEDI, 2026.

Inclui bibliografia

ISBN: 978-65-5274-587-3

Modo de acesso: www.conpedi.org.br em publicações

Tema: Constituição, Processo e Justiça Social: o papel da jurisdição constitucional na efetivação de direitos

2. Direito. 3. Governança e novas tecnologias. IX Encontro Virtual do CONPEDI (2: 2026: Florianópolis, Brasil).

CDU: 34



Conselho Nacional de Pesquisa
e Pós-Graduação em Direito Florianópolis
Santa Catarina – Brasil
www.conpedi.org.br

IX ENCONTRO VIRTUAL DO CONPEDI

DIREITO, GOVERNANÇA E NOVAS TECNOLOGIAS VI

Apresentação

O IX Encontro Virtual do CONPEDI foi realizado entre os dias 22 a 26 de junho de 2026 e teve como temática central “Constituição, processo e justiça social: o papel da jurisdição constitucional na efetivação de direitos”.

No plano das diversas atividades acadêmicas ocorridas neste encontro, destacam-se, além das palestras e oficinas, os grupos de trabalho temáticos, os quais representam um locus de interação entre pesquisadores que apresentam as suas pesquisas temáticas, seguindo-se de debates.

Especificamente, para operacionalizar tal modelo, os coordenadores dos GTs são os responsáveis pela organização dos trabalhos em blocos temáticos, dando coerência à produção e estabelecendo um fio condutor para organizar os debates em subtemas.

No caso concreto, assim aconteceu com o GT Direito, Governança e Novas Tecnologias VI, o qual ocorreu no dia 26 de junho de 2026 das 14h00 às 17h30 e foi coordenado pelos professores Jonathan Barros Vita, Thiago de Mello Azevedo Guilherme e Patricia Eliane da Rosa Sardeto.

O referido GT foi palco de profícuas discussões decorrentes dos trabalhos apresentados, os quais são publicados na presente obra, a qual foi organizada seguindo alguns blocos temáticos específicos, que compreenderam os 21 artigos submetidos ao GT, cujos temas são citados abaixo:

2. O declínio do senso-crítico na era digital: a erosão da autonomia e a superficialidade na sociedade do espetáculo.

Autor(es): Juliana Sales Pavini / Victor Matheus de Campos Leite Neves

3. Sistemas autônomos, vigilância algorítmica, cidadania e os limites do controle humano na governança contemporânea.

Autor(es): Rossana Gemeli Roncato Carloto / Marcio Renan Hamel

4. Transformações legislativas no enfrentamento da desinformação política: uma análise comparativa sobre projetos de lei propostos entre 2018 e 2026.

Autor(es): Bruna Bastos / Leticia Kirinus Danski / Isabela Quartieri da Rosa

5. Proteção integral, opacidade algorítmica e os limites do ECA Digital na governança de plataformas digitais.

Autor(es): Isabela Quartieri da Rosa / Ana Carolina Sassi / Rafael Santos de Oliveira

6. O constitucionalismo digital como limite ao poder privado das Big Techs nas relações digitais.

Autor(es): Samira Bonfim Bernardi / Roberto da Freiria Estevão

7. A possibilidade do desenvolvimento da ação comunicativa entre adolescentes na era digital.

Autor(es): Samira Bonfim Bernardi

9. Governança algorítmica no setor público, participação social e seus impactos jurídicos.

Autor(es): Amadeu dos Anjos Vidonho Junior

10. Inteligência Artificial de Alto Risco e Assurance Contínua: um estudo dogmático à luz do AI Act e do NIST AI RMF.

Autor(es): Daniel David Guimarães Freire / Karine Queiroz de Souza

11. Web agêntica, microtrabalhos e responsabilização jurídica: rastreabilidade e proteção de direitos fundamentais.

Autor(es): Sebastiana Sena Carvalho / Cibele Alexandre Uchoa / João Araújo Monteiro Neto

12. Quality Assurance em Sistemas Automatizados de Compliance e AML: limites jurídicos da confiança em decisões algorítmicas.

Autor(es): Daniel David Guimarães Freire / Karine Queiroz de Souza

13. Governança de IA Agêntica: uma análise do Framework AgentSafe e os desafios de compliance na Web Agêntica.

Autor(es): Gustavo Amaral Lima / Cibele Alexandre Uchoa / João Araújo Monteiro Neto

14. A accountability algorítmica na atividade administrativa do Poder Judiciário brasileiro.

Autor(es): Mateus Marinho Alencar / Cassius Guimaraes Chai

Bloco 3 – Inteligência Artificial, Profissões, Educação e Instituições

16. Educação jurídica em transformação: entre a formação crítica e os desafios do uso da Inteligência Artificial.

Autor(es): Denise Almeida De Andrade / Beatriz de Castro Rosa / Rafaela Mota Holanda

17. Responsabilidade civil do advogado pelos atos praticados com a Inteligência Artificial Generativa no processo judicial eletrônico brasileiro.

Autor(es): Amadeu dos Anjos Vidonho Junior

18. Sala de aula em movimento, uma experiência com o ensino médio: juventude e redes sociais, entre a liberdade de acesso e os riscos da exposição íntima.

Autor(es): Dayana Claudia Tavares Barros De Castro / Aimme Beatrice de Oliveira Dutra Cordeiro / Denise Almeida De Andrade

19. Da imprescindibilidade da subjetividade humana na tessitura do processo decisório: por uma perspectiva ética na utilização da Inteligência Artificial.

Autor(es): Lanaira da Silva / Rayssa Gargano Soares

20. A reconfiguração do standard of care médico na era algorítmica: o dever de diligência tecnológica e os limites da confiança nas decisões automatizadas na saúde.

Autor(es): Sara Lupion dos Santos / Patricia Eliane da Rosa Sardeto

21. Inteligência Artificial e qualificação registral: limites dogmáticos e possibilidades de coexistência no registro de imóveis.

Esse é o contexto que permite a promoção e o incentivo da cultura jurídica no Brasil, consolidando o CONPEDI, cada vez mais, como um importante espaço para discussão e apresentação das pesquisas desenvolvidas nos ambientes acadêmicos da graduação e pós-graduação em direito.

Finalmente, desejamos uma boa leitura, deste Grupo de trabalho que reuniu diversos textos e autores de todo o Brasil para servir como resultado de pesquisas científicas realizadas no âmbito dos cursos de Pós-Graduação Stricto Sensu de nosso país.

Prof. Dr. Jonathan Barros Vita – Unimar

Prof. Dr. Thiago De Mello Azevedo Guilherme - ITE-CEUB

Profa. Dra. Patricia Eliane da Rosa Sardeto – Pontifícia Universidade Católica do Paraná – Câmpus Londrina

INTELIGÊNCIA ARTIFICIAL DE ALTO RISCO E ASSURANCE CONTÍNUA - UM ESTUDO DOGMÁTICO À LUZ DO AI ACT E DO NIST AI RMF

HIGH-RISK ARTIFICIAL INTELLIGENCE AND CONTINUOUS ASSURANCE: A DOGMATIC STUDY IN LIGHT OF THE AI ACT AND THE NIST AI RMF

**Daniel David Guimarães Freire
Karine Queiroz de Souza**

Resumo

O presente artigo analisa os limites dos modelos tradicionais de auditoria e governança aplicados a sistemas de inteligência artificial de alto risco, sustentando que mecanismos baseados em certificações pontuais, auditorias episódicas e compliance formal são insuficientes para assegurar accountability, legitimidade regulatória e proteção efetiva de direitos fundamentais em sistemas adaptativos e não determinísticos. A pesquisa parte da hipótese de que a crise contemporânea da auditabilidade algorítmica decorre da incompatibilidade entre instrumentos clássicos de validação e a natureza dinâmica da inteligência artificial contemporânea. A partir de revisão bibliográfica e análise normativa comparada, o trabalho examina o AI Act da União Europeia e o AI Risk Management Framework do NIST, investigando como esses modelos estruturam mecanismos contínuos de monitoramento, gestão de riscos e supervisão pós-implantação. O artigo também discute os impactos da concentração infraestrutural, do colonialismo de dados, da captura regulatória e da fragmentação da responsabilidade algorítmica sobre os modelos atuais de governança. Como contribuição teórica, propõe a assurance contínua como categoria jurídico-regulatória voltada à produção permanente de evidências de conformidade, rastreabilidade, contestabilidade e supervisão institucional ao longo do ciclo de vida de sistemas de inteligência artificial de alto risco.

Palavras-chave: Inteligência artificial, Assurance contínua, Governança algorítmica, Accountability, Auditabilidade algorítmica, Ai act

Abstract/Resumen/Résumé

these models structure continuous mechanisms of monitoring, risk management, and post-deployment oversight. The paper also addresses the impacts of infrastructural concentration, data colonialism, regulatory capture, and fragmented algorithmic responsibility on current governance frameworks. As a theoretical contribution, it proposes continuous assurance as a legal-regulatory category aimed at the continuous production of evidence regarding compliance, traceability, contestability, and institutional oversight throughout the lifecycle of high-risk artificial intelligence systems.

Keywords/Palabras-claves/Mots-clés: Artificial intelligence, Continuous assurance, Algorithmic governance, Accountability, Algorithmic auditability, Ai act

1. INTRODUÇÃO

A rápida expansão da inteligência artificial tem promovido uma transformação estrutural nos modelos contemporâneos de tomada de decisão automatizada. Sistemas baseados em machine learning e inteligência artificial generativa passaram a ocupar posição central em setores regulados e de alta sensibilidade, como o financeiro, o setor público e a saúde. Embora ampliem significativamente a eficiência operacional, esses sistemas também introduzem riscos relevantes relacionados à transparência, previsibilidade, responsabilização e proteção de direitos fundamentais, especialmente quando produzem inferências, classificam perfis ou executam decisões automatizadas (NIST, 2023).

Esse cenário tensiona diretamente os modelos tradicionais de governança e auditoria tecnológica. Historicamente, a validação de sistemas computacionais foi estruturada sobre pressupostos de estabilidade, previsibilidade e verificabilidade periódica da conformidade. No entanto, sistemas contemporâneos de inteligência artificial, especialmente aqueles classificados como de alto risco, operam por meio de inferência probabilística, aprendizado contínuo e comportamento não determinístico, o que limita a possibilidade de uma validação definitiva. A dificuldade de explicar integralmente a lógica decisória desses sistemas intensifica a crise da auditabilidade algorítmica, marcada pela insuficiência dos mecanismos clássicos de controle para assegurar rastreabilidade, confiabilidade e conformidade jurídica ao longo do tempo (MOKANDER et al., 2021).

Diante desse contexto, emerge a seguinte problemática: em que medida os modelos tradicionais de auditoria e governança são capazes de assegurar legitimidade regulatória, accountability e proteção de direitos fundamentais em sistemas de inteligência artificial de alto risco que se modificam continuamente ao longo de seu ciclo de vida?

No plano regulatório, observa-se um deslocamento relevante na forma de enfrentar esse desafio. O AI Act da União Europeia estabelece, para sistemas de alto risco, a obrigatoriedade de estruturas permanentes de gestão da qualidade, supervisão humana, documentação técnica e monitoramento contínuo pós-implantação (EUROPEAN UNION, 2024). De forma complementar, o AI Risk Management Framework do NIST propõe uma abordagem baseada em governança, mensuração, gestão e monitoramento contínuo de riscos ao longo de todo o ciclo de vida dos sistemas (NIST, 2023). Em conjunto, esses referenciais evidenciam a transição de modelos baseados em validações pontuais para estruturas de controle contínuo.

A hipótese desenvolvida neste artigo sustenta que a crise contemporânea da auditabilidade algorítmica decorre da incompatibilidade entre mecanismos tradicionais de validação e a natureza adaptativa dos sistemas de inteligência artificial. Argumenta-se que formas pontuais de auditoria são insuficientes para assegurar, de maneira contínua, confiabilidade, rastreabilidade e conformidade regulatória, especialmente em aplicações classificadas como de alto risco.

Nesse contexto, o artigo tem como objetivo analisar de que forma modelos de *continuous AI assurance*¹ reconfiguram os limites jurídicos da confiança em sistemas de inteligência artificial. A *assurance* contínua é compreendida não apenas como uma prática técnica de monitoramento, mas como um mecanismo jurídico-regulatório de produção contínua de evidências de conformidade. A análise concentra-se nos impactos dessa transformação sobre accountability, governança algorítmica, proteção de dados pessoais e legitimidade regulatória.

O recorte do artigo concentra-se especificamente nos desafios jurídicos associados à validação contínua de sistemas de inteligência artificial de alto risco, sem examinar a inteligência artificial em sentido amplo ou reduzir a análise à proteção de dados pessoais de forma isolada. Sua contribuição reside na formulação da *assurance* contínua como categoria jurídico-regulatória capaz de reinterpretar os mecanismos tradicionais de controle à luz de sistemas dinâmicos, adaptativos e não determinísticos.

Para tanto, adota-se o método dogmático-jurídico, com revisão bibliográfica e análise normativa comparada, tomando como referenciais centrais o AI Act da União Europeia e o AI Risk Management Framework do NIST, a fim de investigar como diferentes modelos regulatórios vêm estruturando mecanismos contínuos de controle e validação em sistemas de inteligência artificial.

2. RECONFIGURAÇÃO DA REGULAÇÃO NA ERA DA GOVERNANÇA ALGORÍTMICA

O problema regulatório da inteligência artificial não pode ser reduzido ao controle pontual de aplicações “de alto risco”, porque a própria forma contemporânea da regulação passou a operar em ambiente atravessado por sistemas preditivos, arquiteturas de dados e infraestruturas privadas de classificação. A categoria de governança algorítmica designa, por isso, algo mais amplo do que

¹ *Continuous AI assurance* é compreendida como processo contínuo de produção, verificação e atualização de evidências relacionadas à conformidade, gestão de riscos, rastreabilidade e accountability de sistemas de inteligência artificial ao longo de todo o seu ciclo de vida, em oposição a modelos baseados exclusivamente em auditorias pontuais ou validações estáticas (NIST, 2023; OECD, 2023).

a presença de modelos computacionais na administração pública ou no mercado: trata-se de uma mutação no modo de produção da normatividade prática, em que técnicas de correlação e antecipação passam a disputar com as formas jurídicas tradicionais a definição do que deve ser visto, priorizado, bloqueado, recomendado ou sancionado. Nessa chave, a “algorithmic regulation” nomeada (YEUNG, 2018) não é apenas um instrumento de eficiência administrativa, mas um deslocamento do eixo da decisão para dispositivos que atuam pela contínua modulação de condutas. Em paralelo, a “algorithmic governmentality” explicita que o governo por dados tende a operar sem mediação pública suficiente, apostando menos na formação do sujeito de direito do que na gestão probabilística de perfis, tendências e propensões (ROUVROY, 2013). Nesse contexto, quando o mundo social é convertido em material legível por máquina, o desafio deixa de ser apenas técnico e passa a incidir sobre as próprias condições de inteligibilidade e contestação próprias do Estado de Direito (HILDEBRANDT, 2015).

Os modelos regulatórios hoje dominantes respondem a essa mutação com uma gramática centrada na gestão de riscos, na documentação procedimental e na rastreabilidade formal. O Regulamento (UE) 2024/1689, ao instituir regras harmonizadas para a inteligência artificial, consagra uma arquitetura baseada em categorias de risco; de modo distinto, mas convergente quanto à racionalidade administrativa, o AI RMF 1.0 do NIST oferece um repertório voluntário de governança voltado à gestão organizacional de riscos e à promoção de atributos de “trustworthiness”. Em ambos os casos, a ênfase recai sobre processos de identificação, mitigação e monitoramento de danos potenciais. O problema é que esse enquadramento, embora relevante, tende a traduzir conflitos estruturais em déficits operacionais administráveis e, com isso, obscurece a crítica da economia política da IA. A “tradição” regulatória já havia demonstrado que a descentralização regulatória frequentemente desloca poder normativo para atores híbridos, técnicos e privados, sob a justificativa de complexidade, expertise e flexibilidade (SCOTT, 2004). Na IA, essa descentralização é ainda mais aguda, porque o objeto regulado e os meios técnicos para compreendê-lo pertencem, em larga medida, aos próprios agentes econômicos cuja atividade se pretende disciplinar.

Por isso, a questão central não reside na ausência absoluta de normas, mas na assimetria de autoridade que orienta sua elaboração, interpretação e execução. A economia informacional não apenas explora lacunas jurídicas preexistentes, mas reorganiza as próprias estruturas institucionais em torno de novas formas de mediação, extração de valor e concentração de poder (COHEN, 2019).

Nesse contexto, a plataforma digital assume posição central no capitalismo informacional, convertendo fluxos de dados em mecanismos de mercado, hierarquização e dependência estrutural. A expansão dessas infraestruturas privadas para além do âmbito econômico evidencia sua capacidade de penetrar práticas sociais, relações cívicas e esferas de interesse público, influenciando valores coletivos a partir de arquiteturas proprietárias e modelos opacos de governança (VAN DIJCK; POELL; DE WAAL, 2018). Paralelamente, a monetização do excedente comportamental consolida uma racionalidade empresarial fundada não apenas na observação, mas na predição e modulação do comportamento humano (ZUBOFF, 2019). Como consequência, emerge uma forma difusa de regulação social em que atores privados passam a desempenhar funções tradicionalmente associadas ao direito público, embora sem submissão equivalente a mecanismos democráticos de controle, transparência e accountability.

Nesse ambiente, a responsabilidade algorítmica não pode ser compreendida como mero dever posterior de justificar resultados automatizados após a ocorrência do dano. A própria noção de accountability revela limitações quando transplantada para sistemas algorítmicos complexos, sobretudo se permanecer vinculada a modelos lineares de autoria, decisão e causalidade. A combinação entre opacidade técnica, autonomia funcional aparente e multiplicidade de agentes distribuídos ao longo do ciclo de vida dos sistemas dificulta a identificação precisa de centros decisórios e de responsabilização efetiva (BUSUIOC, 2021). Essa problemática torna-se ainda mais evidente quando se abandona a concepção abstrata de sistemas autônomos isolados e se observa a estrutura material de produção contemporânea da inteligência artificial. Modelos algorítmicos são desenvolvidos a partir de ecossistemas fragmentados que envolvem bibliotecas preexistentes, serviços em nuvem, datasets terceirizados, APIs, módulos integrados e sucessivas camadas de desenvolvimento realizadas por diferentes agentes econômicos e técnicos (WIDDER; NAFUS, 2023). Como consequência, a responsabilidade tende a ser continuamente deslocada entre os diversos elos da cadeia produtiva, produzindo estruturas de accountability fragmentadas nas quais nenhum ator detém, simultaneamente, controle, supervisão e conhecimento suficientes para responder integralmente pelos efeitos sistêmicos produzidos. A pulverização da responsabilidade torna-se, assim, funcionalmente conveniente para organizações que operam em ambientes de alta complexidade tecnológica e baixa transparência institucional.

Esse cenário também evidencia, sob perspectiva específica, os riscos de captura regulatória associados à disciplina jurídica da inteligência artificial. A captura não depende necessariamente

de formas clássicas de corrupção ou influência indevida explícita, podendo manifestar-se por meio da dependência estrutural do regulador em relação à expertise técnica, aos padrões normativos e às métricas produzidas pelos próprios setores regulados. Problemas relacionados ao enforcement do AI Act europeu já foram apontados em razão do peso conferido a standards harmonizados e aos efeitos decorrentes da harmonização regulatória em larga escala (VEALE; ZUIDERVEEN BORGESIU, 2021). Em complemento, também se observa que a governança prevista no AI Act permanece marcada por interação incompleta entre mecanismos públicos e privados de regulação, especialmente diante da limitada participação da sociedade civil nos processos de elaboração de padrões técnicos, códigos de conduta e instrumentos de conformidade (MICKLITZ, 2025). A questão central, portanto, não reside apenas na existência formal de regulação, mas na definição concreta de quem possui legitimidade prática para estabelecer os conteúdos técnicos considerados juridicamente válidos. Quando presunções de conformidade passam a depender de standards elaborados em ambientes técnicos pouco acessíveis e fortemente ocupados por agentes econômicos incumbentes, a promessa de neutralidade regulatória converte-se em mecanismo potencial de reprodução e estabilização de assimetrias de poder.

3. COLONIALISMO DE DADOS, INFRAESTRUTURA E DEPENDÊNCIA TECNOLÓGICA

A insuficiência das respostas normativas contemporâneas torna-se ainda mais evidente quando a inteligência artificial é situada em processos amplos de extração de dados e reorganização infraestrutural da vida social. A economia de dados pode ser compreendida como dinâmica estrutural de apropriação que converte interações humanas, comportamentos e experiências cotidianas em matéria-prima para acumulação econômica (COULDRY; MEJIAS, 2019). A coleta massiva de dados deixa, assim, de representar simples etapa preparatória da inovação tecnológica e passa a constituir condição material indispensável para a própria operacionalidade dos sistemas algorítmicos. Paralelamente, consolida-se uma racionalidade jurídico-econômica que naturaliza a percepção dos dados pessoais como recursos permanentemente disponíveis para exploração e processamento comercial (COHEN, 2019). Nesse contexto, a transformação da experiência humana em insumo previsional e comercializável torna-se elemento central de modelos econômicos fundados em vigilância e modulação comportamental (ZUBOFF, 2019). O problema

da governança algorítmica, portanto, não se limita ao funcionamento interno do algoritmo, mas envolve a infraestrutura extrativa que sustenta sua existência.

Essa dinâmica depende de ecossistemas altamente concentrados compostos por plataformas digitais, serviços de computação em nuvem, data centers e cadeias globais de infraestrutura tecnológica. A dataficação das relações sociais opera simultaneamente como paradigma técnico e racionalidade política, naturalizando a conversão da vida social em dados economicamente exploráveis (VAN DIJCK, 2014). As plataformas deixam, assim, de atuar apenas como intermediárias ocasionais e passam a exercer funções infraestruturais centrais na organização de fluxos econômicos, sociais e institucionais (VAN DIJCK; POELL; DE WAAL, 2018). Mesmo sistemas apresentados como “abertos” permanecem dependentes de recursos computacionais, infraestrutura e capacidade de treinamento controlados por poucos atores corporativos, ampliando relações de dependência tecnológica e concentração de poder (WIDDER; WHITTAKER; WEST, 2024). Nessas condições, a soberania regulatória revela-se parcialmente limitada, uma vez que o direito frequentemente alcança usos periféricos da inteligência artificial sem incidir, com igual intensidade, sobre os centros infraestruturais onde se concentram capacidade computacional, segredo industrial e poder econômico.

Esse cenário também evidencia as limitações do individualismo jurídico ainda predominante na proteção de dados pessoais. A economia dos dados opera prioritariamente por meio de inferências relacionais e classificações coletivas capazes de modular riscos, mercados e comportamentos em escala populacional (VILJOEN, 2021). Como consequência, direitos individuais de acesso, consentimento ou oposição permanecem relevantes, mas insuficientes para enfrentar efeitos distributivos produzidos por sistemas algorítmicos. O dano não se restringe ao erro direcionado contra indivíduos identificáveis, abrangendo também a formação de ambientes classificatórios que redistribuem oportunidades, vigilância e vulnerabilidade entre grupos sociais e populações inteiras. A permanência de uma racionalidade regulatória excessivamente centrada no indivíduo favorece, assim, a continuidade de formas sofisticadas de dominação infraestrutural sob a aparência formal de autodeterminação informativa.

4. AUTOMAÇÃO DECISÓRIA, DESIGUALDADE ESTRUTURAL E LIMITES COGNITIVOS DA EXPLICABILIDADE

A dimensão estrutural da desigualdade algorítmica torna-se particularmente evidente nos campos em que a automação decisória foi incorporada de forma mais intensa. Em políticas de assistência social, moradia, crédito, educação, policiamento e relações de trabalho, sistemas automatizados frequentemente deixam de atuar apenas como instrumentos de eficiência administrativa e passam a operar como mecanismos de ampliação da seletividade e da exclusão social (EUBANKS, 2018; O'NEIL, 2016). Processos de busca, classificação e ranqueamento também tendem a reproduzir hierarquias raciais, econômicas e de gênero já existentes, sobretudo quando estruturados por lógicas comerciais e métricas de visibilidade orientadas à maximização de engajamento e lucro (NOBLE, 2018). Nesse contexto, a retórica da neutralidade tecnológica frequentemente encobre a atualização de práticas discriminatórias sob linguagens aparentemente objetivas de inovação, eficiência e racionalidade técnica (BENJAMIN, 2019). A inteligência artificial não apenas reflete desigualdades previamente existentes, mas pode estabilizá-las, aprofundá-las e revesti-las de legitimidade técnica.

Esse diagnóstico impede que a injustiça algorítmica seja tratada como simples falha contingente solucionável por ajustes incrementais de engenharia. O problema ultrapassa a presença de vieses em bases de dados e alcança os próprios critérios institucionais que determinam quais populações serão continuamente monitoradas, avaliadas e submetidas a processos intensivos de classificação. Em diferentes contextos, a automação decisória articula-se a políticas de austeridade, administração da escassez, monetização da visibilidade e mecanismos circulares de classificação e exclusão social (EUBANKS, 2018; NOBLE, 2018; O'NEIL, 2016). Por essa razão, a discussão sobre igualdade algorítmica não pode limitar-se à calibragem de métricas de fairness ou à ampliação de bases de treinamento. A questão envolve os critérios institucionais que definem quem será tornado legível pelos sistemas, em quais condições essa legibilidade ocorrerá e quais consequências materiais decorrerão dessa classificação. Enquanto essa dimensão estrutural permanecer inalterada, soluções estritamente técnicas tenderão a funcionar como mecanismos de despolitização de conflitos distributivos e desigualdades sociais.

Esses elementos também demonstram por que a explicabilidade técnica, embora juridicamente relevante em determinados contextos, não pode ocupar posição exclusiva ou central na crítica à governança algorítmica. A opacidade dos sistemas de aprendizagem de máquina decorre não apenas de segredo empresarial deliberado, mas também da complexidade técnica e da distância existente entre modelos matemáticos e formas humanas de interpretação (BURRELL,

2016). Além disso, a transparência informacional não garante, por si só, compreensão efetiva das relações institucionais, interesses econômicos e estruturas de poder que condicionam o funcionamento dos sistemas algorítmicos (ANANNY; CRAWFORD, 2018). A própria ideia de explicação algorítmica revela limites epistemológicos importantes, uma vez que sistemas computacionais produzem apenas representações parciais da realidade e não encerram, mediante abertura técnica do código ou do modelo, os problemas éticos e políticos relacionados à decisão automatizada (AMOORE, 2020). A juridicidade da inteligência artificial depende, portanto, não apenas de transparência técnica, mas da preservação de condições efetivas de contestação, argumentação e reciprocidade institucional próprias da forma jurídica (HILDEBRANDT, 2015). O limite da explicabilidade não constitui simples dificuldade metodológica; ele evidencia que o direito não pode delegar integralmente a inteligibilidade do poder a uma racionalidade exclusivamente computacional.

5. LIMITES DO PARADIGMA REGULATÓRIO CONTEMPORÂNEO

O problema central dos modelos contemporâneos de governança da inteligência artificial reside no fato de que grande parte das respostas normativas busca administrar efeitos sem alterar de maneira substancial as estruturas que produzem e concentram poder algorítmico. Regimes regulatórios baseados em risco, princípios éticos e mecanismos de compliance introduzem obrigações relevantes, mas frequentemente permanecem inseridos na própria racionalidade que pretendem limitar, operando por meio da classificação, gerenciamento e mitigação de danos sem questionar adequadamente quem controla infraestrutura computacional, fluxos de dados, padrões técnicos e mercados digitais. O AI Act europeu representa avanço importante ao estabelecer obrigações jurídicas para determinadas categorias de sistemas, mas persistem preocupações relacionadas à efetividade de mecanismos de enforcement e ao peso atribuído a standards harmonizados na operacionalização regulatória (VEALE; ZUIDERVEEN BORGESIU, 2021). Também se observam limitações na articulação entre governança pública e privada, especialmente quanto à participação da sociedade civil nos processos de elaboração de padrões técnicos, códigos de conduta e instrumentos de conformidade (MICKLITZ, 2025). Paralelamente, parte significativa do discurso global de “AI ethics” continua apresentando baixa capacidade de fortalecimento de accountability democrática e participação substantiva (FUKUDA-PARR; GIBBONS, 2021). Soma-se a isso a concentração infraestrutural da inteligência artificial, inclusive em modelos

apresentados como “abertos”, dependentes de recursos computacionais e ecossistemas controlados por poucos atores corporativos (WIDDER; WHITTAKER; WEST, 2024). O resultado é um cenário em que a produção normativa avança mais rapidamente na formulação de obrigações periféricas do que na limitação efetiva dos centros privados de poder tecnológico.

Há, ainda, um descompasso estrutural entre a velocidade de expansão das infraestruturas algorítmicas e a capacidade institucional de produzir supervisão pública substantiva. A regulação estatal permanece predominantemente territorial, setorial e reativa, enquanto sistemas de inteligência artificial operam de forma transnacional, modular e articulada por contratos, APIs, serviços em nuvem e componentes reutilizáveis. A figura tradicional do regulador nacional enfrenta, nesse contexto, crescente dificuldade para compreender e supervisionar ecossistemas cuja inteligibilidade depende de informações distribuídas de maneira profundamente assimétrica (SCOTT, 2004). O desafio regulatório deixa de consistir apenas na limitação da atividade econômica e passa a envolver a própria regulação de estruturas privadas de autorregulação tecnológica (BLACK, 2001). No campo algorítmico, essa dificuldade se intensifica porque reguladores frequentemente dependem de documentação, métricas e categorias produzidas pelos próprios agentes regulados para delimitar o objeto de supervisão. A modularidade das cadeias de produção de IA fragmenta noções tradicionais de responsabilidade e agência (WIDDER; NAFUS, 2023), ao mesmo tempo em que as próprias instituições jurídicas passam a ser progressivamente reorganizadas pelas dinâmicas do capitalismo informacional (COHEN, 2019). A fragilidade institucional contemporânea não decorre da inexistência formal de normas ou autoridades regulatórias, mas da dificuldade material de romper a dependência epistêmica, técnica e operacional em relação às infraestruturas privadas que se pretende supervisionar.

Por essa razão, a crítica à governança contemporânea da inteligência artificial não pode permanecer restrita à oposição simplificada entre inovação e regulação. O déficit normativo atual não resulta apenas de lacunas legislativas, mas da permanência de respostas jurídicas profundamente inseridas na lógica do capitalismo informacional, convertendo problemas de dominação estrutural em questões predominantemente administrativas de gerenciamento e conformidade. Quando a extração massiva de dados se naturaliza como condição ordinária da vida conectada, quando a infraestrutura computacional se concentra em poucos conglomerados globais e quando padrões técnicos passam a materializar escolhas políticas sob aparência de neutralidade, a governança da IA corre o risco de transformar-se em mera administração jurídica da assimetria.

Nesse cenário, mecanismos de verificação, avaliação e conformidade institucional somente poderão desempenhar função efetivamente crítica se não perderem de vista as estruturas de concentração de poder, dependência tecnológica, captura regulatória e desigualdade que sustentam a economia política contemporânea da inteligência artificial.

6. CONTINUOUS AI ASSURANCE COMO CATEGORIA JURÍDICO-REGULATÓRIA

A crítica aos modelos contemporâneos de governança da inteligência artificial não conduz à rejeição dos instrumentos de verificação, avaliação e conformidade, mas à necessidade de sua reconstrução conceitual. Se os capítulos anteriores demonstraram que a inteligência artificial de alto risco opera em ambiente marcado por opacidade técnica, dependência infraestrutural, assimetria regulatória e deslocamento de responsabilidades, torna-se insuficiente compreender a auditoria como ato episódico de certificação ou como procedimento retrospectivo de apuração de falhas. O problema central não é apenas verificar se um sistema foi concebido em conformidade com determinados requisitos normativos, mas assegurar que sua operação permaneça juridicamente controlável ao longo do tempo.

É nesse ponto que a *continuous AI assurance* adquire relevância como categoria jurídico-regulatória. Sua função não é apenas acompanhar o desempenho técnico de sistemas de IA, mas organizar a produção contínua de evidências capazes de sustentar conformidade, accountability e contestabilidade. A confiança regulatória deixa de depender exclusivamente de uma validação inicial e passa a exigir demonstrações atualizadas de que os riscos permanecem identificados, mensurados, mitigados e submetidos a mecanismos efetivos de revisão.

A noção de assurance contínua não deve ser confundida com monitoramento técnico em sentido estrito. Monitorar um sistema significa observar seu desempenho, registrar eventos e identificar desvios operacionais. A assurance contínua pressupõe uma camada normativa adicional: a conversão dessas observações em evidências verificáveis, juridicamente relevantes e aptas a sustentar controle por reguladores, auditores, órgãos de supervisão, usuários afetados e demais partes interessadas. Trata-se, portanto, de um mecanismo de mediação entre a mutabilidade técnica dos sistemas de IA e a exigência jurídica de rastreabilidade, responsabilidade e controle institucional.

Essa compreensão dialoga diretamente com a arquitetura do AI Act. O Regulamento (UE) 2024/1689 estrutura, para sistemas de alto risco, um conjunto de obrigações que inclui sistema de

gestão de riscos, governança de dados, documentação técnica, registros automáticos, transparência, supervisão humana, acurácia, robustez, cibersegurança e monitoramento pós-comercialização. A documentação técnica, nesse contexto, não deve ser compreendida como mero arquivo burocrático produzido antes da entrada do sistema no mercado, mas como instrumento de demonstração da conformidade e de inteligibilidade regulatória. Do mesmo modo, a avaliação de conformidade não encerra a vida jurídica do sistema: modificações substanciais, mudanças de finalidade, degradação de desempenho ou identificação de riscos relevantes podem exigir nova apreciação. O monitoramento pós-mercado reforça essa lógica ao impor ao provedor o dever de coletar, documentar e analisar dados sobre o desempenho do sistema ao longo de seu ciclo de vida, inclusive para permitir ações corretivas tempestivas (EUROPEAN UNION, 2024).

O NIST AI RMF converge com essa racionalidade, embora por via não vinculante e mais flexível. Ao organizar a gestão de riscos de IA em torno das funções Govern, Map, Measure e Manage, o framework propõe um ciclo reiterado de governança, contextualização, mensuração e tratamento de riscos. A assurance contínua pode ser compreendida, nesse sentido, como a tradução jurídico-operacional desse ciclo: não basta que a organização declare possuir governança de IA; ela deve produzir evidências de que governa, mapeia, mede e gerencia riscos de forma persistente, proporcional e auditável (NIST, 2023).

A contribuição específica da assurance contínua está no regime temporal e probatório da conformidade. A auditoria tradicional opera, em grande medida, por amostragem, periodicidade e reconstrução retrospectiva, examinando documentos, processos e resultados em determinado intervalo. Esse modelo permanece relevante, sobretudo para avaliações independentes e controles externos, mas revela limitações diante de sistemas sujeitos a drift de dados, alterações de contexto, atualizações de modelo, mudanças de comportamento dos usuários, integração com novos componentes e exposição a riscos adversariais. O compliance formal, por sua vez, tende a demonstrar aderência procedimental a obrigações internas ou externas, frequentemente por meio de políticas, checklists, treinamentos, aprovações e controles organizacionais. Embora tais elementos sejam indispensáveis, eles não bastam para sistemas de IA de alto risco, porque a conformidade relevante não se esgota na existência de procedimentos. É necessário demonstrar que esses procedimentos permanecem efetivos diante do comportamento real do sistema.

A assurance contínua exige, portanto, a passagem do compliance documental para o compliance evidencial. A questão jurídica deixa de consistir apenas na existência formal de

políticas de governança e passa a envolver a demonstração contínua de evidências capazes de indicar que os riscos permanecem controlados ao longo do ciclo de vida do sistema. Nesse contexto, avaliações de impacto, avaliações de risco, testes pré-implantação e procedimentos de conformidade não desaparecem, mas passam a integrar um processo iterativo de monitoramento, mensuração e reavaliação contínua, compatível com abordagens regulatórias orientadas ao ciclo de vida da IA e à gestão permanente de riscos (NIST, 2023; OECD, 2023).

Esse modelo pressupõe que a conformidade não opera como autorização definitiva para funcionamento em condições substancialmente distintas daquelas originalmente avaliadas. Alterações relevantes no modelo, mudanças em bases de dados, expansão de finalidade, integração com novos sistemas, degradação de desempenho, identificação de vieses, incidentes de segurança, aumento de falsos positivos ou mudanças regulatórias podem exigir nova avaliação das condições de legitimidade operacional do sistema. A gestão de riscos em IA, nessa perspectiva, deixa de possuir caráter estático e passa a depender de mecanismos permanentes de monitoramento, rastreabilidade e revisão institucional ao longo de todo o ciclo de vida do sistema (NIST, 2023; OECD, 2023).

A operacionalização da assurance contínua exige, em primeiro lugar, a delimitação do próprio objeto da evidência regulatória. Em sistemas de inteligência artificial classificados como de alto risco, não basta registrar o resultado final produzido pelo modelo ou documentar métricas isoladas de desempenho. A evidência de conformidade precisa abranger o ciclo de vida do sistema de forma mais ampla, incluindo finalidade pretendida, classificação de risco, fundamentos jurídicos aplicáveis, origem e qualidade dos dados utilizados, critérios de treinamento e validação, limites conhecidos do modelo, grupos potencialmente afetados, mecanismos de supervisão humana, registros operacionais, incidentes, atualizações e decisões de mitigação adotadas ao longo do funcionamento do sistema. Essa estrutura de rastreabilidade busca responder à fragmentação contemporânea da responsabilidade em ecossistemas algorítmicos compostos por múltiplos atores, componentes e etapas de desenvolvimento, permitindo reconstruir processos decisórios, identificar critérios utilizados e verificar como determinados riscos foram avaliados e administrados (BUSUIOC, 2021; WIDDER; NAFUS, 2023).

Além disso, a assurance contínua depende da definição dos sujeitos responsáveis pela produção, validação e contestação dessas evidências. O provedor ocupa posição central, especialmente em modelos regulatórios inspirados no AI Act europeu, por concentrar maior

capacidade de controle sobre desenho, treinamento, documentação e avaliação inicial dos sistemas. Entretanto, em ambientes sociotécnicos complexos, a responsabilização não pode permanecer restrita ao provedor isoladamente. Integradores, deployers, operadores, fornecedores de dados, provedores de infraestrutura, auditores independentes, autoridades regulatórias e até grupos potencialmente afetados podem desempenhar funções relevantes na identificação de riscos e na contestação de resultados produzidos pela inteligência artificial. Se reduzida a mecanismo interno de autorregistro corporativo, a assurance contínua tende a reproduzir a dependência epistêmica e a assimetria informacional já presentes nos modelos contemporâneos de governança tecnológica. Sua legitimidade depende, portanto, da possibilidade de verificação externa, da preservação de trilhas auditáveis e da existência de mecanismos efetivos de supervisão e contestação institucional.

Em terceiro lugar, é necessário definir a periodicidade e a intensidade da assurance. A lógica contínua não significa revisão permanente e indistinta de todos os elementos do sistema, o que seria operacionalmente inviável e juridicamente desproporcional. Significa, antes, a adoção de uma proporcionalidade dinâmica: quanto maior o risco, maior a densidade dos controles, a frequência das revisões, a granularidade dos registros e a exigência de validação independente. Sistemas utilizados em saúde, crédito, emprego, educação, segurança pública, migração, benefícios sociais ou infraestrutura crítica não podem ser submetidos ao mesmo regime de evidência de aplicações de baixo impacto. O princípio da proporcionalidade, presente no próprio AI Act, deve ser reinterpretado como critério de calibração da intensidade dos controles conforme a gravidade dos danos possíveis, a escala de uso, a vulnerabilidade dos grupos afetados e o grau de autonomia decisória do sistema.

A assurance contínua também não pode ser reduzida a métricas exclusivamente técnicas. Indicadores como acurácia, precisão, recall, robustez, estabilidade, drift ou desempenho estatístico por subgrupos são relevantes para a avaliação de sistemas de inteligência artificial, mas não esgotam sua análise jurídica e institucional. Um sistema pode apresentar elevado desempenho quantitativo e, ainda assim, produzir efeitos incompatíveis com direitos fundamentais em razão da finalidade de uso, da assimetria informacional envolvida, da ausência de mecanismos efetivos de contestação ou da distribuição desigual de impactos entre grupos sociais. A visibilidade técnica do funcionamento do sistema, por si só, não garante compreensão substantiva das estruturas institucionais e relações de poder que condicionam a decisão automatizada (ANANNY; CRAWFORD, 2018). Por essa razão, a assurance contínua deve articular evidências quantitativas

e qualitativas, incorporando avaliação jurídica, análise institucional, supervisão humana significativa, impacto sobre grupos afetados e mecanismos concretos de reparação e contestação.

Esse aspecto é particularmente relevante para evitar que a assurance contínua se converta em mero instrumento de legitimação regulatória. Quando limitada à produção de dashboards, relatórios internos e documentação padronizada, ela corre o risco de apenas sofisticar práticas de compliance performativo, oferecendo aparência de controle sem alterar as assimetrias estruturais que caracterizam a governança algorítmica contemporânea. A produção contínua de evidências somente possui relevância jurídica efetiva quando associada à possibilidade de questionamento, revisão e imposição de consequências institucionais concretas, como correção de dados, ajustes de modelo, restrição de uso, revisão de finalidade, suspensão operacional, comunicação a autoridades competentes ou reavaliação de conformidade. Sem capacidade de produzir efeitos institucionais, a evidência tende a transformar-se em simples registro documental; vinculada a mecanismos efetivos de supervisão e responsabilização, passa a integrar estruturas materiais de accountability.

A categoria da assurance contínua também permite reinterpretar a proteção de dados pessoais no contexto da inteligência artificial de alto risco. A governança de dados não pode ser tratada apenas como etapa preliminar do treinamento de modelos ou como cumprimento formal de obrigações de privacidade. Em sistemas adaptativos, dados operacionais, feedbacks, logs e informações derivadas influenciam continuamente a produção de inferências e a distribuição de riscos. Isso exige mecanismos permanentes de rastreabilidade relacionados à proveniência dos dados, critérios de retenção, atualização de bases, representatividade, minimização, segurança e compatibilidade de finalidade. Além disso, a proteção individualizada revela-se insuficiente diante de danos produzidos em escala relacional e coletiva (VILJOEN, 2021). A assurance contínua deve, portanto, produzir evidências não apenas sobre a conformidade formal do tratamento de dados pessoais, mas também sobre os efeitos distributivos e populacionais decorrentes das inferências algorítmicas e de sua inserção em ambientes sociais amplos.

No plano dogmático, a principal contribuição da assurance contínua é deslocar a confiança regulatória de uma presunção formal para uma justificabilidade permanente. Sistemas de IA de alto risco não devem ser considerados confiáveis apenas porque foram certificados em determinado momento, porque adotam linguagem ética ou porque se alinham abstratamente a princípios de governança. Devem ser considerados juridicamente toleráveis enquanto houver evidências suficientes, atualizadas e verificáveis de que seus riscos permanecem controlados dentro dos

limites normativos aplicáveis. Essa formulação aproxima a assurance contínua de uma teoria procedimental da legitimidade algorítmica: a validade prática do sistema depende da manutenção de condições de rastreabilidade, contestabilidade, supervisão e revisão.

A relação entre AI Act e NIST AI RMF torna-se, assim, mais produtiva. O AI Act oferece uma moldura jurídica vinculante, com obrigações específicas para sistemas de alto risco, mecanismos de avaliação de conformidade e deveres pós-mercado. O NIST AI RMF, por sua vez, oferece uma gramática operacional mais flexível para organizar práticas de governança e gestão de riscos ao longo do ciclo de vida. A assurance contínua emerge justamente na interseção entre esses dois modelos: do AI Act, retira a densidade normativa da conformidade; do NIST AI RMF, retira a lógica iterativa de governança, mapeamento, mensuração e gestão. O resultado não é a simples soma entre regulação europeia e framework norte-americano, mas a formulação de uma categoria analítica capaz de explicar como a conformidade de sistemas de IA de alto risco deve ser continuamente produzida, documentada e submetida a controle.

Contudo, essa categoria permanece limitada se não enfrentar os problemas estruturais desenvolvidos nos capítulos anteriores. A assurance contínua não resolve, por si só, colonialismo de dados, concentração infraestrutural, captura regulatória, desigualdade algorítmica ou privatização transnacional da normatividade. Ela pode, inclusive, ser absorvida por esses mesmos processos, tornando-se serviço privado de certificação, produto de governança automatizada ou camada reputacional para organizações dominantes. Sua força crítica depende de sua vinculação a controles públicos, padrões transparentes, participação social, independência avaliativa e possibilidade real de sanção. Por isso, deve ser compreendida como condição necessária, mas não suficiente, para a legitimidade regulatória da IA de alto risco. Ela melhora a auditabilidade, mas não substitui escolhas políticas sobre limites de uso, proibições, redistribuição de poder e proteção substantiva de direitos fundamentais.

A consequência dogmática é clara: em sistemas de inteligência artificial de alto risco, conformidade não deve ser tratada como estado, mas como processo. Esse processo exige evidências técnicas, jurídicas e organizacionais permanentemente atualizadas; sujeitos responsáveis por produzi-las, validá-las e contestá-las; gatilhos de revisão proporcionais ao risco; conexão entre monitoramento e consequência institucional; e documentação compreendida não como instrumento defensivo da organização, mas como infraestrutura de accountability. A continuous AI assurance, nesses termos, reconfigura os limites jurídicos da confiança porque

substitui a pergunta “o sistema foi aprovado?” por uma pergunta mais exigente: “quais evidências demonstram, agora, que este sistema continua merecendo autorização social e jurídica para operar?”.

7. CONCLUSÃO

Este artigo analisou os limites dos modelos tradicionais de auditoria e governança aplicados a sistemas de inteligência artificial de alto risco, partindo da hipótese de que mecanismos baseados em certificações pontuais, auditorias episódicas e compliance formal são insuficientes para assegurar legitimidade regulatória, accountability e proteção efetiva de direitos fundamentais. Demonstrou-se que a natureza adaptativa, modular e dinâmica da IA contemporânea produz um descompasso entre a velocidade de transformação desses sistemas e a capacidade dos instrumentos tradicionais de controle de acompanhar seus riscos ao longo do tempo. Nesse contexto, a conformidade não pode mais ser compreendida como resultado definitivo de uma validação inicial, mas como condição permanentemente sujeita à reavaliação.

A análise do AI Act europeu e do NIST AI RMF evidenciou a consolidação de uma racionalidade regulatória orientada por gestão de riscos, documentação técnica, rastreabilidade e monitoramento contínuo do ciclo de vida dos sistemas. Contudo, o trabalho também demonstrou que essas estruturas permanecem tensionadas por assimetrias de poder relacionadas à concentração infraestrutural da inteligência artificial, à dependência tecnológica e ao controle privado sobre dados, padrões técnicos e capacidade computacional. Os riscos da IA não decorrem apenas do momento final da decisão automatizada, mas também das condições econômicas, institucionais e infraestruturais que tornam possível sua operação. Por essa razão, modelos regulatórios centrados exclusivamente na explicação do output algorítmico ou na correção posterior de decisões individuais mostram-se estruturalmente insuficientes.

Nesse cenário, a *continuous AI assurance* foi proposta como categoria jurídico-regulatória voltada à reorganização dos mecanismos tradicionais de controle. Sua contribuição não consiste em substituir auditorias, avaliações de impacto ou regimes de compliance, mas em integrá-los a um processo contínuo de produção, verificação e contestação de evidências relacionadas ao funcionamento dos sistemas de IA. A confiança regulatória deixa, assim, de depender de uma presunção formal baseada em certificações iniciais e passa a exigir justificabilidade permanente,

sustentada por mecanismos contínuos de supervisão, rastreabilidade, revisão e responsabilização institucional.

A pesquisa também demonstrou que a assurance contínua não deve ser idealizada. Se reduzida a relatórios internos, métricas padronizadas ou práticas meramente performativas de compliance, ela pode funcionar apenas como instrumento sofisticado de legitimação das mesmas estruturas de poder que pretende controlar. Sua efetividade depende da possibilidade de verificação externa, da existência de mecanismos de contestação, da produção de consequências institucionais concretas e da preservação de uma orientação material voltada à proteção de direitos fundamentais. A evidência regulatória somente adquire relevância jurídica quando pode ser utilizada para impor revisões, limitações, ajustes ou até a suspensão de sistemas incompatíveis com os parâmetros normativos aplicáveis.

Conclui-se, portanto, que a crise contemporânea da auditabilidade algorítmica exige uma mudança de paradigma: da validação pontual para a assurance contínua; da documentação defensiva para a rastreabilidade permanente; do compliance formal para mecanismos efetivos de accountability. Em sistemas de inteligência artificial de alto risco, a questão regulatória decisiva não é apenas se o sistema foi aprovado em determinado momento, mas quais evidências demonstram, de forma contínua e verificável, que ele permanece compatível com os limites jurídicos e institucionais exigidos pela proteção de direitos fundamentais em sociedades crescentemente organizadas por decisões automatizadas

REFERÊNCIAS BIBLIOGRÁFICAS

EUROPEAN UNION. Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act). Official Journal of the European Union, Brussels, 2024. Disponível em: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>

MOKANDER, Jakob et al. Conformity Assessments and Post-market Monitoring: A Guide to the Role of Auditing in the Proposed European AI Regulation. arXiv, 2021. Disponível em: <https://arxiv.org/abs/2111.05071>

NATIONAL INSTITUTE OF STANDARDS AND TECHNOLOGY (NIST). Artificial Intelligence Risk Management Framework (AI RMF 1.0). Gaithersburg: U.S. Department of Commerce, 2023. Disponível em: <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>

AMOORE, Louise. Cloud Ethics: Algorithms and the Attributes of Ourselves and Others. Durham: Duke University Press, 2020.

ANANNY, Mike; CRAWFORD, Kate. Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *New Media & Society*, v. 20, n. 3, p. 973-989, 2018.

BENJAMIN, Ruha. Race After Technology: Abolitionist Tools for the New Jim Code. Cambridge: Polity, 2019.

BLACK, Julia. Decentring Regulation: Understanding the Role of Regulation and Self-Regulation in a “Post-Regulatory” World. *Current Legal Problems*, v. 54, n. 1, p. 103-146, 2001.

BURRELL, Jenna. How the machine “thinks”: Understanding opacity in machine learning algorithms. *Big Data & Society*, v. 3, n. 1, p. 1-12, 2016.

BUSUIOC, Madalina. Accountable Artificial Intelligence: Holding Algorithms to Account. *Public Administration Review*, v. 81, n. 5, p. 825-836, 2021.

COHEN, Julie E. Between Truth and Power: The Legal Constructions of Informational Capitalism. New York: Oxford University Press, 2019.

COULDRY, Nick; MEJIAS, Ulises A. Data Colonialism: Rethinking Big Data’s Relation to the Contemporary Subject. *Television & New Media*, v. 20, n. 4, p. 336-349, 2019.

COULDRY, Nick; MEJIAS, Ulises A. The Costs of Connection: How Data Is Colonizing Human Life and Appropriating It for Capitalism. Stanford: Stanford University Press, 2019.

EUBANKS, Virginia. Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor. New York: St. Martin’s Press, 2018.

FUKUDA-PARR, Sakiko; GIBBONS, Elizabeth. Emerging Consensus on “Ethical AI”: Human Rights Critique of Stakeholder Guidelines. *Global Policy*, v. 12, supl. 6, p. 32-44, 2021.

HILDEBRANDT, Mireille. *Smart Technologies and the End(s) of Law: Novel Entanglements of Law and Technology*. Cheltenham: Edward Elgar, 2015.

MICKLITZ, Hans-W. Incomplete Interaction of Public–Private Governance in the AI Act. In: BROWNSWORD, Roger; DIMATTEO, Larry A. (eds.). *The Cambridge Handbook of the Governance of Technology*. Cambridge: Cambridge University Press, 2025. p. 265-288.

NOBLE, Safiya Umoja. *Algorithms of Oppression: How Search Engines Reinforce Racism*. New York: New York University Press, 2018.

OECD. *Advancing Accountability in AI: Governing and Managing Risks Throughout the Lifecycle for Trustworthy AI*. Paris: OECD Publishing, 2023.

O’NEIL, Cathy. *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown, 2016.

RAJI, Inioluwa Deborah et al. Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing. *FAT* Conference*, 2020.

ROUVROY, Antoinette; BERNIS, Thomas. Algorithmic Governmentality and Prospects of Emancipation: Disparateness as a Precondition for Individuation through Relationships? *Réseaux*, n. 177, p. 163-196, 2013.

SCOTT, Colin. Regulation in the Age of Governance: The Rise of the Post-Regulatory State. In: JORDANA, Jacint; LEVI-FAUR, David (eds.). *The Politics of Regulation: Institutions and Regulatory Reforms for the Age of Governance*. Cheltenham: Edward Elgar, 2004. p. 145-174.

VAN DIJCK, José. Datafication, dataism and dataveillance: Big Data between scientific paradigm and ideology. *Surveillance & Society*, v. 12, n. 2, p. 197-208, 2014.

VAN DIJCK, José; POELL, Thomas; DE WAAL, Martijn. *The Platform Society: Public Values in a Connective World*. New York: Oxford University Press, 2018.

VEALE, Michael; ZUIDERVEEN BORGESIOUS, Frederik J. Demystifying the Draft EU Artificial Intelligence Act: Analysing the good, the bad, and the unclear elements of the proposed approach. *Computer Law Review International*, v. 22, n. 4, p. 97-112, 2021.

VILJOEN, Salomé. A Relational Theory of Data Governance. *Yale Law Journal*, v. 131, n. 2, p. 573-654, 2021.

WIDDER, David Gray; NAFUS, Dawn. Dislocated accountabilities in the “AI supply chain”: Modularity and developers’ notions of responsibility. *Big Data & Society*, v. 10, n. 1, 2023.

WIDDER, David Gray; WHITTAKER, Meredith; WEST, Sarah Myers. Why “open” AI systems are actually closed, and why this matters. *Nature*, v. 635, n. 8040, p. 827-833, 2024.

YEUNG, Karen. Algorithmic Regulation: A Critical Interrogation. *Regulation & Governance*, v. 12, n. 4, p. 505-523, 2018.

ZUBOFF, Shoshana. *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. New York: PublicAffairs, 2019.