

Fórum: Perspectivas Práticas

Ética algorítmica e regulação econômica: como autoridades antitruste enfrentam os riscos da IA?

Mayla Cristina Costa Maroni Saraiva ^{1*}Fátima de Souza Freire ²

* Autora Correspondente

¹ Universidade de Brasília, Programa de Pós-Graduação em Administração e Programa de Pós-Graduação em Ciências Contábeis, Brasília, DF, Brasil² Universidade de Brasília, Programa de Pós-Graduação em Ciências Contábeis, Brasília, DF, Brasil

A integração da Inteligência Artificial (IA) em órgãos de defesa da concorrência para identificar práticas anticompetitivas acarreta riscos éticos inerentes, como o viés algorítmico e a dependência excessiva de fornecedores de tecnologia. Visando mitigar tais riscos, este artigo investiga como 35 autoridades reguladoras, classificadas pelo *Global Competition Review (GCR) Enforcement Rating* de 2023, abordam esses desafios. Por meio de uma análise documental realizada com o auxílio do ATLAS.ti25, comparam-se as organizações com base em cinco princípios éticos (Transparência, Responsabilidade/*Accountability*, Justiça e Equidade, Robustez e Segurança, e Privacidade), fundamentados em frameworks consolidados de ética da IA e bioética de proteção. Os resultados demonstram uma disparidade crítica na maturidade regulatória: observa-se maior rigor técnico em princípios de ética operacional (Robustez, Privacidade e Segurança) do que em princípios de ética social. Persistem deficiências substanciais em Transparência, Responsabilidade e Justiça/Equidade. Essa lacuna aponta para um déficit no princípio central da Explicabilidade (que engloba a inteligibilidade e a prestação de contas). A causa principal reside na ausência de procedimentos objetivos e na falta de divulgação sobre o uso da IA para a atividade-fim dessas organizações. O estudo conclui que o aprimoramento da liderança ética e o estabelecimento de mecanismos organizacionais de prestação de contas são fundamentais para garantir uma aplicação justa e transparente da IA na regulação econômica.

Palavras-chave: inteligência artificial, regulação econômica, ética algorítmica, explicabilidade, integridade pública.

DOI: <https://doi.org/10.1590/0034-761220250255>

Submetido em 29 de junho de 2025 e aceito para publicação em 05 de março de 2026.

Editora-chefe:Alketa Peci, Fundação Getúlio Vargas, Rio de Janeiro, RJ, Brasil **Editora adjunta:**Gabriela Spanghero Lotta, Fundação Getúlio Vargas, São Paulo, SP, Brasil **Pareceristas:**Fernando Filgueiras, Escola Nacional de Administração Pública, Brasília, DF, Brasil Kleber Cuissi Canuto, Federação das Indústrias do Estado do Paraná, Curitiba, PR, Brasil Maria Irene da Fonseca e Sá, Universidade Federal do Rio de Janeiro, Rio de Janeiro, RJ, Brasil **Relatório de revisão por pares:**O relatório de revisão por pares está disponível em <https://periodicos.fgv.br/rap/article/view/97069/90457>

ISSN: 1982-3134



Algorithmic ethics and economic regulation: how do antitrust authorities address the risks of AI?

The integration of Artificial Intelligence (AI) into competition authorities to detect anti-competitive practices entails inherent ethical risks, such as algorithmic bias and excessive dependence on technology providers. To mitigate these risks, this article investigates how thirty-five regulatory authorities, ranked in the 2023 GCR Enforcement Rating, address these challenges. A document analysis, carried out using ATLAS.ti25, compares these organizations based on five ethical principles (transparency, accountability, fairness and equity, robustness and security, and privacy), grounded in consolidated frameworks in AI ethics and protection bioethics. The results reveal a critical disparity in regulatory maturity: greater technical rigor is observed in operational ethics principles (robustness, privacy, and security) than in social ethics principles. Substantial deficiencies persist in transparency, accountability, and fairness/equity. This gap points to a deficit in the core principle of explainability (encompassing both intelligibility and accountability). The main cause lies in the absence of clear procedures and limited disclosure regarding the use of AI in the core activities of these organizations. The study concludes that strengthening ethical leadership and establishing organizational accountability mechanisms are essential to ensure the fair and transparent application of AI in economic regulation.

Keywords: artificial intelligence, economic regulation, algorithmic ethics, explainability, public integrity.

Ética algorítmica y regulación económica: cómo las autoridades antimonopolio enfrentan los riesgos de la IA?

La integración de la inteligencia artificial (IA) en los órganos de defensa de la competencia para identificar prácticas anticompetitivas conlleva riesgos éticos inherentes, como el sesgo algorítmico y una dependencia excesiva de proveedores tecnológicos. Para mitigar tales riesgos, este artículo investiga cómo treinta y cinco autoridades reguladoras, clasificadas según el Rating de Cumplimiento de la Global Competition Review (GCR) de 2023, abordan estos desafíos. A través de un análisis documental llevado a cabo mediante ATLAS.ti25, se comparan las organizaciones basándose en cinco principios éticos: Transparencia, Responsabilidad (*Accountability*), Justicia y Equidad, Robustez y Seguridad, y Privacidad, fundamentados en marcos consolidados de ética de la IA y bioética de protección. Los resultados revelan una disparidad crítica en la madurez regulatoria: se observa un mayor rigor técnico en los principios de ética operacional (Robustez, Privacidad y Seguridad) que en los de ética social. Persisten deficiencias sustanciales en Transparencia, Responsabilidad y Justicia/Equidad. Esta brecha apunta a un déficit en el principio central de Explicabilidad (que abarca tanto la inteligibilidad como la rendición de cuentas). La causa principal radica en la ausencia de procedimientos claros y en la falta de divulgación sobre el uso de la IA para el propósito central de estas organizaciones. El estudio concluye que el fortalecimiento del liderazgo ético y el establecimiento de mecanismos organizacionales de rendición de cuentas son esenciales para garantizar una aplicación justa y transparente de la IA en la regulación económica.

Palabras clave: inteligencia artificial, regulación económica, ética algorítmica, explicabilidad, integridad pública.

1. INTRODUÇÃO

A Inteligência Artificial (IA) emergiu como um vetor de transformação das práticas de governança pública, com implicações crescentes sobre a regulação econômica global. No campo antitruste, sua incorporação às rotinas de investigação e análise de condutas empresariais tem ampliado significativamente a capacidade de detecção de cartéis, fusões e práticas anticompetitivas (Lu et al., 2019; von Ingersleben-Seip, 2023). No entanto, essa incorporação da IA, ao deslocar a tomada de decisão para sistemas algorítmicos complexos, precipita um conjunto de dilemas éticos e institucionais de alta complexidade. Nesse sentido, questiona-se: como as autoridades de defesa da concorrência podem assegurar transparência, responsabilidade e equidade em decisões automatizadas que afetam direitos econômicos fundamentais?

Essa questão configura o problema de pesquisa central deste artigo: compreender como as autoridades de defesa da concorrência têm incorporado princípios éticos na adoção e no uso da IA, e se a maturidade institucional dessas práticas é suficiente para mitigar os riscos à integridade das decisões regulatórias e à confiança pública.

A relevância do problema se dá devido à movimentação normativa global pela institucionalização da “ética algorítmica” (Floridi & Cowls, 2019; Jobin et al., 2019). Desde 2018, diversas jurisdições têm formulado estratégias nacionais e quadros regulatórios – como a *German AI Strategy*, o *Model AI Governance Framework* de Singapura (Singapore Government, 2019) e o notório *AI Act* da União Europeia (Veale & Borgesius, 2021) –, refletindo um esforço para estabelecer parâmetros de governança digital em consonância com o princípio da dignidade humana.

Contudo, apesar da proliferação de declarações normativas, persiste uma lacuna crítica na evidência empírica sobre como os órgãos reguladores traduzem tais princípios em práticas concretas de governança da IA. A literatura enfatiza riscos como viés algorítmico, opacidade decisória e dependência tecnológica (Raji & Dobbe, 2023; Schmude et al., 2023), mas pouco se sabe sobre a maturidade ética dos sistemas de IA utilizados pelas autoridades antitruste e, crucialmente, se os fatores institucionais (como orçamento e estabilidade de pessoal) realmente influenciam essa maturidade.

Neste estudo, examinamos essa lacuna por meio de uma análise comparativa em 35 autoridades de defesa da concorrência, listadas no *Global Competition Review (GCR) – Enforcement Rating 2023*. A investigação toma como referência os cinco princípios de Floridi e Cowls (2019) – Transparência, Responsabilidade, Justiça e Equidade, Robustez e Segurança, e Privacidade – e os articula com os princípios bioéticos de proteção de Schramm (2008), que enfatizam a prevenção de danos e a proteção de vulneráveis. A integração dessas duas tradições permite verificar se os marcos regulatórios analisados transcendem o formalismo, promovendo uma governança protetiva.

Metodologicamente, a pesquisa emprega uma análise documental com o auxílio do ATLAS.ti25, cruzando achados qualitativos com indicadores de capacidade institucional. Adicionalmente, aplicam-se regressões múltiplas para testar a hipótese de que o perfil institucional influencia a maturidade ética regulatória, além de análises de correlação regional. Ao adotar essa abordagem multidimensional, o artigo oferece uma contribuição empírica original ao debate sobre como as organizações antitruste podem utilizar a IA de forma justa e transparente, atendendo aos requisitos de uma governança democrática e protetiva.

Os dados rejeitam a premissa de que a maturidade ética seja um subproduto linear do recurso orçamentário disponível. Observa-se uma ‘saturação da maturidade ética’, em que o aporte financeiro incremental prioriza a robustez técnica, mas falha em romper a opacidade dos sistemas, mantendo o déficit de explicabilidade mesmo em jurisdições de alta renda.

Além desta introdução, o artigo estrutura-se em quatro seções. A primeira estabelece o referencial teórico, articulando o meta-princípio da explicabilidade à bioética de proteção no contexto da governança algorítmica pelos órgãos de regulação econômica. A segunda detalha os procedimentos metodológicos, com ênfase na análise documental via ATLAS.ti25 e na modelagem estatística aplicada. A terceira apresenta a análise e a discussão dos resultados, contrastando a maturidade ética e a capacidade institucional entre agrupamentos regionais. Por fim, a quarta seção sintetiza os achados e propõe diretrizes para o fortalecimento de uma governança algorítmica democrática e protetiva.

2. REFERENCIAL TEÓRICO

2.1 A Explicabilidade como meta-princípio na governança ética da IA

A expansão da inteligência artificial é um vetor de transformação sociotécnica que reestrutura o trabalho e a gestão pública (Hoch & Engelmann, 2023; Roberts et al., 2021), mas introduz riscos sistêmicos que exigem uma resposta regulatória sofisticada. O problema central reside na assimetria de poder e na concentração de desenvolvimento da IA em um número limitado de organizações (Agar, 2020; Crawford, 2021), o que desafia a capacidade das estruturas regulatórias existentes, consideradas estruturalmente despreparadas para a dinamicidade tecnológica (Bélisle-Pipon et al., 2021; Häußler, 2021).

Esse cenário de risco concentrado e governança incipiente impulsionou o consenso de que a ética deve ser central na IA (Bostrom & Yudkowsky, 2011; Etzioni & Etzioni, 2017; Coeckelbergh, 2020). A resposta global manifestou-se na “proliferação de princípios” – uma vasta produção de diretrizes por governos, indústrias e academia (Fjeld et al., 2020; Hagendorf, 2020; Jobin et al., 2019; Ryan & Stahl, 2020; Zeng et al., 2018) visando orientar a conduta de *stakeholders* e mitigar o uso indevido (Floridi et al., 2018; Greene et al., 2019).

Para transcender a dispersão normativa, o presente estudo adota o arcabouço unificado de Floridi e Cowls (2019), que consolida a ética da IA em cinco princípios, ancorados na tradição bioética principialista (Beauchamp & Childress, 2001):

- (1) Beneficência: foco no bem-estar e no interesse público.
- (2) Não maleficência: dever de prevenir danos, garantir a Robustez e Segurança dos sistemas.
- (3) Autonomia: garantia de que a IA permaneça sob o controle e a capacidade de decisão humana.
- (4) Justiça: compromisso com a Equidade e a mitigação de vieses discriminatórios.

A contribuição fundamental de Floridi e Cowls (2019) reside na elevação da Explicabilidade ao status de meta-princípio. Em sistemas algorítmicos opacos, os princípios éticos clássicos tornam-se inexecutáveis sem um mecanismo que viabilize sua verificação. A explicabilidade preenche essa lacuna, materializando-se em duas dimensões complementares: (1) inteligibilidade epistemológica, que fundamenta a transparência sobre o funcionamento interno do sistema; e (2) responsabilidade ética (accountability), que institucionaliza a rastreabilidade das decisões.

No contexto antitruste, essa exigência ganha contornos críticos. A ausência de explicabilidade nas decisões automatizadas transita o problema de uma ‘vulnerabilidade universal’ para uma ‘vulneração concreta’ (Schramm, 2008). Quando o agente regulado é privado da compreensão dos critérios de uma sanção, o déficit de transparência deixa de ser apenas uma falha técnica para se configurar como uma barreira ao exercício do contraditório e do direito de defesa.

Portanto, o desafio da governança de IA nos órgãos de concorrência – refletido em exemplos de alto risco (Biondi & Cernev, 2023) – transcende a mera contratação de pessoal ou o aumento orçamentário. Exige uma mudança de paradigma que priorize a Transparência e a Responsabilidade sobre a excelência técnica pura, garantindo que o desenvolvimento e a aplicação da IA sejam conduzidos de forma

democrática e protetiva (Stahl, 2021). A insuficiência da maturidade institucional em implementar a Explicabilidade é, em última análise, o risco mais significativo à integridade das decisões regulatórias.

1.2 O uso da IA pelos conselhos de regulação econômica

O interesse na integração da inteligência artificial em órgãos reguladores de competição é crescente (Smuha, 2021), impulsionado pela capacidade da tecnologia de processar grandes volumes de dados e identificar padrões complexos de práticas anticompetitivas, como fixação de preços e colusão (Organization for Economic Co-operation and Development [OECD], 2021). Técnicas avançadas de IA, como análise de dados de *e-commerce* e modelos preditivos, têm se mostrado eficazes para antecipar e investigar comportamentos anticompetitivos (Bodrick et al., 2024), aprimorando a eficácia e eficiência regulatória (Arner et al., 2021; Schrepel, 2021).

Contudo, essa adoção acelerada está intrinsecamente associada a riscos éticos: a coleta e análise massiva de dados suscitam preocupações de Privacidade e Segurança (Biondi & Cernev, 2023), enquanto a dependência de fornecedores externos e a opacidade algorítmica comprometem a Autonomia e a Responsabilidade dos órgãos reguladores (Raji & Dobbe, 2023). A literatura aponta que a falta de Transparência dos algoritmos e o uso de conjuntos de dados enviesados podem resultar em decisões injustas ou discriminatórias (Ezrachi & Stucke, 2019).

Os riscos éticos associados ao uso da IA em qualquer domínio, incluindo o antitruste, foram categorizados por Cave e ÓhÉigearthaigh (2018) em (1) riscos de autonomia, (2) riscos de justiça, (3) riscos de explicabilidade e (4) riscos de colaboração. Tais riscos se alinham diretamente aos princípios do arcabouço unificado de Floridi e Cowls (2019) – Não maleficência; Autonomia; Justiça; Beneficência; e, notavelmente, Explicabilidade –, que servem como a estrutura básica para a análise regulatória.

Embora o arcabouço de Floridi e Cowls (2019) forneça a base universal da ética da IA, a integração da Bioética de Proteção (BP) de Schramm (2008) cumpre um propósito crucial para o ineditismo deste estudo: mover o foco da vulnerabilidade (condição universal e potencial de dano) para a vulneração (a manifestação concreta e atualizada do dano).

Schramm (2008) estabelece uma distinção fundamental, de origem latino-americana, essencial para avaliar a eficácia da regulação:

- (1) Vulnerabilidade (Potencial): é a condição humana universal, inerente à finitude. Todos os agentes regulados são universalmente vulneráveis.
- (2) Vulneração (Ato/Fato): é o estado de ser efetivamente atingido, afetado e ferido por um dano ou uma carência concreta e constatável.

No contexto antitruste, a falha nos princípios éticos de Explicabilidade (Transparência/Responsabilidade) e Justiça resulta na vulneração algorítmica dos agentes econômicos. Se as decisões automatizadas de um órgão regulador forem opacas ou enviesadas, o déficit ético transita do plano da possibilidade (vulnerabilidade) para o plano do dano real (vulneração), justificando a adoção de uma governança protetiva que garanta a integridade e a Justiça no contexto do mercado.

3. PROCEDIMENTOS METODOLÓGICOS

O presente estudo adota um desenho de pesquisa de método misto, predominantemente qualitativo, exploratório e comparativo, complementado por análises estatísticas para testar hipóteses sobre a capacidade institucional. A abordagem qualitativa se baseia na análise documental de 35 autoridades de defesa da concorrência, identificadas no GCR – *Enforcement Rating 2023*. O corpus documental foi construído com base em relatórios anuais, políticas de governança e diretrizes éticas/técnicas de IA disponíveis nos respectivos sites oficiais, categorizados com o auxílio do software de análise qualitativa ATLAS.ti25 (Babbie, 2016).

A utilização do *software* ATLAS.ti25 permitiu a sistematização de grandes volumes de dados normativos, otimizando a identificação de padrões semânticos por meio de técnicas de codificação assistida. Essa abordagem está alinhada à tendência de emprego de modelos computacionais avançados para o processamento eficiente de informações e predição de prioridades em ambientes de alta complexidade técnica (Bani-Salameh et al., 2021).

O arcabouço teórico para a codificação e categorização é composto por um sistema de dois níveis, garantindo o ineditismo e a profundidade da análise. Nesse sentido, as categorias analíticas principais foram operacionalizadas de acordo com os princípios éticos universais de Floridi e Cowls (2019), nos quais a Explicabilidade é tratada como um princípio transversal que sustenta a Transparência (inteligibilidade técnica) e a Responsabilidade (*accountability* institucional), compondo o pilar da ética social do presente estudo, conforme a Tabela 1:

TABELA 1 CORRESPONDÊNCIA ENTRE OS PRINCÍPIOS DE FLORIDI E COWLS E AS CATEGORIAS ANALÍTICAS DO ESTUDO

Princípio ético (Floridi & Cowls, 2019)	Categoria analítica (Presente estudo)	Função ética estratégica
Explicabilidade	Transparência e Responsabilidade	Permite compreender e atribuir responsabilidade às decisões algorítmicas.
Justiça	Justiça e Equidade	Garante decisões não discriminatórias e distribuição equitativa de benefícios.
Não maleficência	Robustez, Segurança e Privacidade	Impede danos e assegura proteção de dados e segurança operacional.
Beneficência	Finalidade regulatória	Direciona a IA para o bem público e o fortalecimento institucional.
Autonomia	Controle humano decisório	Mantém o poder de decisão e reversibilidade sobre sistemas automatizados.

Fonte: Adaptada de Floridi e Cowls (2019).

A estrutura de Floridi e Cowls é complementada pela BP de Schramm (2008). Essa integração aprofunda a avaliação, deslocando o foco da mera vulnerabilidade para o dever institucional de proteção contra a vulneração (dano concreto). Essa abordagem permite avaliar se a maturidade ética das agências transcende o formalismo (Tabela 2).

TABELA 2 PRINCÍPIOS DE PROTEÇÃO APLICADOS À ÉTICA DA IA

Princípio bioético (Schramm, 2008)	Aplicação à ética da IA	Objeto de proteção
Não maleficência	Garantir robustez técnica, segurança e uso ético dos dados.	Protege contra riscos, falhas e violações de privacidade.
Autonomia	Assegurar controle humano e decisão reversível sobre algoritmos.	Protege a liberdade de escolha e a autodeterminação dos usuários.
Justiça	Monitorar e mitigar vieses algorítmicos.	Protege a equidade e evita discriminação.
Beneficência	Direcionar a IA para o bem comum e o interesse público.	Promove o bem-estar coletivo e a integridade de mercado.

Fonte: Adaptada de Schramm (2008).

A Explicabilidade emerge, nessa tabela, como pré-condição da proteção: apenas sistemas inteligíveis e auditáveis permitem identificar danos, garantir autonomia e responsabilizar agentes. O procedimento analítico ocorreu em três etapas principais, visando à triangulação metodológica entre os dados qualitativos (códigos éticos) e os dados quantitativos (indicadores institucionais):

- (1) Codificação dos trechos relevantes à IA com o ATLAS.ti25 e organização dos códigos em torno das 5 categorias analíticas (Tabela 1).
- (2) Realização de múltiplas regressões para testar formalmente a hipótese de que o perfil institucional (recursos, estabilidade) influencia a maturidade regulatória ética, conforme sumariado na Tabela 3.
- (3) Confronto da maturidade ética (escores qualitativos) com os indicadores institucionais, seguido de agrupamento dos dados por regiões mundiais. Observou-se a associação entre as médias regionais por meio do coeficiente de correlação de Pearson para conferir robustez às inferências exploratórias.

Para viabilizar a transição da análise qualitativa para a quantitativa, a maturidade ética foi operacionalizada por meio de uma escala ordinal de 0 a 4, aplicada a cada categoria analítica no ATLAS.ti25. A pontuação seguiu critérios de densidade normativa e institucionalização: (0) ausência

de menção; (1) menção genérica; (2) diretrizes declaratórias sem mecanismos de aplicação; (3) procedimentos técnicos específicos e guias de conduta; e (4) evidência de mecanismos de auditoria ou prestação de contas institucionalizados. Esse rigor na atribuição de escores visa garantir a replicabilidade do estudo e a consistência dos dados que alimentam os modelos estatísticos subsequentes.

TABELA 3 SUMÁRIO DOS MODELOS DE REGRESSÃO: VARIÁVEIS INSTITUCIONAIS VS. MATURIDADE ÉTICA ALGORÍTMICA

Análise	Variável dependente (Y)	Variáveis independentes (X)	Objetivo
Regressão 1	Pontuação total de maturidade ética	Orçamento	Verificar se recursos financeiros estão correlacionados com a maturidade ética geral.
Regressão 2	Transparência	Salário presidente + número de funcionários	Explorar se o investimento em capital humano e o porte da instituição afetam a pontuação no princípio de Transparência (que se liga à Explicabilidade).
Regressão 3	Robustez e Segurança	Tempo médio de permanência (em anos) + % de funcionários com mais de 5 anos de experiência	Verificar se a estabilidade do corpo técnico (experiência) influencia o rigor nos princípios técnicos/operacionais.

Fonte: Dados da Pesquisa.

Por fim, cumpre salientar as limitações inerentes ao desenho desta pesquisa. Primeiramente, o corpus documental restringe-se a documentos públicos, o que pode omitir práticas de governança interna não divulgadas ou, opostamente, superestimar compromissos meramente formais (*ethics washing*). Adicionalmente, o tamanho reduzido da amostra em análises regionais específicas demanda que os resultados das regressões sejam interpretados como tendências indicativas de correlação, e não como generalizações universais definitivas. Tais limitações, contudo, não invalidam o ineditismo do diagnóstico, servindo como base para futuras investigações que possam incluir dados primários ou entrevistas com os reguladores.

4. ANÁLISE E DISCUSSÃO DOS DADOS: A DISSONÂNCIA ENTRE MATURIDADE TÉCNICA E ÉTICA SOCIAL

A análise comparativa entre as 35 autoridades de defesa da concorrência, articulando a avaliação qualitativa dos princípios éticos com os indicadores institucionais, revela uma dissonância crítica: a maturidade é alta nos aspectos técnicos (ética operacional), mas persistentemente baixa nos aspectos sociais (ética social).

Os resultados revelam que a maturidade institucional é insuficiente para mitigar os riscos éticos de justiça e transparência impostos pela IA. Essa conclusão é sustentada pela triangulação entre a análise regional (Tabela 4) e os testes de regressão (Tabelas 5 e 6), que dissociam o avanço ético da mera alocação de recursos financeiros.

TABELA 4 CONTEXTUALIZAÇÃO REGIONAL: MATURIDADE ÉTICA VS. CAPACIDADE INSTITUCIONAL (MÉDIA)

Região	Média orçamento (milhões de €)	Média nº funcionários	Média tempo permanência (anos)	Média pontuação total ética	Média Robustez/ privacidade (operacional)	Média Transparência/ Justiça (Social)
América Latina (AL)	8,86	123,6	6,18	6,2	3,6	2,6
Europa (UE/EEE)	16,29	188,8	6,67	6,75	4	2,75
América do Norte (EUA)	162,70	1.083	N/A	7	4	3
Ásia-Pacífico e África (APAC)	22,84	193,8	7,55	6,8	4	2,8

Fonte: Dados da pesquisa.

A segmentação das jurisdições em quatro agrupamentos regionais – América do Norte (EUA), Europa (União Europeia/Espaço Econômico Europeu), Ásia-Pacífico e África (APAC) e América Latina (AL) – foi fundamental para contextualizar a heterogeneidade da maturidade ética algorítmica. Esta análise buscou transcender a mera descrição, estabelecendo pontes com o referencial teórico que aborda os princípios éticos de Floridi e Cowls (2019) e a bioética de proteção (Schramm, 2008, 2017), permitindo identificar como a capacidade institucional e o capital humano acumulado (tempo de permanência) influenciam a transição da conformidade formal para a robustez ética efetiva.

3.1 América do Norte (EUA): recursos massivos e desafio de transparência

O perfil do *Federal Trade Commission* (FTC) dos EUA serve como um *benchmark* de escala, operando com um orçamento de aproximadamente 162,70 milhões de euros e um corpo de 1.083 funcionários. Apesar da pontuação total de 7 ser a mais alta, essa performance reforça a principal conclusão da Regressão 1 (que pode ser vista na Tabela 5), de que o recurso financeiro é uma condição necessária, mas não suficiente para a excelência ética. A diferença de recursos em relação à América Latina (18 vezes maior) é desproporcional ao ganho marginal na pontuação (apenas 0,8 ponto). Isso indica uma saturação da maturidade ética, na qual o investimento adicional não se traduz em melhoria nos princípios mais complexos.

A pontuação perfeita na ética operacional (Robustez/Privacidade) (4) é compatível com o foco regulatório do FTC em proteger o consumidor contra danos técnicos diretos. Contudo, a persistência de um desafio na ética social (3), especialmente em Transparência, sugere uma dificuldade em implementar o princípio de Explicabilidade de Floridi e Cows (2019) em todo o ciclo decisório.

3.2 Europa (UE/EEE): estabilidade técnica e liderança normativa

A Europa demonstra um perfil de recursos medianos (16,29 milhões de euros e 188,8 funcionários) com uma alta estabilidade institucional, refletida no tempo médio de permanência (6,67 anos), o maior entre os grupos com dados completos. Essa estabilidade é crucial: a excelência consistente na ética operacional (4), que iguala a performance dos EUA apesar da diferença orçamentária, sugere que a estabilidade da força de trabalho e o conhecimento institucional acumulado (conforme indicado pela Regressão 3) são o principal motor para a excelência em *Robustez e Segurança*.

Por sua vez, a maturidade ética elevada (6,75) reflete uma cultura de conformidade consolidada por marcos regulatórios como o Regulamento Geral sobre a Proteção de Dados (RGPD) e o *AI Act* (Veale & Borgesius, 2021) — este último, o primeiro regramento global a categorizar sistemas de IA por níveis de risco. Contudo, a Europa compartilha o desafio comum na ética social (2,75). Essa lacuna aponta para um gargalo estrutural na transposição de princípios abstratos de Justiça e Transparência para mecanismos operacionais que garantam a autonomia e mitiguem a vulneração humana, conforme postula a Bioética de Proteção (Schramm, 2008, 2017).

3.3 Ásia-Pacífico e África – APAC: experiência e rigor técnico

As jurisdições deste grupo demonstram uma abordagem pragmática, amparada por marcos de referência como o *Model AI Governance Framework* (Singapore Government, 2019), onde a ética social (2,8) é considerada, mas o foco recai sobre a segurança e a estabilidade tecnológica. Este grupo apresenta um perfil competitivo de recursos (22,84 milhões de euros) e a maior estabilidade (7,55 anos) na amostra. O tempo de permanência extremamente alto reforça a hipótese da Regressão 3. A maior longevidade e experiência do corpo técnico neste grupo provavelmente capacitam as jurisdições a alcançarem a melhor pontuação em ética operacional (4) e uma pontuação total de 6,8, superando ligeiramente a Europa e demonstrando que o conhecimento institucional é um poderoso ativo ético. A alta pontuação em *Robustez e Privacidade* reflete o alinhamento global com o rigor técnico, mas as jurisdições demonstram uma abordagem pragmática, onde a ética social (2,8) é considerada, mas não priorizada em detrimento da segurança e da estabilidade tecnológica.

3.4 América Latina (AL): Desafios Estruturais e Vulneração Ética

A América Latina enfrenta os desafios mais significativos, operando com os recursos mais limitados (8,86 milhões de euros) e a menor estabilidade (6,18 anos). A pontuação mais baixa em ética operacional (3,6) e a menor estabilidade confirmam o impacto da estabilidade (Regressão 3): a alta rotatividade institucional na região pode estar minando a capacidade de construir e manter políticas de Robustez e Segurança complexas, que exigem retenção de talentos em tecnologia. A pontuação mais baixa em ética social (2,6) é a lacuna mais acentuada. Em um contexto de recursos limitados e menor capacidade técnica, a dificuldade em implementar a Explicabilidade (Transparência/Responsabilidade) acentua o risco de vulneração humana (Schramm, 2017).

A opacidade algorítmica, neste cenário, não é apenas uma falha técnica, mas um problema moral concreto, pois impede que os indivíduos e empresas reguladas exerçam sua autonomia e compreendam as razões para o dano ou a carência que lhes foi imposta pelo sistema de IA. Apesar do desempenho ético ser o mais baixo, a pontuação de 6,2/10, mesmo com um orçamento limitado, sugere que as jurisdições latino-americanas têm buscado adotar políticas de IA de forma eficiente. Contudo, os desafios estruturais e a instabilidade impedem o avanço para a governança mais complexa da ética social.

Os resultados das regressões múltiplas iniciais (Tabela 5) indicaram uma dissonância estatística. As Regressões 1 e 2 (Orçamento vs. Ética Total; Pessoal vs. Transparência) não encontraram significância estatística, sugerindo que o avanço ético não é uma função linear da quantidade de recursos.

TABELA 5 SUMÁRIO DOS MODELOS DE REGRESSÃO (INICIAL)

Modelo de regressão	Variável dependente (Y)	Variáveis independentes (X)	R2 ajustado	P-valor (Modelo F)	Coefficiente mais relevante (sinal)	Conclusão principal
Regressão 1	Pontuação total de maturidade ética	Orçamento institucional	0,098	0,198	Orçamento (+0,0076)	Não significativo. O recurso financeiro não é um preditor estatisticamente robusto da maturidade ética geral.
Regressão 2	Pontuação de transparência	Salário presidente e nº de funcionários	-0,084	0,59	Nenhum	Não significativo. O capital humano quantitativo (tamanho da equipe e remuneração) não explica o rigor no princípio de Transparência/Explicabilidade.
Regressão 3	Pontuação de Robustez e Segurança	Tempo médio de permanência e % de experiência > 5 anos	0,347	0,187	Tempo médio de permanência (+0,076)	Tendência forte. A estabilidade do corpo técnico (tempo de permanência) apresenta a correlação positiva mais forte ($R^2 = 0,347$) com o rigor técnico/operacional, embora a amostra reduzida ($n = 7$) impeça a significância formal.

Fonte: Dados da pesquisa.

A análise estatística inicial das regressões (Regressões 1 e 2), que indicou a falta de significância dos recursos financeiros (orçamento) e humanos (número de funcionários) na maturidade ética individual das jurisdições, não implica a irrelevância desses fatores. Pelo contrário, essa dissonância sugere que a eficácia ética é um desafio qualitativo e cultural, e não meramente quantitativo.

As correlações regionais transformaram as relações que eram estatisticamente insignificantes no nível individual (Regressões 1 e 2) em fortemente positivas no nível agregado. Isso sugere que os fatores institucionais (orçamento, pessoal e estabilidade) realmente influenciam a maturidade ética, mas esse efeito só se manifesta nitidamente quando as diferenças de recursos e cultura (as regiões) são contrastadas, e não quando a variância interna das jurisdições é testada.

A correlação regional deve ser usada para contextualizar e validar a análise qualitativa, e não para substituir a análise de regressão original (que, apesar de suas limitações, é a estatística formal). Assim, as correlações fortes observadas as médias regionais ($r > 0,78$ em todos os modelos) indicam que o orçamento e o pessoal são fatores que influenciam a capacidade institucional regional de absorver os requisitos éticos, mas o efeito só se manifesta nitidamente quando as vastas diferenças de escala são contrastadas, conforme observado na Tabela 6.

TABELA 6 CORRELAÇÃO ENTRE AS REGIÕES (NÍVEL AGREGADO)

Hipótese central	Regressão original (N = 10 ou N = 7)	Coef. de correlação regional (r, N=4)	Variáveis correlacionadas	Achado principal e implicação teórica
1. Orçamento vs. Ética Total	R ² ≈ 0,098 (não significativo)	r ≈ 0,835 (forte e positiva)	Orçamento (X) vs. pontuação total ética (Y)	Dissociação estatística: o orçamento não é um preditor robusto no nível individual. A forte correlação regional sugere que o recurso financeiro é uma condição necessária regionalmente, mas sua eficácia é limitada por uma saturação da maturidade ética (o dinheiro não resolve o déficit de Explicabilidade).
2. Pessoal vs. Ética Social	R ² < 0 (não significativo)	r ≈ 0,781 (forte e positiva)	n.º funcionários (X) vs. ética social (Transparência/Justiça) (Y)	Desafio qualitativo: a insignificância estatística confirma que a Transparência/Explicabilidade não é resolvida pela simples contratação de mais pessoal. A solução é qualitativa e cultural, exigindo mandatos legais e o alinhamento da cultura de <i>accountability</i> para mitigar a vulneração (Schramm, 2008).

(Continua)

Hipótese central	Regressão original (N = 10 ou N = 7)	Coef. de correlação regional (r, N=4)	Variáveis correlacionadas	Achado principal e implicação teórica
3. Estabilidade vs. Ética Operacional	R ² ≈ 0,347 (tendência forte)	r ≈ 0,871 (muito forte e positiva)	Tempo médio de permanência (X) vs. ética operacional (Robustez/Segurança) (Y)	Achado-chave: a Estabilidade do corpo técnico é o principal preditor do cumprimento dos princípios de Não maleficência e Robustez (Floridi & Cowls). A excelência técnica é uma função direta do conhecimento institucional acumulado, superando a importância do orçamento para este pilar ético.

Fonte: Dados da pesquisa.

Como demonstrado na Tabela 6, a correlação entre orçamento e maturidade ética não se mostrou estatisticamente significativa. Nesse sentido, a evidência empírica aponta para a Estabilidade do Corpo Técnico como o preditor mais fidedigno da Não Maleficência algorítmica. Diferente do aporte financeiro isolado, o ‘capital de experiência acumulado’ (tempo de permanência) é o que permite às agências transitarem de diretrizes puramente declaratórias para mecanismos de segurança e robustez técnica institucionalizados. O achado mais expressivo é a correlação positiva e forte entre a estabilidade institucional e a pontuação em ética operacional (Robustez e Segurança), corroborando a tendência observada na Regressão 3. Jurisdições com alta permanência (como Europa e Ásia-Pacífico/África) atingem patamares elevados de Robustez, enquanto a América Latina, marcada pela menor média de permanência, apresenta o desempenho mais fragilizado.

Esse resultado indica que o domínio dos princípios técnicos – Robustez e Segurança – é fundamentalmente uma questão de conhecimento institucional acumulado. Para Floridi e Cowls (2019), a Robustez é um princípio técnico de não maleficência; sua implementação exige a retenção de talentos em *data science* e *machine learning* por longos períodos. A estabilidade do corpo técnico, portanto, é o principal preditor do cumprimento da obrigação de Não Maleficência técnica no uso da IA.

O principal déficit da amostra é global e reside nos princípios de ética social (Transparência, Responsabilidade e Justiça/Equidade), na qual nenhuma região alcança a pontuação máxima. Mesmo o grupo da América do Norte (EUA), com seu orçamento maciço, pontua apenas 3 de 5 neste quesito.

Essa lacuna crítica aponta para a falha na implementação do princípio central de Explicabilidade (Floridi & Cowls, 2019), que engloba a Inteligibilidade e a Responsabilidade. A dificuldade não é tecnológica, mas epistêmica e de governança: a opacidade algorítmica (déficit de Transparência) impede que o regulado (seja consumidor ou empresa) compreenda *por que* foi alvo de uma investigação ou sanção. Tal opacidade, no contexto da regulação econômica, transforma a vulnerabilidade universal em uma vulneração concreta (Schramm, 2008). A falta de Justiça e Equidade (3 é a pontuação mais baixa) impede o exercício do direito de defesa e do contraditório, configurando um dano real e constatável o qual a BP exige que a autoridade reguladora mitigue ativamente (Schramm, 2017).

A insignificância da Regressão 2 (Pessoal vs. Transparência) confirma que resolver essa “crise da explicabilidade” exige mais do que simplesmente contratar mais funcionários; requer uma

mudança cultural e a implementação de mandatos legais que forcem a *Human-in-the-Loop* e a Responsabilidade (princípios exigidos por Floridi & COWLS, 2019) em pontos críticos do processo decisório, independentemente do volume de recursos.

Ao analisar os dados divulgados no *rating* para caracterização das organizações, podemos perceber uma interseção entre as preocupações levantadas por diversos autores e os aspectos observados no órgão regulador brasileiro. Por exemplo, ao discutir a distribuição desigual de recursos no Conselho Administrativo de Defesa Econômica (Cade), o qual possui um orçamento considerável, mas um salário relativamente baixo para seu presidente, podemos relacionar essa observação com as preocupações de autores como Crawford (2021) e Agar (2020) sobre o domínio exercido por organizações restritas no desenvolvimento e aplicação da inteligência artificial. Essa distribuição desigual pode sugerir uma possível assimetria de poder no órgão regulador brasileiro, o que pode afetar sua capacidade de lidar efetivamente com questões relacionadas à IA, incluindo a gestão do comportamento humano (Hoch & Engelmann, 2023; Roberts et al., 2021).

Além disso, a complexa estrutura organizacional do órgão regulador brasileiro, evidenciada pelo considerável contingente de colaboradores e alto número de admissões, também pode ser relacionada às preocupações levantadas por autores como Bélisle-Pipon et al. (2021) sobre a dinamicidade da tecnologia de IA e a preparação inadequada das estruturas regulatórias existentes. Essa complexidade organizacional pode demandar uma supervisão e coordenação internas mais robustas (Häußler, 2021).

Por fim, a falta de transparência em relação à representatividade feminina no órgão regulador brasileiro ecoa as preocupações de autores como Floridi et al. (2018) sobre a importância da diversidade de perspectivas na regulamentação ética da IA. A ausência dessas informações pode indicar uma lacuna na transparência e responsabilização institucional, conforme discutido por autores como Jobin et al. (2019) e Ryan e Stahl (2020), enfraquecendo a capacidade do Brasil de regular de forma efetiva o uso da IA, especialmente no que diz respeito à manipulação do comportamento humano.

Essa disparidade regional reforça que a lacuna identificada nos princípios de ética social não representa apenas um atraso administrativo, mas uma vulneração algorítmica concreta. Ao operar com sistemas de ‘caixa-preta’, as autoridades antitruste deslocam o regulado de uma condição de vulnerabilidade universal para um estado de ferimento ao direito de defesa e à inteligibilidade, tornando o déficit de explicabilidade um problema de integridade pública, e não apenas de eficiência tecnológica (Schramm, 2017). Dessa forma, a regulação econômica, sob a ótica da Bioética de Proteção, deixa de ser um instrumento meramente técnico e passa a ser uma salvaguarda necessária contra o dano algorítmico inevitável em cenários de assimetria de informação.

5. CONSIDERAÇÕES FINAIS

O presente estudo investigou os desafios e riscos éticos decorrentes da aplicação de sistemas de IA em 35 órgãos de defesa da concorrência, utilizando o quadro ético (Floridi & COWLS, 2019). A análise demonstrou que o avanço ético na regulação econômica é regido mais por fatores qualitativos e estruturais do que por recursos financeiros isolados. O achado mais robusto é a correlação entre a estabilidade institucional e o rigor na ética operacional, sugerindo que a prevenção de falhas algorítmicas é um ativo de conhecimento acumulado pela retenção de talentos. Entretanto, a lacuna universal nos princípios de ética social (Transparência e Responsabilidade) revela um déficit crítico de explicabilidade. Sob a lente da bioética de proteção de Schramm (2008), essa opacidade deixa de

ser um desafio técnico para se tornar um risco de vulneração algorítmica, comprometendo o direito de defesa e a integridade do mercado.

5.1 Implicações para a administração pública e regulação

As implicações deste estudo para a administração pública são diretas: a governança da IA no setor público não deve ser tratada como um problema de aquisição de tecnologia, mas de desenvolvimento institucional. A forte correlação entre tempo de permanência e segurança tecnológica indica que políticas de retenção de servidores técnicos são garantidoras de uma IA mais robusta e segura, o que auxilia a área de gestão nas organizações públicas a adotar melhores práticas de retenção de talentos e conhecimento especializado.

Além disso, para evitar a vulneração dos administrados, órgãos públicos devem transitar da “IA de caixa-preta” para sistemas auditáveis. A implicação prática é a necessidade de criar mecanismos de *accountability* que permitam a supervisão humana efetiva e a inteligibilidade dos processos decisórios automatizados.

5.2 Limitações e recomendações

Apesar do ineditismo do diagnóstico apresentado, este desenho de pesquisa possui limitações que devem ser ponderadas para a interpretação dos resultados. Primeiramente, o corpus documental restringiu-se a documentos públicos, o que pode evidenciar mais as intenções normativas das agências do que suas práticas cotidianas efetivas, incorrendo no risco de camuflar lacunas operacionais sob um verniz de conformidade (*ethics washing*). Em segundo lugar, o uso de dados transversais e o tamanho reduzido da amostra em certas clivagens regionais impedem a afirmação de causalidade absoluta, permitindo apenas a identificação de tendências estatísticas e correlações exploratórias.

Para que a integração da IA na regulação econômica avance de forma ética e mitigue o risco de vulneração algorítmica, recomendam-se as seguintes ações institucionais: (1) estabelecer a inteligibilidade algorítmica como requisito mandatório de governança, assegurando o direito de defesa dos administrados; (2) priorizar políticas de retenção das equipes técnicas para garantir a estabilidade do conhecimento necessário à ética operacional; e (3) integrar especialistas em ética, direito e ciências sociais no ciclo de desenvolvimento dos sistemas para equilibrar o rigor técnico com a responsabilidade social.

Quanto à agenda de pesquisas futuras, sugere-se o monitoramento da evolução do *Enforcement Rating* da GCR em um horizonte de três a cinco anos, visando mensurar o impacto real de marcos como o *AI Act* europeu na maturidade das autoridades. Adicionalmente, recomenda-se a realização de estudos qualitativos baseados em entrevistas com gestores para desvelar barreiras culturais à transparência, bem como o desenvolvimento de índices que quantifiquem o custo jurídico e social da vulneração algorítmica, transpondo o debate da ética teórica para uma intervenção regulatória baseada em evidências sólidas.

REFERÊNCIAS

- Agar, N. (2020). *The sceptic and the AI society*. Palgrave Macmillan.
- Arner, D. W., Buckley, R. P., & Zetzsche, D. A. (2021). Finance, technology, and regulation: From hedgehogs to foxes. *European Business Organization Law Review*, 22(1), 11-40. <https://doi.org/10.1007/s40804-020-00201-9>
- Babbie, E. (2016). *The Practice of Social Research* (14th ed.). Cengage Learning. <https://doi.org/10.4324/9781315518014>
- Bani-Salameh, H., Sallam, M., & AlShboul, B. (2021). A deep-learning-based bug priority prediction using RNN-LSTM neural networks. *e-Informatica Software Engineering Journal*, 15(1), 29-45. <https://doi.org/10.37190/e-Inf210102>
- Beauchamp, T. L., & Childress, J. F. (2001). *Principles of biomedical ethics* (5th ed.). Oxford University.
- Bélisle-Pipon, J. C., Couture, V., Roy, M. C., Ganache, I., Goetghebeur, M., & Cohen, I. G. (2021). What makes artificial intelligence exceptional in health technology assessment? *Frontiers in Artificial Intelligence*, 4, 736697. <https://doi.org/10.3389/frai.2021.736697>
- Biondi, G. M. C. B., & Cernev, A. K. (2023). Nuveo: ética digital e inteligência artificial para desafios do mundo real. *Revista de Administração Contemporânea*, 27(3), e220063. <https://doi.org/10.1590/1982-7849rac2023220063>
- Bodrick, M., Alqarni, H., Alsuhaime, M., & Almuways, Y. S. (2024). *Critical appraisal of definitions on intelligence within the organizational context*. *Journal of Learning and Development Studies*, 4(2), 12-20. <https://doi.org/10.32996/jlds.2024.4.2.2>
- Bostrom, N., & Yudkowsky, E. (2011). The ethics of artificial intelligence. In K. Frankish & W. M. Ramsey (Eds.), *The Cambridge handbook of artificial intelligence* (pp. 316-334). Cambridge University.
- Cave, S., Heigeartaigh, S. O. (2019) An AI Race for Strategic Advantage: Rhetoric and Risks. In *Proceedings of the 2019 AAAI/ACM Conference on AI, Ethics, and Society* (AIES '19). Association for Computing Machinery. <https://doi.org/10.1145/3278721.3278780>
- Coeckelbergh, M. (2020). *AI ethics*. MIT Press. <https://doi.org/10.7551/mitpress/12549.001.0001>
- Crawford, K. (2021). *Atlas of AI: power, politics, and the planetary costs of artificial intelligence*. Yale University.
- Etzioni, A., & Etzioni, O. (2017). *Incorporating ethics into artificial intelligence*. *The Journal of Ethics*, 21(4), 403-418. <https://doi.org/10.1007/s10892-017-9252-2>
- Ezrachi, A., & Stucke, M. E. (2019). *Virtual Competition: The Promise and Perils of the AlgorithmDriven Economy*. Harvard University Press. <https://doi.org/10.4159/9780674241596>
- Fjeld, J., Achten, N., Hilligoss, H., Nagy, A., & Srikumar, M. (2020). *Principled artificial intelligence: Mapping consensus in ethical and rights-based approaches to principles for AI*. Berkman Klein Center for Internet & Society. <https://doi.org/10.2139/ssrn.3518482>
- Floridi, L., & Cows, J. (2019). A unified framework of five principles for AI in society. *Harvard Data Science Review*, 1(1). <https://doi.org/10.1162/99608f92.8cd550d1>
- Floridi, L., Cows, J., Beltrametti, M., Chatila, R., Chazerand, P., Dignum, V., Luetge, C., Madelin, R., Pagallo, U., Rossi, F., Schafer, B., Valcke, P., & Vayena, E. (2018). AI4People – An ethical framework for a good AI society: opportunities, risks, principles, and recommendations. *Minds and Machines*, 28, 689-707. <https://doi.org/10.1007/s11023-018-9482-5>
- Greene, D., Hoffmann, A. L., & Stark, L. (2019). Better, nicer, clearer, fairer: A critical assessment of the movement for ethical artificial intelligence and machine learning. In *Proceedings of the 52nd Hawaii International Conference on System Sciences*. <https://doi.org/10.24251/hicss.2019.258>
- Hagendorff, T. (2020). *The ethics of AI ethics: An evaluation of guidelines*. *Minds and Machines*, 30(1), 99-120. <https://doi.org/10.1007/s11023-020-09517-8>
- Häußler, H. (2021). The underlying values of data ethics frameworks: A critical analysis of discourses and power structures. *Libri*, 71(4), 307-319. <https://doi.org/10.1515/libri-2021-0095>
- Hoch, P. A., & Engelmann, W. (2023). Regulação da inteligência artificial no Judiciário brasileiro e

- européu. *Pensar – Revista de Ciências Jurídicas*, 28(4), 1-18. <https://doi.org/10.5020/2317-2150.2023.14263>
- Jobin, A., Ienca, M., & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1, 389-399. <https://doi.org/10.1038/s42256-019-0088-2>
- Mayson, S. G. (2020). The limits of machine learning in adjudicating employment discrimination claims. *California Law Review*, 108(2), 387-442.
- Motta, M. (2020). Competition policy in the digital era. *Economic Journal*, 130(629), 1999-2023.
- Organization for Economic Co-operation and Development. (2021). *Unlocking the potential of artificial intelligence in competition policy* [OECD Digital Economy Papers, No. 306]. <https://doi.org/10.1787/000306-en>
- Raji, I. D., & Dobbe, R. (2023). Concrete problems in AI safety, revisited. *arXiv preprint arXiv:2401.10899*. <https://arxiv.org/abs/2401.10899>
- Roberts, H., Cows, J., Morley, J., Taddeo, M., Wang, V., & Floridi, L. (2021). The chinese approach to artificial intelligence: an analysis of policy, ethics, and regulation. *AI & Society*, 36, 59-77. <https://doi.org/10.1007/s00146-020-00992-2>
- Ryan, M., & Stahl, B. C. (2020). Artificial intelligence ethics guidelines for developers and users: Clarifying their content and normative implications. *Journal of Information, Communication and Ethics in Society*, 19(1), 61-86. <https://doi.org/10.1108/JICES-12-2019-0138>
- Schmude, T., Koesten, L., Möller, T., & Tschischek, S. (2023, June). On the impact of explanations on understanding of algorithmic decision-making. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, Nova York, NY, EUA. <https://doi.org/10.1145/3593013.3594054>
- Schramm, F. R. (2008). Bioética da proteção: ferramenta válida para enfrentar problemas morais na era da globalização. *Revista Bioética*, 16(1), 11-23. https://revistabioetica.cfm.org.br/revista_bioetica/article/view/52
- Schramm, F. R. (2017). A bioética de proteção: uma ferramenta para a avaliação das práticas sanitárias? *Ciência & Saúde Coletiva*, 22, 1531-1538. <https://doi.org/10.1590/1413-81232017225.04532017>
- Schrepel, T. (2021). *Computational antitrust: An introduction and research agenda*. Stanford University CodeX Blockchain Group Research Paper. <http://dx.doi.org/10.2139/ssrn.3766960>
- Singapore Government. (2019). *Model AI governance framework*. <https://www.pdpc.gov.sg/Help-and-Resources/2020/01/Model-AI-Gov-Framework>
- Smuha, N. A. (2021). From a ‘race to AI’ to a ‘race to AI regulation’: regulatory competition for artificial intelligence. *Law, Innovation and Technology*, 13(1), 57-84. <https://doi.org/10.1080/17579961.2021.1898300>
- Stahl, B. C. (2021). *Artificial Intelligence for a Better Future: An Ecosystem Perspective on the Ethics of AI and Emerging Digital Technologies*. Springer.
- Veale, M., & Borgesius, F.Z. (2021). Demystifying the draft EU Artificial Intelligence Act – analysing the good, the bad, and the unclear elements of the proposed approach. *Computer Law Review International*, 22(4), 97-112. <https://doi.org/10.9785/cri-2021-220402>
- von Ingersleben-Seip, N. (2023). Competition and cooperation in artificial intelligence standard setting: explaining emergent patterns. *Review of Policy Research*, 40(5), 781-810. <https://doi.org/10.1111/ropr.12538>
- Zekos, G. I. (2024). Digital economy and competition. In *Artificial intelligence and competition: economic and legal perspectives in the digital age* (pp. 293-359). Springer.
- Zeng, Y., Lu, E., & Huangfu, C. (2018). *Linking artificial intelligence principles*. arXiv. <https://doi.org/10.48550/arXiv.1812.04814>

Mayla Cristina Costa Maroni Saraiva

Doutora em Administração pela Universidade Positivo; Professora associada na Universidade de Brasília (UnB). E-mail: mayla.saraiva@unb.br

Fátima de Souza Freire

Doutora em Economia pela Université des Sciences Sociales Toulouse I; Professora titular na Universidade de Brasília (UnB). E-mail: ffreire@unb.br

CONTRIBUIÇÃO DAS AUTORAS

Mayla Cristina Costa Maroni Saraiva: Análise formal (Liderança); Investigação (Liderança); Metodologia (Liderança); Administração do projeto (Suporte); Escrita – rascunho original; Escrita – revisão e edição (Igual).

Fátima de Souza Freire: Aquisição de financiamento; Investigação (Suporte); Metodologia (Suporte); Administração de projeto (Liderança); Recursos (Liderança); Escrita – revisão e edição (Igual).

DISPONIBILIDADE DE DADOS DA PESQUISA

Os dados não estão disponíveis publicamente devido a restrições institucionais e éticas, podendo ser disponibilizados pelos autores mediante solicitação formal.

CONFLITO DE INTERESSE

Os autores não possuem conflitos de interesse.

USO DE INTELIGÊNCIA ARTIFICIAL

A ferramenta de inteligência artificial Gemini (Google) foi utilizada estritamente para suporte linguístico, incluindo revisão técnico-gramatical e auxílio na tradução para o inglês acadêmico. Os autores assumem total responsabilidade pela integridade, precisão e originalidade do conteúdo final do manuscrito.

FINANCIAMENTO

Agradecemos ao financiamento da Fundação de Apoio à Pesquisa do Distrito Federal (FAP/DF), Edital n. 25/2023 – FAP/DFPRES/GAB, Chamada n. N. 03/2023 GOV LEARNING, vinculado ao Edital n. 10/2023 – Programa FAPDF Learning.

AGRADECIMENTOS

As autoras expressam seu reconhecimento aos pareceristas anônimos da RAP pelas sugestões e críticas construtivas, que foram fundamentais para o aprimoramento e fortalecimento deste manuscrito. Agradecemos igualmente à equipe editorial pelo suporte e pela condução do processo de revisão. Agradecemos ao CADE, pelo acesso aos dados.