

Contribuições - Prado Vidigal Advogados

Carolina Giovanini | PVA <carolina@pradovidigal.com.br>

sex 10/06/2022 12:40

Para: CJSUBIA <CJSUBIA@senado.leg.br>;

Cc: Paulo Vidigal | PVA <paulo@pradovidigal.com.br>; Luis Fernando Prado | PVA <luis@pradovidigal.com.br>;

 1 anexo

Contribuições_Prado Vidigal Advogados_062022.pdf;

Você não costuma receber emails de carolina@pradovidigal.com.br. [Saiba por que isso é importante](#)

À Comissão de Juristas responsável por subsidiar elaboração de substitutivo sobre inteligência artificial no Brasil (CJSUBIA)

A/C Presidente Ministro Ricardo Villas Boas Cueva

C/C Relatora Prof. Dra. Laura Schertel Mendes

Conforme chamamento público divulgado pelo Senado Federal, para prestar contribuição aos trabalhos desta Comissão de Juristas, encaminhamos as contribuições do escritório Prado Vidigal Advogados para a regulação de inteligência artificial.

Permanecemos à disposição para quaisquer esclarecimentos e complementações.

Cordialmente,



Este e-mail contém informação confidencial. Caso tenha recebido por engano, avise a pessoa remetente e apague-o imediatamente.

This message contains privileged and confidential information. If you received it in error, please notify the sender by replying to this email and delete it immediately.

Contribuições para a Comissão de Juristas do Senado Federal responsável elaborar o marco normativo sobre inteligência artificial

Sumário

Sumário.....	1
1. Premissas da regulação.....	1
2. Inteligência artificial e riscos	3
2.1. Riscos para IA: possibilidades de classificação e gradação	3
2.2. Da inaplicabilidade de previsão legislativa de riscos inaceitáveis	5
3. Transparência: parâmetros para fornecimento de informações e possibilidades de implementação de medidas voltadas ao público	7
4. Explicabilidade e revisão	8
4.1. Registro de metodologias utilizadas e parâmetros para explicação ao público	8
4.2. Revisão de processos realizados por sistemas de IA e intervenção humana	9
5. Avaliações de impacto.....	10
6. Desafios para a responsabilização	12
7. Considerações finais.....	13

1. Premissas da regulação

Considerado a crescente utilização de sistemas de inteligência artificial (IA) em diferentes setores (como, saúde, trabalho, justiça, transporte, vendas etc.), as particularidades do funcionamento de cada tipo de tecnologia e o fato de que a criação de sistemas de IA e seus respectivos usos, potencialidades e limitações no contexto brasileiro ainda são aspectos em desenvolvimento, entendemos que a regulação deve partir das seguintes premissas: (i) elaboração de uma norma geral e principiológica; (ii) fundamentada na promoção do desenvolvimento econômico e tecnológico; e (iii) pautada em uma abordagem baseada no risco.

Em primeiro lugar, faz-se necessário ressaltar que, dado o momento marcado pelo recente desenvolvimento do uso de sistemas de IA em diversos setores, entendemos que eventual

regulação afetará diretamente o desenvolvimento tecnológico de diversas iniciativas e, por tal razão, deve ter características de norma geral, prevendo princípios que orientarão a criação e o uso de tais sistemas.

A partir da elaboração de uma norma geral, seria possível que cada setor estabelecesse normas e orientações específicas e de acordo com suas particulares, seja por meio de instrumentos de correção, seja via regulamentação setorial por parte de autoridades competentes. Em síntese, é necessário reconhecer as particulares das diversas áreas de aplicação de sistemas de IA, evitando a unificação em um modelo jurídico engessado.

Além disso, entendemos que os fundamentos da regulação devem abranger a promoção da inovação, da livre iniciativa e da livre concorrência, de modo que a proteção da pessoa humana e a inovação pautada no desenvolvimento econômico e tecnológico convivam de forma harmônica na futura norma e não sejam vistas como um *trade-off*. É essencial que as normas que regularão esta matéria não afetem negativamente os rumos do desenvolvimento econômico, especialmente diante da escolha legislativa sobre quem deverá suportar os riscos decorrentes do uso de sistemas de IA¹.

Por fim, consideramos que uma das premissas da regulação deve ser a abordagem baseada no risco, de modo a incentivar condutas pautadas em precaução e transparência. A partir da precaução, em situações nas quais existam ameaças de danos graves ou irreversíveis, ainda que em potencial, é necessário tomar medidas de proteção sem esperar que esses riscos se tornem plenamente aparentes².

Nesse sentido, é possível observar que a avaliação de risco e o princípio da precaução andam juntos, pois são instrumentos que determinam conjuntamente a criticidade das atividades, produtos ou serviços e o custo dos potenciais danos identificados. A abordagem baseada no risco contribui para a lógica da precaução na medida em que o risco sinaliza a ameaça de dano de modo mensurável e pode se tornar uma ferramenta para a tomada de decisões, possibilitando que eventos futuros sejam gerenciados para se tornarem certos e controláveis³.

¹ GALASSO, Alberto; LUO, Hong. Punishing Robots: Issues in the Economics of Tort Liability and Innovation in Artificial Intelligence, In: **The Economics of Artificial Intelligence: An Agenda**, p.502. National Bureau of Economic Research. Disponível em: <https://www.nber.org/books-and-chapters/economics-artificial-intelligence-agenda/punishing-robots-issues-economics-tort-liability-and-innovation-artificial-intelligence>.

² COSTA, Luiz. Privacy and the precautionary principle. **Computer Law & Security Review**, [s. l], v. 28, n. 1, p. 14-24, fev. 2012. Disponível em: <https://www.sciencedirect.com/science/article/abs/pii/S0267364911001804?via%3Dihub>.

³ GELLERT, Raphael. Data protection: a risk regulation?: between the risk management of everything and the precautionary alternative. *International Data Privacy Law*, [s. l], v. 5, n. 1, p. 3-19, mar. 2015. Disponível em: <https://academic.oup.com/idpl/article-abstract/5/1/3/622981?redirectedFrom=fulltext>.

2. Inteligência artificial e riscos

2.1. Riscos para IA: possibilidades de classificação e gradação

Em abril de 2021, a Comissão Europeia - instituição da União Europeia competente para propor normas regulatórias – apresentou a proposta de Regulamento sobre inteligência artificial (*Artificial Intelligence Act* – COM/2021/206)⁴, estabelecendo regras para sistemas de inteligência artificial, provedores e usuários. A lógica regulatória está pautada na abordagem baseada no risco, pois as obrigações surgem de acordo com os usos e riscos associados ao sistema de IA em questão.

Assim, a proposta europeia distingue entre (i) risco inaceitável; (ii) risco alto; e (iii) risco baixo ou mínimo. O objetivo de tal estratégia é trazer tratamentos normativos distintos para os diferentes tipos de aplicação de IA existentes, evitando que usos que ofereçam risco baixo ou mínimo sejam afetados ou limitados excessivamente.

Como forma de sintetizar a classificação de riscos realizada pela proposta de Regulamento, apresentamos a tabela abaixo:

Risco	Descrição	Observações
Risco inaceitável	Sistemas de IA cuja utilização seja considerada inaceitável por violar os valores da União Europeia.	São sistemas proibidos. Abrangem práticas com potencial significativo para manipular as pessoas por meio de técnicas de difícil percepção ou para explorar as vulnerabilidades de grupos específicos, distorcendo o seu comportamento de forma que seja suscetível de causar danos psicológicos ou físicos. A proposta também proíbe a classificação social para uso geral por parte das autoridades públicas, bem como a utilização de sistemas de identificação biométrica à distância em tempo real em espaços acessíveis ao público para efeitos de manutenção da ordem pública, exceto em caso de cumprimento de obrigação legal específica.

⁴ UNIÃO EUROPEIA. Proposal for a Regulation of the European Parliament and the Council Laying Down Harmonized Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts (COM/2021/206). Bruxelas: Comissão Europeia, 2021.

Risco alto	Criam risco elevado para saúde, segurança ou direitos fundamentais de pessoas naturais.	São sistemas autorizados, mas estão sujeitos ao cumprimento de determinados requisitos obrigatórios e a uma avaliação da conformidade <i>ex ante</i> . A proposta identifica duas categorias principais de sistemas de IA de risco elevado: (i) sistemas de IA concebidos para serem utilizados como componentes de segurança de produtos que estão sujeitos a uma avaliação da conformidade <i>ex ante</i> por terceiros; ou (ii) outros sistemas de IA autônomos com implicações em matéria de direitos fundamentais que são previstos no anexo III da referida proposta.
Risco baixo ou mínimo	Criam risco limitado ou mínimo para pessoas físicas (ex: <i>chatbots</i>).	Tais sistemas são autorizados e dispensados da maioria das obrigações previstas na proposta. No caso de sistemas de risco baixo, é necessário cumprir com as obrigações de transparência e, no caso de sistemas de risco mínimo, não há restrições por parte do Regulamento.

Desse modo, observa-se que uma possível estratégia regulatória é a vinculação da intensidade de intervenção legislativa ao nível de ameaça aos direitos fundamentais e à segurança de pessoas naturais, adotando uma classificação de riscos.

Entendemos que a regulação a partir de classificação e gradação de riscos é uma estratégia importante e passível de aplicação, pois afasta uma solução jurídica geral e abstrata. No entanto, ao contrário da classificação de riscos em debate na União Europeia – que elenca especificamente usos de sistemas de IA e vincula determinados setores para cada nível de risco – entendemos que tal classificação deve ser feita a partir da aferição de probabilidade e impacto de eventuais danos aos direitos de pessoas naturais, conforme construção abaixo:

Risco	Descrição	Fatores para avaliação
Risco alto	Sistemas de IA que muito provavelmente irão gerar	• Uso de dados pessoais sensíveis⁵;

⁵ Art. 5º, II, da LGPD: dado pessoal sensível: dado pessoal sobre origem racial ou étnica, convicção religiosa, opinião política, filiação a sindicato ou a organização de caráter religioso, filosófico ou político, dado referente à saúde ou à vida sexual, dado genético ou biométrico, quando vinculado a uma pessoa natural;

	consequências graves e, até mesmo, irreversíveis para pessoas naturais (ex: danos físicos, impedimento de exercício de direitos etc.).	• Público-alvo em situação de vulnerabilidade (ex: crianças, adolescentes, idosos, grupos minoritários etc.); • Uso para classificação, perfilização, monitoramento sistemático e/ou tomada de decisões automatizadas; • Uso para gerenciamento e operação de infraestrutura crítica ou voltada para componentes de segurança de determinado produto/serviço; • Acesso aos serviços públicos e direitos fundamentais; • Administração da justiça e processos democráticos.
Risco médio	Sistemas de IA que possuem potencial para gerar consequências significativas, que podem gerar dificuldades relevantes e impactar direitos de pessoas naturais (ex: danos materiais).	
Risco baixo	Sistemas de IA que excepcionalmente podem afetar direitos de pessoas naturais de forma mínima (ex: aborrecimentos e entraves), porém, não impactam o exercício de direitos e liberdades fundamentais.	

Em síntese, entendemos que, embora os níveis de risco e os parâmetros para gradação possam estar previstos em lei, a classificação de risco não deve estar determinada previamente na norma por meio de vinculação a determinados usos (ex: identificação biométrica) ou a determinadas áreas (ex: educação e formação profissional), sob pena de fatores contextuais relevantes serem desconsiderados.

2.2. Da inaplicabilidade de previsão legislativa de riscos inaceitáveis

O debate acerca da regulação de tecnologias passa por correntes que defendem a identificação de determinadas atividades que deveriam ser banidas, em razão de não serem suficientemente precisas e maduras. No entanto, entendemos que, para efetivar a proteção dos direitos de pessoas naturais, a estratégia regulatória pautada no banimento de tecnologias não é a ideal, pois a proibição do desenvolvimento de produtos, serviços e tecnologias não é compatível com os valores de desenvolvimento econômico e livre iniciativa.

Nesse sentido, destaca-se o art. 170 da Constituição Federal de 1988:

Art. 170. A ordem econômica, fundada na valorização do trabalho humano e na livre iniciativa, tem por fim assegurar a todos existência digna, conforme os ditames da justiça social (...).

Parágrafo único. É assegurado a todos o livre exercício de qualquer atividade econômica, independentemente de autorização de órgãos públicos, salvo nos casos previstos em lei.

Evidentemente, o uso de sistemas de IA e novas tecnologias deve ser acompanhado de determinados controles a depender do nível de risco oferecido, porém, tais controles não devem impedir que agentes do mercado testem tecnologias, ainda que em ambientes controlados (como *sandboxes*).

É importante notar que, embora existam movimentações favoráveis ao banimento de determinadas tecnologias, há de se reconhecer que uma política de proibição indefinida provavelmente envolveria custos de oportunidade significativos para a sociedade como um todo, sendo necessário um debate público substancial para estabelecer as condições nas quais as tecnologias podem ser implementadas apropriadamente e quais soluções técnicas podem ser utilizadas diante de problemas de precisão e parcialidade⁶.

A título de exemplificação, o Fórum Econômico Mundial, ao tratar da busca por equilíbrio no uso de tecnologias de reconhecimento facial, esclarece que a tecnologia possui o potencial de ajudar a resolver, impedir e prevenir crimes, sendo necessário estabelecer parâmetros para definir o que constitui o uso responsável do reconhecimento facial⁷.

No mais, há de se reconhecer que propostas regulatórias que de saída trazem “riscos inaceitáveis” e o consequente banimento do uso da tecnologia em nada contribuem para a evolução tecnológica, a qual passa pelo aperfeiçoamento de seu uso e mitigação de riscos conhecidos. Inclusive, referidas propostas partem da premissa de se regular com base no estado da tecnologia que se conhece hoje, sem reconhecer os esforços da sociedade para que riscos hoje existentes sejam mitigados amanhã. Em outras palavras, ainda que, no momento atual, determinado uso de sistemas de AI gere elevado risco à sociedade (a ponto de ser considerado por lei como “inaceitável”), nada impede que, no futuro breve, a tecnologia evolua a ponto de reduzir referido risco, sendo que tal evolução somente ocorrerá caso o Direito abra espaço para tanto. Por fim, importante lembrar que a estratégia regulatória de não barrar em lei determinados usos de sistemas de AI em nada afasta a incidência das consequências civis e penais aos responsáveis pelos males causados à sociedade.

⁶ ERLER, Alexander. Regulate or Ban? Artificial Intelligence Will Lead to Tough Choices. Green European Journal. Disponível em: <https://www.greeneuropeanjournal.eu/regulate-or-ban-artificial-intelligence-will-lead-to-tough-choices/>.

⁷ MADZOU, Lofred *et al.* This is best practice for using facial recognition in law enforcement. Disponível em: <https://www.weforum.org/agenda/2021/10/facial-recognition-technology-law-enforcement-human-rights/>.

3. Transparência: parâmetros para fornecimento de informações e possibilidades de implementação de medidas voltadas ao público

Para que os usuários de sistemas de IA e terceiros interessados consigam efetivamente exercer eventuais controles e compreender o funcionamento de tais sistemas, é essencial que informações claras, precisas e de fácil acesso sejam prestadas.

É necessário considerar que, atualmente, os usuários são bombardeados com diversas informações sobre diferentes aspectos de produtos/serviços e, conseqüentemente, acabam não se interessando por leituras extensas e técnicas referentes aos termos e políticas disponibilizados. Tendo em vista o contexto mencionado, é importante que a norma forneça margem para que os agentes de mercado criem e utilizem diferentes estratégias de comunicação, como “*just-in-time notices*” (mensagens curtas que aparecem no momento do preenchimento de informações), abordagens em camadas, FAQs, recursos gráficos e vídeos.

A partir dos critérios apresentados na proposta de Regulamento sobre inteligência artificial (*Artificial Intelligence Act* – COM/2021/206)⁸ e nos critérios já presentes na Lei Geral de Proteção de Dados (Lei nº 13.709/2018, abreviada por “LGPD”, art. 9º), entendemos que os seguintes tópicos devem ser endereçados para fins de fornecimento de informações ao público:

- Identificação clara do agente responsável pelo desenvolvimento do sistema de IA;
- Identificação clara do agente responsável pelo fornecimento do sistema de IA;
- Disponibilização de informações sobre o canal de atendimento de solicitações;
- Condições gerais de uso do sistema de IA;
- Informações sobre deveres e vedações aos usuários do sistema de IA; e
- Informações sobre situações de risco e interrupção do funcionamento do sistema de IA.

⁸ UNIÃO EUROPEIA. Proposal for a Regulation of the European Parliament and the Council Laying Down Harmonized Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts (COM/2021/206). Bruxelas: Comissão Europeia, 2021.

4. Explicabilidade e revisão

4.1. Registro de metodologias utilizadas e parâmetros para explicação ao público

O direito à explicação está intrinsecamente relacionado à transparência⁹, pois possibilita o entendimento acerca do funcionamento de sistemas de IA, o que poderá permitir que sejam contestados e/ou revisados¹⁰.

Nesse ponto, vale destacar que a Lei Geral de Proteção de Dados (Lei nº 13.709/2018) prevê, em seu art. 20, §1º, que o controlador deverá fornecer, sempre que solicitadas, informações claras e adequadas a respeito dos critérios e dos procedimentos utilizados para tomada de decisão automatizada, observados os segredos comercial e industrial. Além disso, a Lei do Cadastro Positivo (Lei nº 12.414/2011) prevê, em seu art. 5º, IV, que é direito do cadastrado conhecer os principais elementos e critérios considerados para a análise de risco, resguardado o segredo empresarial.

Desse modo, é possível observar que o direito à explicação – em certa medida – já foi introduzido no ordenamento jurídico brasileiro e, em matéria de inteligência artificial, é essencial para construir relações de confiança entre usuários, desenvolvedores e fornecedores, bem como para impactar positivamente a forma como produtos e serviços que utilizam sistemas de IA são desenvolvidos.

Nesse sentido, entendemos que a regulação, ao prever um direito à explicação sobre o funcionamento de sistemas de IA, deverá traçar parâmetros e limitações para o fornecimento de informações referentes a metodologia, critérios e bancos de dados utilizados. A previsão de ressalva acerca da divulgação de segredo comercial e industrial é essencial para a proteção da propriedade intelectual, uma vez que o desenvolvimento de sistemas de IA demandam investimentos em P&D.

Ademais, é necessário considerar que o fornecimento ilimitado de informações acerca do funcionamento de determinado sistema de IA pode representar uma violação à segurança dos próprios usuários, além de ameaça ao instituto do segredo de negócio. Uma vez que as informações são disponibilizadas para terceiros, cria-se um elo frágil no sistema, com a possibilidade de vazamentos de detalhes que carregam o potencial de colocar em risco a segurança e robustez dos sistemas de IA¹¹.

⁹ SMITH, L. (2016). "Algorithmic transparency: Examining from within and without". IAPP Privacy Perspectives. Disponível em: <https://iapp.org/news/a/algorithmic-transparency-examining-from-withinand-without/>.

¹⁰ SELBST, A. D.; POWLES, J. (2017). "Meaningful information and the right to explanation". International Data Privacy Law, vol. 7, nº 4, p. 233-242. Disponível em: <https://ssrn.com/abstract=3039125>.

¹¹ GUTIERREZ, Andriei. É possível confiar em um sistema de Inteligência Artificial? Práticas em torno da melhoria da sua confiança, segurança e evidências de accountability. In: FRAZÃO, Ana; MULHOLLAND, Caitlin (Org.). Inteligência Artificial e Direito. São Paulo: Revista dos Tribunais, 2019. p. 89.

Para assegurar o atendimento ao direito à explicação, entendemos que seria interessante que se avaliasse a utilidade e relevância de se instituir, via regulação, obrigações de registro de aspectos gerais dos sistemas de IA, a exemplo daqueles dispostos a seguir, com a ressalva de que quaisquer informações levadas registros devem ser suficientemente amplas, preservando os detalhes e especificidades dos sistemas, em estrita observância aos segredos comercial e industrial:

- Identificação do tipo de algoritmo e abordagem/técnica utilizada;
- Registro das características do(s) banco(s) de dado(s) utilizados;
- Registro da metodologia e critérios adotados (parâmetros e variáveis utilizados e influência de tais parâmetros e variáveis)¹².

4.2. Revisão de processos realizados por sistemas de IA e intervenção humana

A LGPD prevê, em seu art. 20, que o titular dos dados tem direito a solicitar a revisão de decisões tomadas unicamente com base em tratamento automatizado de dados pessoais que afetem seus interesses, incluídas as decisões destinadas a definir o seu perfil pessoal, profissional, de consumo e de crédito ou os aspectos de sua personalidade. Assim, observa-se que o direito à revisão já está inicialmente inserido no ordenamento jurídico brasileiro e, conforme é possível extrair da redação de tal dispositivo, não há obrigação de intervenção humana no processo de revisão.

Nessa direção, entendemos que, caso a futura regulação estabeleça um direito geral de revisão de processos realizados por sistemas de IA, não deveria haver obrigação de interferência humana em tal processo. Ocorre que o processo de decisão humana também é marcado por dois conceitos estatísticos: (i) o viés, isto é, quando previsões ou julgamentos sempre tendem para um mesmo resultado; e (ii) o ruído, ou seja, a dispersão dos resultados a partir de (conjunto de elementos aleatórios que são inconscientemente levados em consideração em um processo de análise)¹³.

Por exemplo, um médico que avalia o mesmo caso, sem saber, com diferença de meses entre a primeira e a segunda análise tem mais de 60% de chance de apresentar um diagnóstico diferente em ambas as ocasiões. Além disso, em um conjunto de indivíduos, um mesmo caso pode gerar diagnósticos e tratamentos diferentes em mais de 60% dos médicos consultados (KAHNEMAN; SIBONY; SUNSTEIN).

Assim, faz-se necessário questionar a real utilidade e os benefícios da participação humana no processo de tomada ou mesmo revisão da decisão. Observa-se que as decisões tomadas por humanos não devem ser encaradas como de alto nível de confiabilidade, pelo contrário, são passíveis de falhas já identificadas por pesquisas científicas no âmbito comportamental. Assim, o uso de algoritmos de decisão sem intervenção humana pode ser benéfico, pois o sistema irá operar

¹² CJEU. C-203/22 - Dun & Bradstreet Austria. Disponível em: <https://curia.europa.eu/juris/liste.jsf?num=C-203/22&language=en>

¹³ KAHNEMAN, Daniel; SIBONY, Olivier; SUNSTEIN, Cass. Noise: A Flaw in Human Judgment. Little, Brown Spark's, 2021.

sem influência de ruídos em sem processo decisório, minimizando possíveis erros e apresentando maior precisão.

No entanto, é importante ressaltar que, embora não vislumbremos um dever expresso de viabilizar a revisão humana, é possível prever o direito à revisão (não necessariamente humana), quando solicitada pelo titular. Nesse caso, o agente responsável deverá promover a atualização, correção ou complementação (se o caso) de algum dos fatores que estruturam o comando realizado (fonte de dados, dados e algoritmos).

Em síntese, entendemos que a necessidade ou não de revisão humana deve ficar a cargo do agente responsável em questão (avaliação caso a caso), não devendo ser imposição regulatória. Nesse sentido, sustentamos que os esforços para se garantir decisões automatizadas baseadas em dados pessoais o mais livre possível de vieses podem estar melhor empregues se concentrados na fomentação de guias e recomendações que se dirijam à fase de concepção de novos produtos ou serviços, reforçando a tônica do *privacy by design*, do que se apegados à tentativa de convencimento do legislador sobre a necessidade/importância da adoção do remédio da revisão humana, o qual, por sua vez, é acompanhado de vários efeitos colaterais para lá de indesejados¹⁴.

5. Avaliações de impacto

Partindo da premissa de que a regulação será pautada em uma abordagem baseada no risco (considerando, inclusive, classificações e gradações para os sistemas de IA), verifica-se que a gestão de riscos e a avaliação de conformidade legal são elementos centrais de avaliações de impacto. A partir de uma metodologia de gestão de risco adequada, todos os riscos possíveis e outros impactos negativos podem ser identificados e idealmente mitigados, de modo que eventuais riscos residuais são justificados e registrados¹⁵.

Os objetivos da realização de uma avaliação de impacto em sistemas de IA são:

- estabelecer limites, usos e prazos de determinado sistema;
- construir confiança entre as partes interessadas; e
- registrar o funcionamento de determinado sistema de IA e funcionar como documentação passível de responsabilizar os desenvolvedores do sistema pelos resultados de sua aplicação.

Assim, a avaliação de impacto deve considerar se o sistema é (i) robusto, isto é, seguro, de modo que não seja alterado ou comprometido; (ii) justo, evitando tratamentos discriminatórios e

¹⁴ PRADO, Luis Fernando. Algoritmos e decisões automatizadas: buscando a conformidade com a LGPD. In PALHARES, Felipe. Estudos Sobre Privacidade e Proteção de Dados. São Paulo: Thomson Reuters, 2021, p. 122-123.

¹⁵ KLOZA, Dariusz. Privacy Impact Assessments as a Means to Achieve the Objectives of Procedural Justice. Jusletter IT. Die Zeitschrift für IT und Recht, 2014. Disponível em: <https://researchportal.vub.be/en/publications/privacy-impact-assessments-as-a-means-to-achieve-the-objectives-of-procedural-justice>.

observando as consequências para grupos minoritários historicamente marginalizados; e (iii) explicável, ou seja, se o funcionamento do sistema pode ser compreendido por usuários e desenvolvedores.¹⁶

Nesse sentido, a partir de orientações e guias práticos formulado por autoridades estrangeiras, entendemos que a avaliação de impacto de determinado sistema de IA deve envolver aspectos técnicos, jurídicos e éticos/sociais, endereçando tópicos como:

- Como o sistema funciona, incluindo taxas de precisão¹⁷;
- Quais dados pessoais são utilizados pelo sistema de IA¹⁸;
- Como o sistema performa em situações adversas (*worst-case scenario*), incluindo taxas de precisão¹⁹;
- Como o sistema reage diante de diferentes tipos e naturezas de vulnerabilidades, tais como poluição de dados, indisponibilidades, relações com infra-estrutura física, ataques cibernéticos etc.²⁰;
- Verificação de que o sistema possui um plano de ação adequado para momentos adversos ou outras situações inesperadas (por exemplo, procedimentos técnicos de troca ou acionamento de um operador humano)²¹;
- Avaliação do nível de acurácia do sistema conforme contexto de utilização²²;
- Identificação de riscos para direitos e liberdades fundamentais de pessoais naturais²³;
- Medidas mitigatórias e planos de ação para endereçamentos dos riscos.

¹⁶ KOSHIYAMA, Adriano; ENGIN, Zeynep. Algorithmic Impact Assessment: Fairness, Robustness and Explainability in Automated Decision-Making. Data for Policy 2019: Digital Trust and Personal Data (Data for Policy 2019) (DFP), London, 10-12 June 2019. Disponível em: <https://zenodo.org/record/3361708#YnCBUtrMKiM>;

¹⁷ FEDERAL TRADE COMMISSION. Facing Facts: Best Practices for Common Uses of Facial Recognition Technologies. Federal Trade Commission. [S.l.]. 2012. Disponível em: <https://www.ftc.gov/sites/default/files/documents/reports/facing-facts-best-practices-common-uses-facial-recognition-technologies/121022facialtechrpt.pdf>

¹⁸ KAMINSKI, Margot e; MALGIERI, Gianclaudio. Algorithmic impact assessments under the GDPR: producing multi-layered explanations. International Data Privacy Law, [S.l.], v. 11, n. 2, p. 125-144, 6 dez. 2020. Oxford University Press (OUP). <http://dx.doi.org/10.1093/idpl/ipaa020>.

¹⁹ FEDERAL TRADE COMMISSION. Facing Facts: Best Practices for Common Uses of Facial Recognition Technologies. Federal Trade Commission. [S.l.]. 2012. Disponível em: <https://www.ftc.gov/sites/default/files/documents/reports/facing-facts-best-practices-common-uses-facial-recognition-technologies/121022facialtechrpt.pdf>

²⁰ COUNCIL OF EUROPE. Guidelines on Facial Recognition. Disponível em: <https://rm.coe.int/guidelines-on-facial-recognition/1680a134f3>

²¹ COUNCIL OF EUROPE. Guidelines on Facial Recognition. Disponível em: <https://rm.coe.int/guidelines-on-facial-recognition/1680a134f3>

²² COUNCIL OF EUROPE. Guidelines on Facial Recognition. Disponível em: <https://rm.coe.int/guidelines-on-facial-recognition/1680a134f3>

²³ EUROPEIA, C. Orientações éticas para uma IA de confiança. Grupo de peritos de alto nível sobre a inteligência artificial. Bruxelas. 2019. Disponível em: <https://eur-lex.europa.eu/legal-content/PT/TXT/PDF/?uri=CELEX:52019DC0168&from=EN#:~:text=As%20aplica%C3%A7%C3%B5es%20de%20IA%20devem%20capacitar%20os%20cidad%C3%A3os%20e%20respeitar%20o%20acesso%20de%20pessoas%20com%20defici%C3%A7%C3%A3o.&text=contributo%20para%20o%20trabalho%20do,alto%20n%C3%ADvel%20sobre%20a%20IA>

6. Desafios para a responsabilização

A utilização de sistemas de IA atrai intenso debate sobre qual deveria ser o regime de responsabilidade civil previsto pela futura regulação, especialmente em situações nas quais não há relação de consumo.

Evidentemente, o debate e as diferentes correntes surgem porque sistemas de IA possuem diferentes tipologias, com maior ou menor grau de autonomia e de delegação de responsabilidades humanas. Além disso, a cadeia de desenvolvimento e de uso de tais sistemas passa por diferentes atores (desde desenvolvedores até usuários), afetando a aferição do nexo de causalidade.

Assim, é necessário avaliar a responsabilização a partir da participação de cada agente na cadeia de desenvolvimento e utilização de sistemas de IA, modulando a responsabilização de acordo com a participação em cada cenário concreto. Conforme aponta Lopes (2020)²⁴, responsabilidade e autonomia constituem duas faces de uma mesma moeda, pois para que determinada pessoa seja obrigada a reparar um dano injusto, é fundamental que ela possua autonomia de atuação.

Nesse contexto, entendemos que não é possível prever uma solução geral e abstrata para a responsabilidade civil por danos causados por sistemas de IA. Em primeiro lugar, antes de decidir pela criação de um regime de responsabilidade específico, é necessário avaliar as normas já existentes no ordenamento jurídico brasileiro, por exemplo:

- Código de Defesa do Consumidor (Lei nº 8.078/90): responsabilidade civil objetiva pelo fato do produto e do serviço (arts. 12 e 14) e responsabilidade civil objetiva pelo vício do produto e do serviço (art. 18 e 20);
- Código Civil (Lei nº 10.406/02): responsabilidade civil subjetiva decorrente da prática de ato ilícito (art. 186 c/c art. 927);
- Código Civil (Lei nº 10.406/02): responsabilidade civil objetiva quando a atividade normalmente desenvolvida pelo autor do dano implicar, por sua natureza, risco para os direitos de outrem (art. 927, § único);
- Código Civil (Lei nº 10.406/02): responsabilidade civil por fato de terceiro, decorrente do dever de guarda, vigilância e cuidado (art. 932).

Na situação em concreto, é necessário avaliar, em primeiro lugar, se o dano causado pelo sistema de IA resulta diretamente de erro humano (como falha de dever de cuidado, defeito de fabricação, ausência ou insuficiência de orientações etc.). Nesses casos, deve-se buscar a identificação das diretrizes humanas responsáveis pelo funcionamento, por exemplo, por meio do retorno ao processo de programação, ao banco de dados utilizado etc. A aplicação de regimes já existentes de

²⁴ LOPES, Giovana Figueiredo Peluso. **INTELIGÊNCIA ARTIFICIAL (IA):** considerações sobre personalidade, imputação e responsabilidade. 2020. 145 f. Dissertação (Mestrado) - Curso de Programa de Pós-Graduação em Direito, Universidade Federal de Minas Gerais, Belo Horizonte, 2020. Disponível em: <http://hdl.handle.net/1843/34056>.

responsabilidade civil pode ser feita a uma inteligência artificial autônoma da mesma forma como o é para outra máquina ou ferramenta utilizada por um ser humano.

Por outro lado, em casos de danos causados por sistemas de IA que estejam atuando de maneira verdadeiramente autônoma (isto é, independentemente de uma diretiva humana), surge a questão acerca de quem deveria arcar com os custos da reparação: o desenvolvedor, o distribuidor, o operador ou ambos os agentes?

Diante de tal cenário, entendemos que um regime único e geral de responsabilidade civil não seria suficiente para abranger todas as possibilidades de aplicação de sistemas de IA e todas as especificidades que podem estar presentes na situação que gerou o dano a ser reparado. Portanto, as soluções para a responsabilização devem ser avaliadas de acordo com as particularidades de cada caso concreto, evitando que agentes sofram imputação de atos em relação aos quais não possuem ingerência ou participação.

7. Considerações finais

Por fim, colocamo-nos à disposição para os desdobramentos desta contribuição, certos de que a participação da sociedade na construção da norma é o melhor caminho para termos normatizações equilibradas, eficazes e plenamente exequíveis.

Responsáveis pelas contribuições:

Carolina Giovanini*

Luis Fernando Prado**

Paulo Vidigal***

PRADO VIDIGAL ADVOGADOS

**Advogada no escritório Prado Vidigal Advogados, profissional de privacidade certificada pela International Association of Privacy Professionals (CIPP/E) e mestranda em Direito e Inovação pela Universidade Federal de Juiz de Fora (UFJF).*

***Sócio do escritório Prado Vidigal Advogados, profissional de privacidade certificado pela International Association of Privacy Professionals (CIPP/E), mestre (LLM) em Direito Digital e Sociedade da Informação pela Universidade de Barcelona e especialista em Propriedade Intelectual e Novos Negócios pela FGV Direito SP.*

****Sócio do escritório Prado Vidigal Advogados, profissional de privacidade certificado pela International Association of Privacy Professionals (CIPP/E), pós-graduado em MBA em Direito Eletrônico pela Escola Paulista de Direito, especialista em Direito Processual Civil pela PUC-SP, possui extensão em Privacy by Design pela Ryerson University e em Privacidade e Proteção de Dados pela Universidade Presbiteriana Mackenzie.*