

A ADOÇÃO DE UMA ABORDAGEM BASEADA EM RISCOS PARA REGULAÇÃO DA INTELIGÊNCIA ARTIFICIAL

The adoption of a risk-based approach to artificial intelligence regulation
Revista de Direito e as Novas Tecnologias | vol. 21/2023 | Out - Dez / 2023
DTR\2023\10402

Andressa Giroto Vargas

Especialista em Direito do Estado pela Universidade Federal do Rio Grande do Sul (UFRGS).
Graduada em Ciências Jurídicas e Sociais pela Universidade Federal do Rio Grande do Sul
(UFRGS). Membro da International Association for Artificial Intelligence and Law. Servidora pública
federal. andressa.giroto@hotmail.com

Área do Direito: Digital

Resumo: O avanço exponencial da tecnologia de Inteligência Artificial (IA) na última década proporcionou inegável incremento em diferentes setores da sociedade. Todavia, à medida que sistemas dotados de IA se tornam mais complexos, suscitam questionamentos quanto aos possíveis riscos que podem ocasionar. Diante desse cenário de incertezas, discute-se qual abordagem regulatória seria mais adequada para regulá-los.

Palavras-chave: Inteligência artificial – Regulação baseada em risco – Regulação da inteligência artificial – Direito regulatório – Direito e tecnologia

Abstract: The exponential advancement of Artificial Intelligence (AI) technology in the last decade has provided undeniable improvements in different sectors of society. However, as AI-powered systems become more complex, questions are raised about the possible risks they can cause. Along with this context of uncertainties, there comes discussions about which regulatory approach would be most appropriate to regulate them.

Keywords: Artificial Intelligence - Risk based approach regulation - Artificial Intelligence regulation - Regulatory Law - Law and Technology

Para citar este artigo: VARGAS, Andressa Giroto. A adoção de uma abordagem baseada em riscos para regulação da inteligência artificial. *Revista de Direito e as Novas Tecnologias*. vol. 21. ano 6. São Paulo: Ed. RT, out.-dez. 2023. Disponível em: [inserir link consultado](#). Acesso em: DD.MM.AAAA.

Assista neste link, comentários de autora sobre este artigo

Sumário:

1. Introdução - 2. Riscos associados a sistemas dotados de IA - 3. Abordagem regulatória baseada em riscos - 4. Conclusão - 5. Referências - 6. Legislação

1. Introdução

A despeito das oportunidades que a IA possa promover, questiona-se quanto aos riscos que ela poderá provocar à humanidade. Embora haja notáveis casos de sucesso de sua aplicação, em outros, foram alvo de severas críticas, em razão de discriminações provocadas, enviesamento em decisões e falta de transparência no processo decisório. Somado a isso, problematiza-se a utilização de tal tecnologia para fins de manipulação de discurso, armas autônomas letais, reconhecimento facial em espaços públicos e a vigilância estatal permanente por meio de sistemas de crédito social.

Em meio a esse cenário, são debatidas hipóteses quanto à necessidade de regulação da referida tecnologia. Ainda, quando e qual a melhor alternativa a ser adotada.

Na vanguarda, a Comissão Europeia¹, depois de dois anos de análise, publicou, em 2021, uma proposta de regulação de IA de forma horizontal, a qual deverá ser discutida até o final de 2023 junto ao Conselho da Europa e seu Parlamento. Ao longo da elaboração da proposta, a Comissão Europeia realizou avaliação de impacto, a partir da qual diversas opções regulatórias com diferentes graus de intervenção foram analisadas, considerando os impactos econômicos e sociais e, principalmente, os direitos fundamentais. A partir de tal análise, foi escolhida a opção regulatória que

adotava uma abordagem baseada no risco de forma proporcional, a qual poderia ser complementada por códigos de conduta para os sistemas de IA que não fossem classificados como de alto risco.²

Assim como ocorreu com o *General Data Protection Regulation* (GDPR) em 2018, acredita-se que a proposta de regulamento europeu para a IA possa se tornar um padrão global.³

Na mesma esteira, no Brasil, foram submetidos Projetos de Lei com o objetivo de regular a temática, sendo o de maior destaque o PL 21/2020 e seu substitutivo, o PL 2338/2023. Entre os diversos debates acerca da regulação de IA, discute-se qual seria o modelo e a abordagem regulatória mais adequada para tanto.

2. Riscos associados à sistemas dotados de IA

Os sistemas dotados de IA são capazes de realizar inferências e predições, todavia, tais decisões não são imunes a erros.⁴ Em casos como o de acidentes provocados por veículos autônomos ou até mesmo diagnósticos médicos por sistemas de IA, por exemplo, é muito provável que os indivíduos envolvidos desejarem obter maiores informações acerca do processo decisório realizado.⁵ Como mencionado anteriormente, sistemas de aprendizado profundo podem se revelar opacos, de modo a dificultar explicações acerca de suas decisões, sendo usualmente denominados de “caixas-pretas” (*black boxes*). Tal característica pode ser considerada problemática e consistir em um obstáculo para eventual responsabilização nos casos em que tais sistemas causem danos⁶ e na consecução do direito à revisão de decisões tomadas unicamente com base em tratamento automatizado de dados pessoais, previsto no art. 20 da Lei 13.709, de 14 de agosto de 2018 (LGL\2018\7222) – Lei Geral de Proteção de Dados (LGPD), assim como no art. 5º da Lei 12.414, de 09 de junho de 2011 (LGL\2011\1883) – Lei do Cadastro Positivo.

Como bem destacado por Isabela Ferrari, a problemática em torno de algoritmos que empregam técnicas de *machine learning* não é facilmente solucionável considerando que dada a complexidade de sua estrutura, mesmo que houvesse a exibição de código-fonte, para fins de explicação da tomada de decisões automatizadas, por exemplo, é muito provável que não auxiliaria a compreensão da forma que operam, uma vez que o “o referido código só expõe o método de aprendizado de máquinas usado, e não a regra de decisão, que emerge automaticamente a partir dos dados específicos sob análise.”⁷

Todavia, pesquisas têm sido realizadas no âmbito da área de Inteligência Explicável (*Explainable Artificial Intelligence – XAI*), a fim de que sejam desenvolvidos sistemas de IA mais transparentes, e consequentemente, mais explicáveis.⁸

Além da opacidade, outro risco associado aos sistemas de IA são vieses discriminatórios. Frequentemente são noticiados casos em que algoritmos incorrem em tais decisões⁹⁻¹⁰. Tais vieses podem se originar em diferentes estágios ao longo do desenvolvimento dos sistemas: no conjunto de dados, na construção do próprio sistema ou na interpretação da decisão.¹¹ Muitas vezes, os vieses estão associados a uma baixa representatividade nos dados utilizados para treinamento do sistema de IA¹² ou a comportamentos e práticas discriminatórias presentes até então na sociedade, refletindo, assim, vieses estruturais.¹³

Por sua vez, ao considerar o surgimento de uma IA Geral, Irving Good, em 1965, sinalizou para a ameaça de riscos existenciais à humanidade, podendo ser essa a “última invenção que o ser humano precisará fazer, contanto que a máquina seja dócil o suficiente para nos dizer como mantê-la sobre controle”.¹⁴

Ainda que cientistas da computação, como Stuart Russell, acreditem que uma IA Geral possa existir até o ano de 2045¹⁵, é necessário que se reflita igualmente sobre eventuais riscos associados a esse tipo de IA, dada a velocidade com que tais sistemas têm evoluído.¹⁶

No que tange à probabilidade de automação de empregos, a título de exemplo, no Brasil, segundo pesquisa realizada¹⁷, estima-se que cerca de 58% dos empregos brasileiros podem desaparecer nos próximos 10 ou 20 anos, em função da automação. Na mesma linha, pesquisa conduzida pelo Instituto de Pesquisa e Economia Aplicada (IPEA) prevê que caso empresas optem por automatizar as profissões com altas chances de automação, aproximadamente 30 milhões de empregos estariam em risco até 2026.¹⁸

Outro risco frequentemente associado à IA está relacionado com a mitigação da autodeterminação humana, a partir da personalização de conteúdo, por meio de técnicas de manipulação indireta de comportamento¹⁹ e a criação de “filtros de bolha”²⁰.

De fato, diferentes riscos podem ser associados à utilização de IA. No âmbito individual, sinaliza-se para riscos à direitos e garantias fundamentais, com destaque para vida, liberdade, igualdade, segurança, e à proteção dos dados pessoais. Quanto ao último ponto, destaca-se, por exemplo, a utilização de reconhecimento facial em espaços públicos²¹ e sua utilização para fins de vigilância social²² e discriminação.²³

No âmbito coletivo, aponta-se para ameaças à segurança nacional, à estabilidade política, à manipulação de áudios e vídeos (*deepfake*) e a ataques cibernéticos²⁴.

A respeito do tema, o Parlamento Europeu elaborou relatório por meio do qual foram elencados potenciais riscos em diferentes setores, como saúde, meio ambiente, política externa e mercado de trabalho.²⁵ O relatório destaca, por exemplo, que a utilização de sistemas de IA no diagnóstico de doenças pode representar riscos para os pacientes, embora a IA já supere os diagnósticos dos médicos em vários instâncias, como câncer de mama, o relatório considera que os atuais regimes de responsabilidade não proporcionam segurança jurídica suficiente sobre quem seria o responsável em caso de diagnóstico incorreto.²⁶ Aponta-se, também, para o desenvolvimento de armas letais autônomas, observando que esses tais armas já são utilizadas em conflitos militares e advertindo que mesmo grupos armados não estatais poderão em breve equipar drones com *software* de IA para navegação e reconhecimento facial reconhecimento e assim transformá-los em armas ofensivas letais de baixo custo capazes de agir inteiramente sem supervisão humana.²⁷

Relativamente aos riscos éticos relacionados aos sistemas de IA, em novembro de 2021, a Unesco publicou a “Recomendação sobre a ética da Inteligência Artificial”. Entre os principais objetivos da recomendação, estão o fornecimento de um marco universal de valores, princípios e ações para orientar os Estados na formulação de suas legislações, políticas ou outros instrumentos relativos à IA, em conformidade com o direito internacional, assim como a orientação de ações de indivíduos, grupos, comunidades, instituições e empresas do setor privado para garantia da incorporação da ética em todas as etapas do ciclo de vida dos sistemas de IA.²⁸

Em relação à avaliação de impacto ético, a Recomendação sugere que governos adotem marcos normativos que estabeleçam avaliações de impacto ético em sistemas de IA para previsão de consequências, mitigação de riscos e consequências prejudiciais e estabelecimento de mecanismos de supervisão que incluam auditabilidade, rastreabilidade e explicabilidade de algoritmos, incluindo a revisão externa dos sistemas de IA.²⁹

No que tange à mitigação desses riscos, uma das possíveis soluções a ser endereçada é a regulação *ex ante*, haja vista a potencialidade de riscos existenciais à humanidade, como já endereçado por diversos especialistas.³⁰

Todavia, conforme mencionado, há diferentes vertentes e graus de autonomia que a IA possa apresentar. Consequentemente, sistemas de IA são capazes de ensejar riscos sob diferentes gradações. Nesse sentido, para que as regulações sejam realizadas de forma proporcional, verifica-se a relevância de que elas incidam de forma distinta a depender do risco que tais sistemas possam acarretar.

2.1. Gradação de riscos

A ideia de gradação de riscos está diretamente associada a uma abordagem baseada em riscos. Ou seja, a partir da classificação deles, é possível que a regulação incida de forma distinta e proporcional. Na proposta da Comissão Europeia há inclusive a aplicação de regimes diferenciados de responsabilidade civil, por exemplo, conforme o risco que o sistema de IA possa provocar.

Conforme mencionado no item anterior, diferentes riscos podem ser associados a depender da aplicação de IA. Uma vez que há certa dificuldade de se prever e avaliar todas as suas possíveis utilizações, a OCDE, por exemplo, recomenda que os sistemas sejam agrupados em níveis de risco³¹, em sintonia com o que foi adotado pela Comissão Europeia na proposta o EU AI Act³², que classificou os riscos de sistemas de IA em: baixo risco, risco limitado alto risco e risco inaceitável.³³

Em 2019, a Comissão Alemã de Ética de Dados já propunha cinco níveis para classificação de riscos causados por sistemas de IA.³⁴ Além disso, outras instituições como o Federation of European Risk Management Associations (FERMA)³⁵, a International Organization for Standardization (ISO), o National Institute of Standards and Technology (NIST)³⁶ e o Institute of Electrical and Electronics Engineers (IEEE)³⁷ têm desenvolvido documentos para avaliação e gerenciamento de risco de sistemas de IA.

De modo semelhante, a OCDE, em 2022, publicou um *framework* para a classificação de sistemas de IA e definiu que entre os principais critérios para determinação do nível de risco que um sistema de IA possa ocasionar, estão: a escala, ou seja, a gravidade dos impactos adversos; e o escopo, ou seja, a amplitude de aplicação, como o número de indivíduos que são ou serão afetados e o grau de escolha quanto a estar sujeito aos efeitos de um sistema de IA.³⁸

Para a avaliação do risco do sistema de IA, a Avaliação de Impacto Algorítmico (AIA) pode ser uma alternativa eficaz. Na elaboração da AIA, segundo documento produzido pela ECP, sugerem-se oito etapas para condução da avaliação. Na primeira etapa, seria determinada a necessidade da realização de uma AIA, a partir de questões como: o sistema de IA possui um alto nível de autonomia? Há a utilização de dados pessoais sensíveis? O sistema de IA realiza tomada de decisões que possam acarretar grave impacto a indivíduos?³⁹

Por sua vez, na segunda etapa, o documento sugere que seja descrito o sistema de IA em si. Para tanto, são sugeridas questões como: qual o objetivo do sistema? Qual a tecnologia utilizada para o alcance do objetivo? Quais dados são utilizados no sistema? Quem está envolvido no desenvolvimento do sistema?⁴⁰

Na terceira etapa, o documento sugere que sejam descritos os benefícios que o sistema de IA possa oferecer para a instituição em que será aplicado, para cada indivíduo e para sociedade em geral.⁴¹

Como quarta etapa, é sugerido que se reflita se o objetivo do sistema é justificável ética e legalmente. Para tanto, questiona-se quais indivíduos estão envolvidos e/ou são afetados pelo sistema de IA; os valores e interesses estão previstos em leis e regulamentos, isto é, são legítimos?⁴²

Na quinta etapa, avalia-se se o sistema é seguro e transparente. Nesse sentido, questiona-se quais medidas foram adotadas para garantir a confiabilidade, segurança e transparência do sistema?⁴³

Nas etapas posteriores são previstas considerações e avaliações, reunião de documentos e, ao final, sugere-se que a avaliação seja revisada periodicamente.⁴⁴

Ainda que diferentes instituições tenham produzido via autorregulação orientações quanto à Análise de Impacto Algorítmico, entende-se que esse tema, juntamente com a gradação de riscos de sistemas de IA possam ser abordadas por meio de regulação estatal ou heterorregulação, ou até mesmo por meio da autorregulação regulada.

3. Abordagem regulatória baseada em riscos

A abordagem regulatória baseada em riscos (*risk based approach*) pode ser apontada como a mais indicada para regulação de IA.⁴⁵ Destaca-se que tal abordagem foi adotada na proposta de regulação de IA pela Comissão Europeia. Em que pese o desenvolvimento de novas tecnologias possa promover notáveis benefícios à sociedade em geral, ele é frequentemente associado a riscos os quais não são passíveis muitas vezes de serem objetivamente estimados.⁴⁶ Ulrich Beck já sinalizava em 1992, que diferentemente das décadas anteriores, os riscos futuros seriam globais, que poderiam se estender às próximas gerações e afetar a todos indistintamente, independentemente de classe, cultura ou cidadania.⁴⁷ Na mesma linha, Anthony Giddens defendia que a sociedade do risco seria aquela em que se viveria em uma fronteira altamente tecnológica e que absolutamente ninguém a compreenderia completamente, de modo a gerar uma diversidade de futuros possíveis.⁴⁸

Nesse sentido, a abordagem regulatória baseada em risco tornou-se popular nos últimos anos⁴⁹, sendo mais frequentemente discutida nos Estados Unidos, na Austrália, na Nova Zelândia e na Europa Continental⁵⁰ e comumente aplicada nos setores financeiro, de aviação, de segurança alimentar e o de nanotecnologia.⁵¹

A despeito das diferentes definições de risco na literatura,⁵² entre as mais recorrentes estão aquelas que definem risco como uma combinação entre a probabilidade de um evento e suas consequências⁵³, a probabilidade de um dano físico devido a determinado processo tecnológico ou outro processo⁵⁴. Segundo a OCDE, risco pode ser considerado como a combinação da probabilidade de dano e da potencial gravidade desse dano. Sob esse aspecto, a referida Organização recomenda que a regulação deve considerar os riscos, ser proporcional a eles e somente vir a existir quando esses forem de fato, significativos.⁵⁵

Na mesma linha, a abordagem regulatória baseada em riscos (*risk-based approach*) preceitua que as intervenções regulatórias devem se concentrar no controle das maiores ameaças potenciais para atingir os objetivos regulatórios, por meio de avaliações *ex ante* de sua probabilidade e consequências⁵⁶, a fim de tornar a regulação mais eficaz e proporcional. Como destacado por Julia Black, o foco da referida abordagem está nos riscos, e não em regras. Por essa razão, os reguladores precisam inicialmente identificar quais são os riscos que precisam ser gerenciados e não quais regras eles têm que aplicar⁵⁷.

Assim, a referida abordagem surgiu como alternativa regulatória mais flexível⁵⁸ em comparação a de comando e controle⁵⁹, para tornar os sistemas regulatórios mais eficientes, mais eficazes, mais resilientes e responsivos em tempos de crise, e também mais capazes de se comunicar sobre seus objetivos, capacidade e resultados.⁶⁰ No Reino Unido, desde 2005⁶¹, a regulação baseada em risco tem sido apontada como um princípio central da agenda para “melhor regulação”, e assim igualmente considerada pela OCDE e por muitos de seus estados membros.⁶² Na mesma linha, o professor Cary Coglianese sustenta que “em um mundo onde riscos são onipresentes, complexos e potencialmente extremamente onerosos, adotar uma abordagem baseada em riscos na regulação é essencial”.⁶³

Destaca-se que no ordenamento jurídico brasileiro, a referida abordagem pode ser observada, para além do Direito Ambiental, em leis como: LGPD, Lei 12.965, de 23 de abril de 2014 (LGL\2014\3339) (Marco Civil da Internet) e na Lei 8.078, de 11 de setembro de 1990 (LGL\1990\40) (Código de Defesa do Consumidor).⁶⁴

Insta salientar que uma ferramenta relevante para a gestão dos riscos é o Relatório de Análise de Impacto Regulatório, que é o documento resultante da Análise de Impacto Regulatório (AIR). Na análise dos riscos associados, analisa-se a relação causal entre os perigos e os efeitos que eles podem apresentar para os atores envolvidos e a sociedade em conjunto; a probabilidade de ocorrência (pesquisando dados históricos do fenômeno ou experiências internacionais sobre a matéria) e a caracterização do risco, descrevendo a sua severidade e o dano que ele pode trazer ao se materializar.⁶⁵

No Brasil, por força do art. 6º da Lei 13.848, de 25 de junho de 2019 (LGL\2019\5137) (Nova Lei das Agências Reguladoras) e do art. 5º Lei 13.874, de 20 de setembro de 2019 (LGL\2019\8262) (Lei da Liberdade Econômica), há previsão para que as propostas de edição e de alteração de atos normativos de interesse geral por órgão ou entidade da administração pública federal, incluídas as autarquias e as fundações públicas, devem ser precedidas da realização de AIR.

Na mesma linha, a LGPD previu a edição do Relatório de Impacto à Proteção de Dados pessoais (RIPD), o qual consiste em “documentação do controlador que contém a descrição dos processos de tratamento de dados pessoais que podem gerar riscos às liberdades civis e aos direitos fundamentais, bem como medidas, salvaguardas e mecanismos de mitigação de risco”.⁶⁶

Reflete-se que tal ferramenta é de suma importância no desenvolvimento de sistemas de Inteligência Artificial e deva ser igualmente endereçada a previsão de um “Relatório de Impacto Algorítmico”, à semelhança daquele exigido pelo Governo do Canadá.⁶⁷

Todavia a abordagem baseada em riscos para regulação de IA não é imune a críticas. Há quem defenda⁶⁸ que se deve adotar uma abordagem baseada em direitos (*rights-based approach*), por meio de uma avaliação de impacto nos Direitos Humanos, de modo que o ônus da prova recaia sobre a organização que deseja desenvolver ou implantar o sistema de IA.

A respeito do tema, em novembro de 2021, cerca de 114 (cento e quatorze) organizações da sociedade civil editaram uma espécie de manifesto intitulado “An EU Artificial Intelligence Act for

Fundamental Rights”⁶⁹, por meio do qual foram endereçadas recomendações ao Parlamento e ao Conselho Europeu, entre as quais quanto à abordagem baseada em riscos. Segundo os autores, a definição *ex ante* de categorias de risco não levaria em consideração que o nível de risco também depende do contexto em que o sistema está inserido, de modo que não pode ser totalmente determinado com antecedência. Outro ponto assinalado pelos signatários é que, diferentemente da previsão de atualização para lista de sistemas de IA classificados como de alto risco, não houve igual previsão para aqueles elencados como de riscos inaceitáveis (art. 5º do EU AI Act), tampouco aos de baixo risco (art. 52 do EU AI Act). Nesse sentido, sustentam que a rigidez de tais aspectos vai de encontro à relevância da IA, no que tange a sua capacidade de responder a desenvolvimentos futuros e riscos emergentes para os Direitos Fundamentais.⁷⁰ No sentido da proteção dos direitos humanos, destaca-se a posição de Ana Frazão, apresentada ao longo do seminário internacional organizado pela Comissão de Juristas responsável por subsidiar a elaboração de substitutivo sobre Inteligência Artificial no Brasil:

“[...] há uma certa preocupação de que essa abordagem possa ser muito formalista, muito procedimental que, em alguns aspectos, ela desconsidere ou não trate suficientemente alguns aspectos substantivos e que dizem respeito, realmente, à proteção dos direitos humanos, dos direitos fundamentais envolvidos. Há, também, um certo receio de que isso acabe deixando as empresas ainda com um grau de discricionariedade muito grande.”

Em contrapartida, defensores da abordagem baseada em riscos argumentam que ela viabiliza uma regulação mais justa e proporcional e que possibilita o uso eficiente e efetivo dos limitados recursos regulatórios.⁷¹

A despeito da sinalização de uma aparente dicotomia entre ambas as abordagens regulatórias, Ivar Hartmann, em Audiência Pública realizada no Senado Federal, ressaltou que, diferentemente do debate internacional, em que comumente são colocadas de forma antagônica, no Brasil não devem ser assim tratadas considerando que já há previsão de direitos constitucionais nesse contexto, de modo que qualquer escolha que renunciasse a direitos sequer poderia ser considerada.⁷² Sugeriu, inclusive, que deva ser dosada a presença de característica e elementos de ambos os modelos.⁷³

Na mesma linha, pesquisa realizada pela Ernst & Young Global Limited destaca que ambas as abordagens não são necessariamente excludentes entre si, de modo que uma abordagem baseada em riscos pode tratar de riscos à violação de direitos, assim como previsão constante na proposta do EU AI Act, em que o risco de dano aos direitos fundamentais encontra-se entre os critérios para a classificação de sistemas de alto risco.⁷⁴

Segundo Paola Cantarini, a abordagem baseada em riscos ou “via risquificação” possui tanto um aspecto individual quanto um coletivo e social dos direitos fundamentais, possuindo, assim, uma múltipla dimensionalidade:

“Na abordagem via risquificação ocorre a reformatação jurídica a partir da ampliação da tutela coletiva, a disseminação de instrumentos regulatórios *ex ante* e o uso intensivo de metodologias de gestão de risco e calibragem entre riscos, inovações, aproximando-se de um processo de “negociação coletiva”, trazendo semelhanças com o direito ambiental, como se percebe pela ideia de poluição de dados, ou seja, quando há um vazamento de dados não se está afrontando apenas direitos individuais do titular dos dados, mas muitas vezes valores democráticos, e o próprio Estado Democrático de Direito. Tal abordagem possui aproximação com a característica essencial de todo e qualquer direito fundamental, qual seja, sua múltipla dimensionalidade, ou seja, há um aspecto individual mas também coletivo e social dos direitos fundamentais.”⁷⁵

Ademais, Hartmann aduziu que considerando que o risco é inerente a tal tecnologia, poderia ser problemático cogitar um sistema focado exclusivamente em direitos, e que em se tratando de uma tecnologia que possui impactos em grupos da população, não seria aplicável o princípio da precaução, sob pena da criação de restrições imprevisíveis e desproporcionais de direitos, principalmente às empresas que desenvolvem e colocam no mercado tais aplicações.⁷⁶

4. Conclusão

Conforme demonstrado, ainda que muito se discuta em âmbito internacional e até mesmo nacional quanto a uma possível dicotomia entre abordagens regulatórias baseadas em direitos e baseadas

em risco e que essas seriam naturalmente excludentes entre si, devendo optar-se por uma em detrimento de outra, verifica-se que é possível que elas sejam conjugadas em um mesmo normativo. O fato de um regulamento estabelecer obrigações de forma escalonada sob a lógica de riscos, sendo mais brandas ou mais gravosas na medida em que os riscos são menores ou maiores não obsta que simultaneamente sejam considerados impactos aos direitos fundamentais daqueles que farão uso de tais sistemas de IA.

Tal premissa pôde ser observada tanto na proposta europeia quanto nas brasileiras, restando corroborada tal tese uma vez que ambas as propostas são orientadas para o desenvolvimento de benefícios aos seres humanos (*human-centered-AI*), acima de tudo.

5. Referências

ADADI, Amina; BERRADA, Mohammed. Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI). *IEEE Xplore*, v. 6, p. 52138-52160, 2018. Disponível em: [https://ieeexplore.ieee.org/abstract/document/8466590]. Acesso em: 07.08.2022.

ALBUQUERQUE, Pedro Henrique Melo *et al.* Na era das máquinas, o emprego é de quem? Estimativa da probabilidade de automação de ocupações no Brasil. Brasília/Rio de Janeiro: Ipea, 2019. Disponível em: [http://repositorio.ipea.gov.br/bitstream/11058/9116/1/td_2457.pdf]. Acesso em: 09.08.2022

BEAUSSIER, Anne-Laure et al. Accounting for failure: risk-based regulation and the problems of ensuring healthcare quality in the NHS. *Health, Risk & Society*, v. 18, n. 3, 2016. Disponível em: [www.ncbi.nlm.nih.gov/pmc/articles/PMC4950452]. Acesso em: 23.07.2022.

BECK, Ulrich. *Risk Society: Towards a New Modernity*. Sage Publications, 1992.

BLACK, Julia; BALDWIN, Robert. Really Responsive Risk-Based Regulation. *Law & Policy*, v. 32, n. 2, 2010. Disponível em: [https://onlinelibrary.wiley.com/doi/epdf/10.1111/j.1467-9930.2010.00318.x]. Acesso em: 23.07.2022.

BLEICHER, Ariel. Demystifying the Black Box That is AI. *Scientific American*. 09 ago. 2017. Disponível em: [www.scientificamerican.com/article/demystifying-the-black-box-that-is-ai]. Acesso em: 07.08.2022.

BONAVENTURE, Lionel. Engenheiro do Google é colocado em licença após afirmar que robô de bate-papo se tornou consciente. *Gaúcha Zero Hora*, Porto Alegre, 13 jun. 2022. Disponível em: [https://gauchazh.clicrbs.com.br/tecnologia/noticia/2022/06/engenheiro-do-google-e-colocado-em-licenca-apos-afirmar]. Acesso em: 08.08.2022.

BRASIL. *Diretrizes Gerais e Guia Orientativo para Elaboração de Análise de Impacto Regulatório – AIR*. Brasília: Presidência da República, 2018. Disponível em: [www.gov.br/economia/pt-br/assuntos/air/guias-e-documentos/diretrizesgeraisguiaorientativo_AIR_semlogo.pdf]. Acesso em: 08.08.2022.

BRASIL. Senado Federal. *Projeto de Lei 21, de 2020*. Estabelece fundamentos, princípios e diretrizes para o desenvolvimento e a aplicação da inteligência artificial no Brasil; e dá outras providências. Brasília: Senado Federal, 2020. Disponível em: [www25.senado.leg.br/web/atividade/materias/-/materia/151547].

BRASIL. Senado Federal. *Projeto de Lei 2338, de 2023*. Dispõe sobre o uso da Inteligência Artificial. Brasília: Senado Federal, 2023. Disponível em: [www25.senado.leg.br/web/atividade/materias/-/materia/15723].

CANTARINI, Paola. Inteligência artificial e abordagem via risquificação. *Migalhas*, 19 abr. 2022. Disponível em: [www.migalhas.com.br/arquivos/2022/4/40AB87D6286F11_IAemmovimento.pdf]. Acesso em: 15.07.2022.

CENTRE FOR INFORMATION POLICY LEADERSHIP – CIPL. *CIPL Recommendations on Adopting a Risk-Based Approach to Regulating Artificial Intelligence in the EU*. 2021. Disponível em: [www.informationpolicycentre.com/uploads/5/7/1/0/57104281/cipl_risk-based_approach_to_regulating_ai__22_march_

Acesso em: 22.07.2022.

COGLIANESE, Cary. What Does Risk-Based Regulation Mean? *The Regulatory Review*, 08 jul. 2019. Disponível em: [www.theregreview.org/2019/07/08/coglianese-what-does-risk-based-regulation-mean]. Acesso em: 08.08.2022.

COMISSÃO EUROPEIA. *Regulamento do Parlamento Europeu e do Conselho que estabelece regras harmonizadas em matéria de Inteligência Artificial (Regulamento Inteligência Artificial) e altera determinados atos legislativos da União – COM(2021) 206 final*. Bruxelas: Comissão Europeia, 2021. Disponível em: [https://eur-lex.europa.eu/legal-content/PT/TXT/HTML/?uri=CELEX:52021PC0206&from=EM]. Acesso em: 06.08.2022.

DASTIN, Jeffrey. Amazon scraps secret AI recruiting tool that showed bias against women. *Reuters*, Londres, 10 out. 2018. Disponível em: [www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G]. Acesso em: 07.08.2022.

DATA ETHICS COMMISSION. *Opinion of the Data Ethics Commission*. Berlin, 2019. Disponível em: [www.bmj.de/SharedDocs/Downloads/DE/Themen/Fokusthemen/Gutachten_DEK_EN.pdf?__blob=publicationFile&v=2]. Acesso em: 13.08.2022.

DÉFENSEUR DES DROITS; COMMISSION NATIONALE INFORMATIQUE & LIBERTÉS. *Algorithms: preventing automated discrimination*. 2020. Disponível em: [www.defenseurdesdroits.fr/sites/default/files/atoms/files/synth-algos-en-num-16.07.20.pdf]. Acesso em: 07.08.2022.

DIRK, Helbing, et al. Will democracy survive big data and artificial intelligence? In: HELBING, D. (Ed.). *Towards Digital Enlightenment: Essays on the Dark and Light Sides of the Digital Revolution*. Springer International Publishing, 2018. p. 73-98. Disponível em: [https://pure.rug.nl/ws/portalfiles/portal/111453647/Helbing2019_Chapter_WillDemocracySurviveBigDataAnd.pdf]. Acesso em: 09.08.2022.

EDWARDS, Lilian. *The EU AI Act proposal*. Ada Lovelace Institute. 2022. Disponível em: [www.adalovelaceinstitute.org/wp-content/uploads/2022/04/Expert-explainer-The-EU-AI-Act-11-April-2022.pdf]. Acesso em: 13.08.2022.

ELECTRONIC COMMERCE PLATFORM NETHERLANDS. *Artificial Intelligence Impact Assessment*. 2019. Disponível em: [https://ecp.nl/wp-content/uploads/2019/01/Artificial-Intelligence-Impact-Assessment-English.pdf]. Acesso em: 13.08.2022.

ERNST & YOUNG GLOBAL LIMITED. *A survey of artificial intelligence risk assessment methodologies: The global state of play and leading practices identified*, 2022. Disponível em: [www.trilateralresearch.com/wp-content/uploads/2022/01/A-survey-of-AI-Risk-Assessment-Methodologies-full-report.pdf]. Acesso em: 07.06.2022.

EUROPEAN DIGITAL RIGHTS et al. *An EU Artificial Intelligence Act for Fundamental Rights a Civil Society Statement*, 2021. Disponível em: [https://edri.org/wp-content/uploads/2021/12/Political-statement-on-AI-Act.pdf]. Acesso em: 08.06.2022.

EUROPEAN PARLIAMENT. *Report on artificial intelligence in a digital age (2020/2266(INI))*. Estrasburgo, 2022. Disponível em: [www.europarl.europa.eu/cmsdata/246872/A9-0088_2022_EN.pdf]. Acesso em: 07.08.2022.

EUROPEAN PARLIMENT. *Artificial Intelligence: threats and opportunities*. 04 maio 2022. Disponível em: [www.europarl.europa.eu/news/en/headlines/society/20200918STO87404/artificial-intelligence-threats-and-opportunities]. Acesso em: 07.08.2022.

FEDERATION OF EUROPEAN RISK MANAGEMENT ASSOCIATIONS (FERMA). *Artificial*
Página 8

Intelligence applied to risk management. Bruxelas, 2019. Disponível em: [www.ferma.eu/app/uploads/2019/11/FERMA-AI-applied-to-RM-FINAL.pdf]. Acesso em: 13.08.2022.

FERRARI, Isabela. *Accountability de Algoritmos: a falácia do acesso ao código e caminhos para uma explicabilidade efetiva*. 2019. Disponível em: [https://itsrio.org/wp-content/uploads/2019/03/Isabela-Ferrari.pdf]. Acesso em: 07.08.2022.

GIBBS, Samuel. Elon Musk: regulate AI to combat 'existential threat' before it's too late. *The Guardian*, Londres, 17 jul. 2017. Disponível em: [www.theguardian.com/technology/2017/jul/17/elon-musk-regulation-ai-combat-existential-threat-tesla-spacex-ceo]. Acesso em: 08.08.2022.

GIDDENS, Anthony. Risk and Responsibility. *The Modern Law Review*, v. 62 n. 1, 1999. Disponível em: [https://courses.washington.edu/sales09/Handouts/Giddens_Risk_Responsibility.pdf]. Acesso em: 22.07.2022.

GOOD, Irving John. Speculations Concerning the First Ultraintelligent Machine. *Advances in Computers*, Nova York, v. 6, 1965. p. 33. Disponível em: [www.sciencedirect.com/science/article/abs/pii/S0065245808604180]. Acesso em: 08.08.2022.

GOVERNMENT OF CANADA. *Algorithmic Impact Assessment tool*. 2022. Disponível em: [www.canada.ca/en/government/system/digital-government/digital-government-innovations/responsible-use-ai/algorithms.html]. Acesso em: 08.08.2022.

HAMPTON, Philip. *Reducing administrative burdens: effective inspection and enforcement*, 2005. Disponível em: [www.regulation.org.uk/library/2005_hampton_report.pdf]. Acesso em: 23.07.2022.

HEIKKILÄ, Melissa. A quick guide to the most important AI law you've never heard of the European Union is planning new legislation aimed at curbing the worst harms associated with artificial intelligence. *MIT Technology Review*, Cambridge, 13 maio 2022. Disponível em: [www.technologyreview.com/2022/05/13/1052223/guide-ai-act-europe]. Acesso em: 06.08.2022.

HIDVEGI, Fanny; LEUFER Daniel; MASSÉ Estelle. The EU should regulate AI on the basis of rights, not risks. *Access Now*, 17 fev. 2021. Disponível em: [www.accessnow.org/eu-regulation-ai-risk-based-approach]. Acesso em: 07.06.2022.

HIGÍDIO, José. ViaQuatro deve indenizar por implantar sistema de detecção facial nas estações. *Consultor Jurídico*, São Paulo, 10 maio 2021. Disponível em: [www.conjur.com.br/2021-mai-10/viaquatro-indenizar-implantar-sistema-deteccao-facial]. Acesso em: 07.08.2022.

INTERNATIONAL ORGANIZATION FOR STANDARDIZATION. *Risk management vocabulary*. 2008. Disponível em: [https://bambangkesit.files.wordpress.com/2015/12/iso-73_2009_risk-management-vocabulary.pdf]. Acesso em: 22.07.2022.

MONIRUZZAMAN, Mohammad. Risk of regulatory failure of "risk-based regulation" while using enterprise risk management as a meta-regulatory toolkit. *Asian Journal of Economics and Banking*, v. 6, n. 1, 2022. p. 103-121. Disponível em: [www.emerald.com/insight/content/doi/10.1108/AJEB-05-2021-0067/full/html]. Acesso em: 23.07.2022.

NATIONAL INSTITUTE OF STANDARDS AND TECHNOLOGY. *AI Risk Management Framework: Initial Draft*. Gaithersburg, 2022. Disponível em: [www.nist.gov/system/files/documents/2022/03/17/AI-RMF-1stdraft.pdf]. Acesso em: 13.08.2022.

ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *OECD Framework for the classification of AI systems*. Paris, 2022. Disponível em: [www.oecd-ilibrary.org/deliver/cb6d9eca-en.pdf?itemId=%2Fcontent%2Fpaper%2Fcb6d9eca-en&mimeType=pdf]. Acesso em: 13.08.2022.

ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *Risk-based regulation*. 2021. Disponível em:

[www.oecd-ilibrary.org/sites/9d082a11-en/index.html?itemId=/content/component/9d082a11-en]. Acesso em: 22.07.2022.

ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *Recommendation of the Council on Regulatory Policy and Governance*, 2012. Disponível em: [www.oecd.org/gov/regulatory-policy/49990817.pdf]. Acesso em: 23.07.2022.

ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *Risk-based regulation: Making sure that rules are science-based, targeted, effective and efficient*, OECD Regulatory Policy Outlook 2021, 2021. Disponível em: [www.oecd.org/gov/regulatory-policy/chapter-six-risk-based-regulation.pdf]. Acesso em: 23.07.2022.

ORGANIZAÇÃO DAS NAÇÕES UNIDAS PARA A EDUCAÇÃO, A CIÊNCIA E A CULTURA (UNESCO). *Recomendação sobre a Ética da Inteligência Artificial*. Paris, 2022. Disponível em: [https://unesdoc.unesco.org/ark:/48223/pf0000381137_por]. Acesso em: 07.08.2022.

OTTONI, Bruno et al. Automation

POMEROY, Robin. The promises and perils of AI – Stuart Russell on Radio Davos. *World Economic Forum*, Genebra, 06 jan. 2022. Disponível em: [www.weforum.org/agenda/2022/01/artificial-intelligence-stuart-russell-radio-davos]. Acesso em: 08.08.2022.

QIAN, Isabelle et al. Four takeaways from a Times investigation into China's expanding surveillance state. *The New York Times*, Nova York, 26 jul. 2022. Disponível em: [www.nytimes.com/2022/06/21/world/asia/china-surveillance-investigation.html]. Acesso em: 08.08.2022.

RAO, Anand. Five Views of AI Risk: Understanding the darker side of AI. *Towards Data Science*, 28 nov. 2020. Disponível em: [<https://towardsdatascience.com/five-views-of-ai-risk-eddb2fcea3c2>]. Acesso em: 07.08.2022.

REJMANIAK, Rafał. Bias in Artificial Intelligence Systems. *Białostockie Studia Prawnicze*, v. 26, n. 3, 2021. p. 25-42. Disponível em: [<https://heinonline.org/HOL/Page?handle=hein.journals/bialspw26&collection=journals&id=385&startid=&end=402>]. Acesso em: 07.08.2022.

ROTHSTEIN, Henry et al. The risks of risk-based regulation: Insights from the environmental policy domain. *Environment International*, v. 32, n. 8, 2006. Disponível em: [www.sciencedirect.com/science/article/pii/S0160412006000882]. Acesso em: 22.07.2022.

SAMPLE, Ian. What is facial recognition – and how sinister is it? *The Guardian*, Londres, 29 jul. 2019. Disponível em: [www.theguardian.com/technology/2019/jul/29/what-is-facial-recognition-and-how-sinister-is-it]. Acesso em: 08.08.2022.

SENADO FEDERAL. 2ª Reunião Ordinária – 2ª parte – Audiência Pública: Painel 2 – Inteligência artificial e regulação: modelos de regulação e abordagem. Disponível em: [<https://legis.senado.leg.br/comissoes/reuniao?3&reuniao=10701&codcol=250>]. Acesso em: 28.06.2022.

TELFORD, Taylor. Apple Card algorithm sparks gender bias allegations against Goldman Sachs. *Washington Post*, Washington, 11 nov. 2019. Disponível em: [www.washingtonpost.com/business/2019/11/11/apple-card-algorithm-sparks-gender-bias-allegations-against-goldman]. Acesso em: 07.08.2022.

ULRICH, Beck. *Risk Society: Towards a New Modernity*. Sage Publications, 1992.

UNITED NATIONS. *The right to privacy in the digital age*. Report of the United Nations High Commissioner for Human Rights, Nova York, 2021. Disponível em: [www.ohchr.org/en/2021/09/artificial-intelligence-risks-privacy-demand-urgent-action-bachelet]. Acesso em: 07.08.2022.

VAN DER HEIJDEN, Jeroen. Risk governance and risk-based regulation: A review of the international academic literature. *State of the Art in Regulatory Governance Research Paper*, 2019. Disponível em: [https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3406998]. Acesso em: 22.07.2022.

ZANATTA, Rafael. Proteção de dados pessoais como regulação de risco: uma nova moldura teórica? *Anais Rede*, 2017. Disponível em: [www.redegovernanca.net.br/public/conferences/1/anais/Anais_REDE_2017-1.pdf#page=179]. Acesso em: 08.08.2022.

6. Legislação

BRASIL. Lei 13.709, de 14 de agosto de 2018 (LGL\2018\7222). Disponível em: [www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/l13709.htm]. Acesso em: 08.08.2022.

1 COMISSÃO EUROPEIA. *Regulamento do Parlamento Europeu e do Conselho que estabelece regras harmonizadas em matéria de Inteligência Artificial (Regulamento Inteligência Artificial) e altera determinados atos legislativos da União – COM(2021) 206 final*. Bruxelas: Comissão Europeia, 2021. Disponível em: [https://eur-lex.europa.eu/legal-content/PT/TXT/HTML/?uri=CELEX:52021PC0206&from=EM]. Acesso em: 06.08.2022.

2 Idem.

3 HEIKKILÄ, Melissa. A quick guide to the most important AI law you've never heard of the European Union is planning new legislation aimed at curbing the worst harms associated with artificial intelligence. *MIT Technology Review*, Cambridge, 13 maio 2022. Disponível em: [www.technologyreview.com/2022/05/13/1052223/guide-ai-act-europe]. Acesso em: 06.08.2022.

4 UNITED NATIONS. *The right to privacy in the digital age*. Report of the United Nations High Commissioner for Human Rights, Nova York, 2021. Disponível em: [www.ohchr.org/en/2021/09/artificial-intelligence-risks-privacy-demand-urgent-action-bachelet]. Acesso em: 07.08.2022.

5 BLEICHER, Ariel. Demystifying the Black Box That is AI. *Scientific American*. 09 ago. 2017. Disponível em: [www.scientificamerican.com/article/demystifying-the-black-box-that-is-ai]. Acesso em: 07.08.2022.

6 Idem.

7 FERRARI, Isabela. *Accountability de Algoritmos: a falácia do acesso ao código e caminhos para uma explicabilidade efetiva*. 2019. Disponível em: [https://itsrio.org/wp-content/uploads/2019/03/Isabela-Ferrari.pdf]. Acesso em: 07.08.2022.

8 ADADI, Amina; BERRADA, Mohammed. Peeking Inside the Black-Box: A Survey on Explainable Artificial Intelligence (XAI). *IEEE Xplore*, v. 6, p. 52138-52160, 2018. Disponível em: [https://ieeexplore.ieee.org/abstract/document/8466590]. Acesso em: 07.08.2022.

9 TELFORD, Taylor. Apple Card algorithm sparks gender bias allegations against Goldman Sachs. *Washington Post*, Washington, 11 nov. 2019. Disponível em: [www.washingtonpost.com/business/2019/11/11/apple-card-algorithm-sparks-gender-bias-allegations-against-goldman-sachs/]. Acesso em: 07.08.2022.

10 DASTIN, Jeffrey. Amazon scraps secret AI recruiting tool that showed bias against women. *Reuters*, Londres, 10 out. 2018. Disponível em: [www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G]. Acesso em: 07.08.2022.

11 REJMANIAK, Rafał. Bias in Artificial Intelligence Systems. *Białostockie Studia Prawnicze*, v. 26,

n. 3, 2021. p. 25-42. Disponível em:

[<https://heinonline.org/HOL/Page?handle=hein.journals/bialspw26&collection=journals&id=385&startid=&end=402>]. Acesso em: 07.08.2022.

12 DÉFENSEUR DES DROITS; COMMISSION NATIONALE INFORMATIQUE & LIBERTÉS.

Algorithms: preventing automated discrimination. 2020. Disponível em:

[www.defenseurdesdroits.fr/sites/default/files/atoms/files/synth-algos-en-num-16.07.20.pdf]. Acesso em: 07.08.2022.

13 EUROPEAN PARLIAMENT. *Artificial Intelligence: threats and opportunities*. 04 maio 2022.

Disponível em:

[www.europarl.europa.eu/news/en/headlines/society/20200918STO87404/artificial-intelligence-threats-and-opportunities]. Acesso em: 07.08.2022.

14 "The first ultraintelligent machine is the last invention that man need ever make, provided the machine is docile enough to tell us how to keep it under control" (GOOD, Irving John. *Speculations Concerning the First Ultraintelligent Machine*. *Advances in Computers*, Nova York, v. 6, 1965. p. 33. Disponível em: [www.sciencedirect.com/science/article/abs/pii/S0065245808604180]. Acesso em: 08.08.2022 – tradução nossa).

15 POMEROY, Robin. The promises and perils of AI – Stuart Russell on Radio Davos. *World Economic Forum*, Genebra, 06 jan. 2022. Disponível em:

[www.weforum.org/agenda/2022/01/artificial-intelligence-stuart-russell-radio-davos]. Acesso em: 08.08.2022.

16 BONAVENTURE, Lionel. Engenheiro do Google é colocado em licença após afirmar que robô de bate-papo se tornou consciente. *Gaúcha Zero Hora*, Porto Alegre, 13 jun. 2022. Disponível em:

[<https://gauchazh.clicrbs.com.br/tecnologia/noticia/2022/06/engenheiro-do-google-e-colocado-em-licenca-apos-afirmar>]. Acesso em: 08.08.2022.

17 OTTONI, Bruno et al. Automation and job loss: the Brazilian case. *Nova Economia*, v. 32, n. 01, p. 157-180, 2022. Disponível em:

[www.scielo.br/j/neco/a/tHdRWS8KqNWZHHJq9LMPKyN/?lang=en#ModalArticles]. Acesso em: 09.08.2022.

18 ALBUQUERQUE, Pedro Henrique Melo et al. *Na era das máquinas, o emprego é de quem?*

Estimação da probabilidade de automação de ocupações no Brasil. Brasília/Rio de Janeiro: Ipea, 2019. Disponível em: [http://repositorio.ipea.gov.br/bitstream/11058/9116/1/td_2457.pdf]. Acesso em: 09.08.2022.

19 DIRK, Helbing, et al. Will democracy survive big data and artificial intelligence? In: HELBING, D. (Ed.). *Towards Digital Enlightenment: Essays on the Dark and Light Sides of the Digital Revolution*. Springer International Publishing, 2018. p. 73-98. Disponível em:

[https://pure.rug.nl/ws/portalfiles/portal/111453647/Helbing2019_Chapter_WillDemocracySurviveBigDataAnd.pdf]. Acesso em: 09.08.2022.

20 Estado de isolamento intelectual que supostamente pode resultar de buscas personalizadas quando um algoritmo adivinha seletivamente quais informações um usuário gostaria de ver com base nas informações sobre o usuário, como localização, comportamento de cliques anteriores e históricos de pesquisas. Como resultado, os usuários ficam separados das informações que discordam de seus pontos de vista, isolando-os efetivamente em suas próprias bolhas culturais ou ideológicas.

21 HIGÍDIO, José. ViaQuatro deve indenizar por implantar sistema de detecção facial nas estações. *Consultor Jurídico*, São Paulo, 10 maio 2021. Disponível em:

[www.conjur.com.br/2021-mai-10/viaquatro-indenizar-implantar-sistema-deteccao-facial]. Acesso em: 07.08.2022.

22 QIAN, Isabelle et al. Four takeaways from a Times investigation into China's expanding surveillance state. *The New York Times*, Nova York, 26 jul. 2022. Disponível em:

[www.nytimes.com/2022/06/21/world/asia/china-surveillance-investigation.html]. Acesso em: 08.08.2022.

23 SAMPLE, Ian. What is facial recognition – and how sinister is it? *The Guardian*, Londres, 29 jul. 2019. Disponível em: [www.theguardian.com/technology/2019/jul/29/what-is-facial-recognition-and-how-sinister-is-it]. Acesso em: 08.08.2022.

24 RAO, Anand. Five Views of AI Risk: Understanding the darker side of AI. *Towards Data Science*, 28 nov. 2020. Disponível em: [<https://towardsdatascience.com/five-views-of-ai-risk-eddb2fcea3c2>]. Acesso em: 07.08.2022.

25 EUROPEAN PARLIAMENT. *Report on artificial intelligence in a digital age (2020/2266(INI))*. Estrasburgo, 2022. Disponível em: [www.europarl.europa.eu/cmsdata/246872/A9-0088_2022_EN.pdf]. Acesso em: 07.08.2022.

26 Idem.

27 Idem.

28 ORGANIZAÇÃO DAS NAÇÕES UNIDAS PARA A EDUCAÇÃO, A CIÊNCIA E A CULTURA (UNESCO). *Recomendação sobre a Ética da Inteligência Artificial*. Paris, 2022. Disponível em: [https://unesdoc.unesco.org/ark:/48223/pf0000381137_por]. Acesso em: 07.08.2022.

29 ORGANIZAÇÃO DAS NAÇÕES UNIDAS PARA A EDUCAÇÃO, A CIÊNCIA E A CULTURA (UNESCO). *Recomendação sobre a Ética da Inteligência Artificial*. Paris, 2022. Disponível em: [https://unesdoc.unesco.org/ark:/48223/pf0000381137_por]. Acesso em: 07.08.2022.

30 GIBBS, Samuel. Elon Musk: regulate AI to combat 'existential threat' before it's too late. *The Guardian*, Londres, 17 jul. 2017. Disponível em: [www.theguardian.com/technology/2017/jul/17/elon-musk-regulation-ai-combat-existential-threat-tesla-spacex-ceo]. Acesso em: 08.08.2022.

31 ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *OECD Framework for the classification of AI systems*. Paris, 2022. Disponível em: [www.oecd-ilibrary.org/deliver/cb6d9eca-en.pdf?itemId=%2Fcontent%2Fpaper%2Fcb6d9eca-en&mimeType=pdf]. Acesso em: 13.08.2022.

32 COMISSÃO EUROPEIA. *Regulamento do Parlamento Europeu e do Conselho que estabelece regras harmonizadas em matéria de Inteligência Artificial (Regulamento Inteligência Artificial) e altera determinados atos legislativos da União*. Bruxelas, 2021. Disponível em: [https://eurlex.europa.eu/resource.html?uri=cellar:e0649735-a372-11eb-9585-01aa75ed71a1.0004.02/DOC_1&format=HTML]. Acesso em: 07.06.2022.

33 EDWARDS, Lilian. *The EU AI Act proposal*. Ada Lovelace Institute. 2022. Disponível em: [www.adalovelaceinstitute.org/wp-content/uploads/2022/04/Expert-explainer-The-EU-AI-Act-11-April-2022.pdf]. Acesso em: 13.08.2022.

34 DATA ETHICS COMMISSION. *Opinion of the Data Ethics Commission*. Berlin, 2019. Disponível em: [www.bmj.de/SharedDocs/Downloads/DE/Themen/Fokusthemen/Gutachten_DEK_EN.pdf?__blob=publicationFile&v=2]. Acesso em: 13.08.2022.

35 FEDERATION OF EUROPEAN RISK MANAGEMENT ASSOCIATIONS (FERMA). *Artificial Intelligence applied to risk management*. Bruxelas, 2019. Disponível em: [www.ferma.eu/app/uploads/2019/11/FERMA-AI-applied-to-RM-FINAL.pdf]. Acesso em: 13.08.2022.

36 NATIONAL INSTITUTE OF STANDARDS AND TECHNOLOGY. *AI Risk Management Framework: Initial Draft*. Gaithersburg, 2022. Disponível em: [www.nist.gov/system/files/documents/2022/03/17/AI-RMF-1stdraft.pdf]. Acesso em: 13.08.2022.

37 ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT, loc. cit.

38 ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *OECD Framework for the classification of AI systems*. Paris, 2022. Disponível em: [www.oecd-ilibrary.org/deliver/cb6d9eca-en.pdf?itemId=%2Fcontent%2Fpaper%2Fcb6d9eca-en&mimeType=pdf]. Acesso em: 13.08.2022.

39 ELECTRONIC COMMERCE PLATFORM NETHERLANDS. *Artificial Intelligence Impact Assessment*. 2019. Disponível em: [https://ecp.nl/wp-content/uploads/2019/01/Artificial-Intelligence-Impact-Assessment-English.pdf]. Acesso em: 13.08.2022.

40 Idem.

41 ELECTRONIC COMMERCE PLATFORM NETHERLANDS. *Artificial Intelligence Impact Assessment*. 2019. Disponível em: [https://ecp.nl/wp-content/uploads/2019/01/Artificial-Intelligence-Impact-Assessment-English.pdf]. Acesso em: 13.08.2022.

42 Idem.

43 Idem.

44 Idem.

45 CENTRE FOR INFORMATION POLICY LEADERSHIP – CIPL. *CIPL Recommendations on Adopting a Risk-Based Approach to Regulating Artificial Intelligence in the EU*. 2021. Disponível em: [www.informationpolicycentre.com/uploads/5/7/1/0/57104281/cipl_risk-based_approach_to_regulating_ai__22_march_2021.pdf]. Acesso em: 22.07.2022.

46 VAN DER HEIJDEN, Jeroen. Risk governance and risk-based regulation: A review of the international academic literature. *State of the Art in Regulatory Governance Research Paper*, 2019. Disponível em: [https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3406998]. Acesso em: 22.07.2022.

47 ULRICH, Beck. *Risk Society: Towards a New Modernity*. Sage Publications, 1992.

48 GIDDENS, Anthony. Risk and Responsibility. *The Modern Law Review*, v. 62 n. 1, 1999. Disponível em: [https://courses.washington.edu/sales09/Handouts/Giddens_Risk_Responsibility.pdf]. Acesso em: 22.07.2022.

49 ROTHSTEIN, Henry et al. The risks of risk-based regulation: Insights from the environmental policy domain. *Environment International*, v. 32, n. 8, 2006. Disponível em: [www.sciencedirect.com/science/article/pii/S0160412006000882]. Acesso em: 22.07.2022.

50 VAN DER HEIJDEN, loc. cit.

51 Idem.

52 Idem.

53 INTERNATIONAL ORGANIZATION FOR STANDARDIZATION. *Risk management vocabulary*. 2008. Disponível em: [https://bambangkesit.files.wordpress.com/2015/12/iso-73_2009_risk-management-vocabulary.pdf]. Acesso em: 22.07.2022.

54 BECK, Ulrich. *Risk Society: Towards a New Modernity*. Sage Publications, 1992. p. 4.

55 ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *Risk-based*

regulation: Making sure that rules are science-based, targeted, effective and efficient, OECD Regulatory Policy Outlook 2021, 2021. Disponível em: [www.oecd.org/gov/regulatory-policy/chapter-six-risk-based-regulation.pdf]. Acesso em: 23.07.2022.

56 BEAUSSIER, Anne-Laure et al. Accounting for failure: risk-based regulation and the problems of ensuring healthcare quality in the NHS. *Health, Risk & Society*, v. 18, n. 3, 2016. Disponível em: [www.ncbi.nlm.nih.gov/pmc/articles/PMC4950452]. Acesso em: 23.07.2022.

57 BLACK, Julia; BALDWIN, Robert. Really Responsive Risk-Based Regulation. *Law & Policy*, v. 32, n. 2, 2010. Disponível em: [https://onlinelibrary.wiley.com/doi/epdf/10.1111/j.1467-9930.2010.00318.x]. Acesso em: 23.07.2022.

58 MONIRUZZAMAN, Mohammad. Risk of regulatory failure of “risk-based regulation” while using enterprise risk management as a meta-regulatory toolkit. *Asian Journal of Economics and Banking*, v. 6, n. 1, 2022. p. 103-121. Disponível em: [www.emerald.com/insight/content/doi/10.1108/AJEB-05-2021-0067/full/html]. Acesso em: 23.07.2022.

59 A regulação por normas de comando e controle caracteriza-se nos casos em que a estrutura normativa opera a partir do binômio prescrição-sanção, de modo que o particular é compelido a adotar determinado comportamento, sob pena de aplicação de uma pena.

60 ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *Risk-based regulation*. 2021. Disponível em: [www.oecd-ilibrary.org/sites/9d082a11-en/index.html?itemId=/content/component/9d082a11-en]. Acesso em: 22.07.2022.

61 HAMPTON, Philip. *Reducing administrative burdens: effective inspection and enforcement*, 2005. Disponível em: [www.regulation.org.uk/library/2005_hampton_report.pdf]. Acesso em: 23.07.2022.

62 ORGANISATION FOR ECONOMIC CO-OPERATION AND DEVELOPMENT (OECD). *Recommendation of the Council on Regulatory Policy and Governance*, 2012. Disponível em: [www.oecd.org/gov/regulatory-policy/49990817.pdf]. Acesso em: 23.07.2022.

63 “In a world where risks are omnipresent, complex, and potentially extremely costly, taking a risk-based approach to regulation is essential” (COGLIANESE, Cary. What Does Risk-Based Regulation Mean? *The Regulatory Review*, 08 jul. 2019. Disponível em: [www.theregreview.org/2019/07/08/coglianesse-what-does-risk-based-regulation-mean]. Acesso em: 08.08.2022 – tradução nossa).

64 ZANATTA, Rafael. Proteção de dados pessoais como regulação de risco: uma nova moldura teórica? *Anais Rede*, 2017. Disponível em: [www.redegovernanca.net.br/public/conferences/1/anais/Anais_REDE_2017-1.pdf#page=179]. Acesso em: 08.08.2022.

65 BRASIL. *Diretrizes Gerais e Guia Orientativo para Elaboração de Análise de Impacto Regulatório – AIR*. Brasília: Presidência da República, 2018. Disponível em: [www.gov.br/economia/pt-br/assuntos/air/guias-e-documentos/diretrizesgeraisguiaorientativo_AIR_semlogo.pdf]. Acesso em: 08.08.2022.

66 BRASIL. Lei 13.709, de 14 de agosto de 2018 (LGL\2018\7222). Disponível em: [www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/l13709.htm]. Acesso em: 08.08.2022.

67 GOVERNMENT OF CANADA. *Algorithmic Impact Assessment tool*. 2022. Disponível em: [www.canada.ca/en/government/system/digital-government/digital-government-innovations/responsible-use-ai/algorithmic-impact-assessment-tool.html]. Acesso em: 08.08.2022.

68 HIDVEGI, Fanny; LEUFER Daniel; MASSÉ Estelle. The EU should regulate AI on the basis of rights, not risks. *Access Now*, 17 fev. 2021. Disponível em: [www.accessnow.org/eu-regulation-ai-risk-based-approach]. Acesso em: 07.06.2022.

69 EUROPEAN DIGITAL RIGHTS et al. *An EU Artificial Intelligence Act for Fundamental Rights a Civil Society Statement*, 2021. Disponível em: [https://edri.org/wp-content/uploads/2021/12/Political-statement-on-AI-Act.pdf]. Acesso em: 08.06.2022.

70 Idem.

71 ERNST & YOUNG GLOBAL LIMITED. *A survey of artificial intelligence risk assessment methodologies: The global state of play and leading practices identified*, 2022. Disponível em: [www.trilateralresearch.com/wp-content/uploads/2022/01/A-survey-of-AI-Risk-Assessment-Methodologies-full-report.pdf]. Acesso em: 07.06.2022.

72 SENADO FEDERAL. *2ª Reunião Ordinária – 2ª parte – Audiência Pública: Painel 2 – Inteligência artificial e regulação: modelos de regulação e abordagem*. Disponível em: [https://legis.senado.leg.br/comissoes/reuniao?3&reuniao=10701&codcol=250]. Acesso em: 28.06.2022.

73 SENADO FEDERAL. *2ª Reunião Ordinária – 2ª parte – Audiência Pública: Painel 2 – Inteligência artificial e regulação: modelos de regulação e abordagem*. Disponível em: [https://legis.senado.leg.br/comissoes/reuniao?3&reuniao=10701&codcol=250]. Acesso em: 28.06.2022.

74 ERNST & YOUNG GLOBAL LIMITED, op. cit., p. 11.

75 CANTARINI, Paola. Inteligência artificial e abordagem via risquificação. *Migalhas*, 19 abr. 2022. Disponível em: [www.migalhas.com.br/arquivos/2022/4/40AB87D6286F11_IAemmovimento.pdf]. Acesso em: 15.07.2022.

76 SENADO FEDERAL. *2ª Reunião Ordinária – 2ª parte – Audiência Pública: Painel 2 – Inteligência artificial e regulação: modelos de regulação e abordagem*. Disponível em: [https://legis.senado.leg.br/comissoes/reuniao?3&reuniao=10701&codcol=250]. Acesso em: 28.06.2022.