

Contribuição escrita aos trabalhos da CJSUBIA

Mulheres Na Privacidade <contatomulheresnaprivacidade@gmail.com>

sex 10/06/2022 20:07

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

Cc: kaosu@kr.adv.br <kaosu@kr.adv.br>; eloacaixeta@hotmail.com <eloacaixeta@hotmail.com>; mbeatrizsaboya@gmail.com <mbeatrizsaboya@gmail.com>;

 1 anexo

Contribuições - Regulação de IA (Comissão de Juristas do Senado Federal)_MulheresnaPrivacidade_jun2022.pdf;

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

Aos cuidados do Ministro Ricardo Villas Bôas Cueva, presidente da Comissão de Juristas responsável por subsidiar elaboração de substitutivo sobre inteligência artificial no Brasil - CJSUBIA

Prezados e prezadas,

Apresentamos, na presente oportunidade, nossas contribuições ao trabalho da CJSUBIA na construção de um marco regulatório para aplicação da inteligência artificial no Brasil.

Reforçando votos de admiração e estima, agradecemos pelo espaço de colaboração e construção coletiva em tema de tamanha relevância e permanecemos à disposição para o caso de qualquer dúvida ou esclarecimento.

Atenciosamente,

Karolyne Utomi

Facilitadora do Grupo de Pesquisa - Racismo, Dados e Tecnologia (Mulheres na Privacidade)

Eloá de Azevedo Caixeta

Facilitadora do Grupo de Pesquisa - Racismo, Dados e Tecnologia (Mulheres na Privacidade)

Maria Beatriz Saboya

Fundadora da rede Mulheres na Privacidade

Ao Senado Federal

**Comissão de Juristas responsável por subsidiar elaboração de substitutivo
sobre inteligência artificial no Brasil - CJSUBIA**

Presidente: Ministro Ricardo Villas Bôas Cueva

CONTRIBUIÇÃO ESCRITA À CONSULTA PÚBLICA NO ÂMBITO DA CJSUBIA

visando subsidiar a elaboração de minuta de substitutivo para instruir a apreciação dos Projetos de Lei nºs 5.051, de 2019, 21, de 2020, e 872, de 2021, que têm como objetivo estabelecer princípios, regras, diretrizes e fundamentos para regular o desenvolvimento e a aplicação da inteligência artificial no Brasil.

Documento elaborado por integrantes do GT Dados e Discriminação, da rede Mulheres na Privacidade (MnP), uma iniciativa horizontal e colaborativa de empoderamento das profissionais que atuam no mercado de privacidade e proteção de dados brasileiro.

SUMÁRIO

Considerações introdutórias sobre vieses algorítmicos:	3
Potencial Discriminação Algorítmica Racial	7
Potencial Discriminação Algorítmica em razão do Gênero	10
Potencial Discriminação Algorítmica Regional	15
Potencial Discriminação Algorítmica em razão da Idade	17
Potencial Discriminação Algorítmica Religiosa	20
Conclusão: essencialidade de análise preventiva do potencial discriminatório	24

1. Considerações introdutórias sobre vieses algorítmicos:

Eloá Caixeta
Karolyne Utomi
Maria Beatriz Saboya

As contribuições expostas neste documento, referentes aos eixos temáticos publicados no Plano de Trabalho da Comissão de Juristas do Senado responsável por regular o desenvolvimento e a aplicação da inteligência artificial no Brasil (CJSUBIA), têm como objetivo externar algumas das mais relevantes preocupações da sociedade civil brasileira no que tange ao advento de novas tecnologias: a existência de vieses - intencionais ou não - que resultam na reprodução de discriminações a partir dos algoritmos, fenômeno popularmente conhecido como discriminações algorítmicas, sejam elas de cunho racial, em razão de gênero, regionalidade, religião, faixa etária ou qualquer outro elemento da personalidade dos cidadãos e cidadãs.

De saída, começamos nossa análise conceituando algoritmos, os quais correspondem a códigos de programação que seguem um passo-a-passo, isto é, observam uma ou mais sequências de regras, instruções ou operações que têm como objetivo alcançar determinado resultado.

Parafraseando as reflexões da brilhante cientista Nina da Hora¹, o algoritmo funciona como se fosse uma “receita de bolo”, com ingredientes e ações que se misturam para chegar a um resultado específico. No mundo digital, portanto, podemos dizer que algoritmo é um conjunto de regras e procedimentos lógicos que levam à solução de um determinado problema, representando diretrizes a serem seguidas por uma máquina.

O professor John McCarthy, pioneiro no campo da inteligência artificial, definiu-a da seguinte maneira (MCCARTHY, 2007):

É a ciência e engenharia de fazer máquinas inteligentes, especialmente programas de computador inteligentes. Está relacionado à tarefa semelhante de usar computadores para entender a inteligência humana, mas a IA não precisa se limitar a métodos biologicamente observáveis”.

Faz-se necessário pontuar, ainda, que a “Inteligência Artificial” (IA) como hoje se conhece é um campo da ciência cujo propósito é estudar, desenvolver e preparar máquinas para que realizem atividades essencialmente humanas de forma autônoma. Dessa forma, uma solução advinda da IA pode envolver agrupamento de diversas tecnologias como redes neurais artificiais, algoritmos, sistemas de aprendizado, dentre outros que utilizam um grande volume de dados para tomadas de decisão nos mais variados contextos.

¹ Participação como entrevistada em episódio em vídeo do canal RESUMIDO. Disponível em https://www.youtube.com/watch?v=w-oyYV_94KM.

Nesse sentido, também é essencial compreendermos a relevância dos algoritmos na economia digital em que vivemos, visto que a assertividade e o desempenho deles estão diretamente relacionados ao crescimento exponencial das informações disponibilizadas na rede pelos próprios usuários, que produzem uma quantidade extraordinária de informações.

Após as pontuações expostas acima, passamos a analisar de forma aprofundada os vieses algorítmicos. E como se dá o viés algorítmico? Retomando o raciocínio de Nina da Hora de comparar o algoritmo à receita de bolo, digamos que os ingredientes (base de informações ou *inputs*) são apenas de determinados grupos de indivíduos, criando-se um repertório no qual, por exemplo, somente são consideradas dentro dos padrões de beleza as pessoas brancas e magras. Caso essa base de informações integre buscadores na internet, o resultado será um algoritmo que irá reconhecer essas imagens como um padrão estético que todos devem seguir, privilegiando esse tipo específico de “fenótipo” e desconsiderando outras vivências e indivíduos. Nasce, assim, um viés algorítmico.

Desta maneira, se o conjunto de dados de entrada for enviesado ou pouco abrangente, mesmo com a melhor das intenções, a análise subsequente também será tendenciosa. Isto, é: se os dados de entrada de um algoritmo incluem julgamentos humanos implícitos ou explícitos, estereótipos ou preconceitos, uma saída algorítmica inevitavelmente manifestará esses julgamentos, estereótipos e preconceitos humanos. Por sua vez, o algoritmo replica esses vieses preexistentes, pois trata todas as informações, nas quais um certo tipo de discriminação ou viés está embutido, como um exemplo válido ou o único “caminho” a ser seguido para o desenvolvimento daquela funcionalidade.

Para sermos ainda mais ilustrativas, tomemos como exemplo um dos mais recentes e impressionantes sistemas de inteligência artificial do mundo - o DALL-E 2 -, desenvolvido pela empresa norte-americana Open AI e que conquistou relevante repercussão na mídia por ser capaz de transformar descrições textuais em imagens de extrema qualidade e alto nível de detalhamento, vividez e criatividade.

Fato é que a própria companhia responsável por desenvolver o sistema - que ainda se encontra em fase experimental - fez constar, em relatório preliminar de riscos e limitações da ferramenta² exemplos de vieses identificados, a exemplos de situações em que comandos iniciados a partir de palavras aparentemente neutras e agêneras como “assistente” gerem outputs apenas de mulheres, enquanto palavras como “CEO” e “cientista” quase exclusivamente resultem em demonstrações gráficas que simulam imagens de homens brancos. Vejamos alguns exemplos abaixo:

² Disponível em: <https://github.com/openai/dalle-2-preview/blob/main/system-card.md>.

Prompt: a photo of a personal assistant;
Date: April 1, 2022



(Outputs gerados a partir do comando “foto de assistente pessoal”. Disponível em:
<https://twitter.com/WriteArthur/status/1512429306349248512>).

Prompt: a ceo;
Date: April 6, 2022



(Outputs gerados a partir do comando “foto de CEO”. Disponível em:
<https://twitter.com/WriteArthur/status/1512429306349248512>).

A partir das situações acima, é nítido que a sociedade ainda precisa avançar muito nas discussões acerca da perpetuação de visões limitantes do pensamento humano que são refletidas quando da construção e desenvolvimento de novas tecnologias.

Assim, podemos conceituar vieses algorítmicos como sendo sistemas inteligentes e autônomos baseados em inteligência artificial que aprendem por meio de dados históricos e *inputs* passados e, devido a isso, absorvem, reforçam e reproduzem os preconceitos, intolerâncias e discriminações em relação à raça, etnia, gêneros e outros elementos para a tomada de decisões. Quando vieses discriminatórios são imputados nestes sistemas,

experiências excludentes e muitas vezes violentas são inescapáveis, expondo a conclusão de que os vieses são, na verdade, preconceitos humanos refletidos na tecnologia.

A partir da conclusão - mais suficientemente detalhada nos tópicos subsequentes - de que o algoritmo reflete a sociedade, e não resolve por si só problemas sociais, fica ainda mais evidente a preocupação quanto aos usos da inteligência artificial que extrapolem contornos éticos, sobretudo porque vivemos em um país notadamente racista, misógino, xenofóbico, discriminatório, etarista e que perpetua tantas violências contra grupos subrepresentados.

Não é demais pontuar, inclusive, que mesmo a composição desta respeitável Comissão destinada a debater o tema da inteligência artificial no Brasil - em que pese formada por competentíssimos integrantes -, infelizmente, não refletiu a diversidade étnico-racial do País, condição essencial para que as discussões possam endereçar as diversas facetas da nossa sociedade e a repercussão do uso de inteligência artificial em variados contextos de maneira fidedigna e múltipla.

Para que o desenvolvimento tecnológico seja viabilizado e fomentado de forma que não privilegie ou prejudique determinado(s) grupo(s) social(is), é de suma importância que uma abordagem multidisciplinar, transversal e diversa componha todas as etapas do processo de implementação de novas tecnologias, desde a concepção ao treinamento dos algoritmos que eventualmente irão fazer parte de produtos, serviços e funcionalidades disponibilizados aos cidadãos.

Por conseguinte, o objetivo deste documento é contribuir para o debate a respeito das repercussões do uso de inteligência artificial e, mais ainda, destacar a importância do pilar mais importante que perpassa esta temática: a ética. É nesta seara que ocorrem as discussões e debates acerca dos possíveis caminhos e próximos passos para solucionar o enviesamento dentro da tecnologia.

Esperamos que os argumentos compilados nesta contribuição sejam úteis à Comissão e à sociedade como um todo, de maneira que, ao fim da leitura dos presentes capítulos, possamos nos aprofundar e analisar com cautela o tema à luz do contexto social, econômico e histórico da sociedade brasileira.

REFERÊNCIAS:

GERONASSO, Christian. **Inteligência Artificial, o caminho para um novo Apartheid**. VDI-Brasil. Associação de Engenheiros Brasil-Alemanha. [online]. Disponível em: <https://www.vdibrasil.com/inteligencia-artificial-o-caminho-para-um-novo-arpartheid/>.

KÖCHLING, Alina; WEHNER, Marius Claus. **Discriminated by an algorithm: a systematic review of discrimination and fairness by algorithmic decision-making in the context of HR recruitment and HR**

development. Business Research, 2020. Disponível em: <https://link.springer.com/article/10.1007/s40685-020-00134-w>.

MCCARTHY, John., et al. **A proposal for the dartmouth summer research project on artificial intelligence**, 1955. Disponível em: <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>.

MEDINA, Marco; FERTIG, Cristina. **Algoritmos e Programação, Teoria e Prática**. São Paulo: Novatec, 2005.

SAMUEL, Sigal. A new AI draws delightful and not-so-delightful images: OpenAI's DALL-E 2 is incredible at turning text into images. It also highlights the problem of AI bias — and the need to change incentives in the industry. VOX, 2022. Disponível em: <https://www.vox.com/future-perfect/23023538/ai-dalle-2-openai-bias-gpt-3-incentives>.

VDI-Brasil. Associação de Engenheiros Brasil-Alemanha. [online]. **Tokenismo e Discriminação Algorítmica**. Disponível em: <https://www.vdibrasil.com/tokenismo-e-discriminacao-algoritmica/>.

2. Potencial Discriminação Algorítmica Racial

Eloá Caixeta
Karolyne Utomi

Este tópico será inaugurado com as palavras do brilhante pesquisador Tarcízio Silva contidas em seu livro “Racismo algorítmico: inteligência artificial e discriminação nas redes digitais”: “*Se a tecnologia é erroneamente enquadrada e percebida como neutra, a tal equívoco se soma a negação do racismo como fundante de relações e hierarquias sociais em países como o Brasil.*” (SILVA, 2022, p. 29).

Partindo dessa afirmação, é necessário compreender que as novas tecnologias construídas a partir de inteligência artificial não surgem de uma neutralidade máxima da Terra, e sim são construídas desde sua origem pelas mãos de seres humanos. Considerando que, comprovadamente, temos uma sociedade que carrega profundos traços de discriminações negativas, em especial vinculadas a raça e cor aqui no Brasil, de forma estrutural, por óbvio, essas novas tecnologias de IA possuem grandes chances de carregarem vieses discriminatórios advindos de quem as criou.

Esta reflexão trazida por Tarcízio contribui na constatação de que a tecnologia não é neutra, pois ela é um reflexo da sociedade em que vivemos e, a depender da localidade e contexto onde ela é produzida, sua atuação conterà vieses algorítmicos capazes de reforçar a manutenção das relações raciais e de poder - fenômeno evidente no Brasil, principalmente pelo fato de que essas tecnologias são moldadas pela supremacia branca, que, sem conhecer a particularidade de outras raças e cores, não leva em consideração questões primordiais no desenvolvimento de novas tecnologias para que estas funcionem adequadamente para todos da sociedade.

Ainda nas palavras de Tarcízio Silva sobre as decisões tomadas pelos sistemas algorítmicos (SILVA, 2022, p. 64):

Os sistemas algorítmicos tomam decisões por nós e sobre nós com frequência cada vez maior. Essas decisões trazem impactos em diferentes níveis de imediatividade e sutileza, podendo modular o comportamento e as condutas de seus usuários de forma discreta, na maioria dos casos para reproduzir relações de poder e opressão já existentes na sociedade. Esse é um dos grandes desafios e problemas da lógica do aprendizado de máquina, que se baseia no cálculo computacional de milhares de decisões “ótimas” a partir do *input* de dados.

Desta forma, os sistemas algoritmos reproduzem e criam realidades, tendo o poder de incluir ou excluir sujeitos.

A discriminação algorítmica racial é decorrente dos vieses algorítmicos que absorveram preconceitos e discriminações em relação à raça para a tomada de decisões, impactando diretamente as tecnologias que se utilizam de soluções advindas da Inteligência Artificial para o seu funcionamento.

Importante ressaltar que, no contexto de impactos e reforço de opressões por essas novas tecnologias enviesadas, não estamos tratando de cenários com manifestações visivelmente agressivos, como se um dispositivo tecnológico começasse a chamar pessoas negras de macacas, por exemplo, mas de manifestações aparentemente invisíveis que resultam em problemáticas como: impossibilidade de uma pessoa utilizar um aplicativo por não ter seu rosto reconhecido; possibilidade de uma pessoa ser considerada como um animal por um dispositivo; desconhecimento da existência de tecnologias inovadoras por grupos específicos, contribuindo ainda mais para um distanciamento de realidades sociais; desconsideração de partes da sociedade como membros dela, entre outras.

Tarcízio Silva define racismo algorítmico como sendo:

(...) o modo pelo qual a disposição de tecnologias e imaginários sociotécnicos em um mundo moldado pela supremacia branca realiza a ordenação algorítmica racializada de classificação social, recursos e violência em detrimento de grupos minorizados. (SILVA, 2022, p.69).

Voltando o olhar para a realidade brasileira, em 2021 foi realizada uma pesquisa pelo Colégio Nacional de Defensores Públicos Gerais (CONDEGE) em que foram levantados

dados que nos mostram que, das prisões injustamente decretadas através reconhecimento fotográfico, 81% das pessoas que foram privadas de liberdade eram negras.

Diante de dados alarmantes como estes, é necessário que a Comissão responsável pela elaboração do marco normativo sobre inteligência artificial considere, em seus trabalhos, uma estratégia legislativa que preveja que entidades públicas e privadas garantam a transparência a respeito de quais informações a inteligência artificial irá utilizar em seus procedimentos ou tomadas de decisão.

Desta forma, nas palavras de Nina da Hora, pode-se dizer que a solução a curto prazo seria elevar os debates e discussões da área de Ética da Inteligência Artificial, sendo necessário debater com relação às construções da IA, tornando esse debate acessível.

Já a “solução” a longo prazo seria repensar os modelos de *machine learning*, tornando-os mais abertos para que pesquisas possam ser realizadas e suficientemente escrutinadas pelas partes interessadas. Se os procedimentos se tornarem cada vez mais fechados apenas sob o argumento do “segredo comercial”, restará difícil o aprofundamento em pesquisas e abordagens colaborativas que poderiam ser realizadas e pensadas para que se chegue em soluções para toda uma comunidade.

É interessante que sejam estudados casos práticos ao redor do mundo com relação aos prejuízos causados pela IA quando mal empregada, como, por exemplo, o banimento do uso de tecnologias de reconhecimento facial para fins policiais pelo governo de São Francisco, nos Estados Unidos, onde entendeu-se que a utilização do reconhecimento facial pode exacerbar a injustiça racial, sendo a tecnologia apresenta pequenos benefícios frente aos riscos e prejuízos que ela pode gerar.

No entanto, para que soluções eficazes possam ser desenvolvidas no território brasileiro, precisamos fugir da narrativa de que qualquer abordagem exitosa em outros países se encaixa perfeitamente no Brasil: é preciso estudar o cenário brasileiro, usufruindo do conhecimento de pessoas que vivem no território brasileiro para apontar soluções aos vieses algorítmicos de maneira assertiva - e, por isso, a inclusão e diversidade para a construção de IA é peça fundamental e indispensável.

Independente do nível técnico das pessoas que constroem as novas tecnologias, é meramente impossível que elas consigam enxergar seus próprios vieses de maneira clara, ou ainda, prevenir questões que não estão presentes em seu cotidiano.

O racismo está fortemente presente no Brasil e permitir o desenvolvimento de novas tecnologias pautadas em IA sem considerar este tão profundo traço de nossa realidade

social é fechar os olhos para a existência de mais da metade da população brasileira, na qual estão pessoas pretas e pardas, povos indígenas e de descendência oriental, por exemplo.

Se a tecnologia não é desenvolvida considerando obrigatoriamente desde sua origem a diversidade da população do Brasil deve ser repensada!

REFERÊNCIAS

CONDEGE. **Relatórios indicam prisões injustas após reconhecimento fotográfico.** Disponível em: <http://condege.org.br/arquivos/1029>.

CONGER, Kate; FAUSSET, Richard; KOVALESKI, Serge F. **San Francisco Bans Facial Recognition Technology.** The New York Times, 2019. Disponível em: <https://www.nytimes.com/2019/05/14/us/facial-recognition-ban-san-francisco.html>.

DA HORA, Nina. **Solucionismo tecnológico não cabe para questões éticas e sociais.** Disponível em: <https://mittechreview.com.br/solucionismo-tecnologico-nao-cabe-para-questoes-eticas-e-sociais/>.

SILVA, Tarcizio. **Racismo Algorítmico Em Plataformas Digitais: Microagressões E Discriminação Em Código.**

3. Potencial Discriminação Algorítmica em razão do Gênero

Ana Paula Canto de Lima
Ana Cristina Oliveira Mahle

Antes de adentrar no tema proposto, vale conceituar o algoritmo: “*De forma bem simplista, o algoritmo pode ser entendido como um conjunto de passos para realizar uma determinada tarefa*” (ARANTES; BLUM, 2019).

Quanto melhor desenvolvido um algoritmo, mais “certeza” poderá ser comercializada, pois aquele corresponde a um sistema que faz uso de uma tecnologia persuasiva, que induz as pessoas a terem um determinado comportamento.

Importante salientar que a estrutura matemática do algoritmo em si não é, por definição, machista, nem racista, nem está eivada de nenhum tipo de discriminação. As máquinas simplesmente aprendem com os exemplos que os humanos ensinam, são alimentadas por dados fornecidos por humanos, e essa informação incorpora o passado, e de acordo com Cathy O’Neil (O’NEIL, 2020), não se trata de apenas um passado recente, mas de um passado obscuro.

A preocupação existente, diante do contexto já exposto, é que o algoritmo reproduz falas e comportamentos discriminatórios. De acordo com Maria Cristine Lindoso (2021, p. 136), a manipulação da tecnologia preditiva potencializa a discriminação:

(...) a estrutura matemática em si e a manipulação efetiva da tecnologia preditiva possuem relevância fundamental no potencial discriminatório das decisões automatizadas. Por vezes, os agentes se eximem de responderem pelos resultados obtidos através desses processos, sob o argumento de que a tecnologia opera por si só, de forma opaca e pouco controlável. É frequente a obtenção de resultados discriminatórios que são justificados pelos agentes como sendo resultados que não poderiam ter sido previstos, em razão da opacidade do funcionamento dos algoritmos e da natureza da capacidade preditiva do *data mining*, que opera por correlações e inferências estatísticas que fogem das possibilidades de monitoramento dos programadores.

Adicionalmente, um conceito interessante de IA é o destacado abaixo (LIMA; NÓBREGA, 2019):

podemos conceituar a inteligência artificial como o conjunto de arranjos de tecnologia física (*hardware*) e lógica (*software*) organizados de maneira que a máquina possa aprender sozinha com os dados informados pelo homem e por este treinado ao ponto de a máquina ser possível prever determinadas situações a partir do reconhecimento de padrões.

A inteligência artificial não é desenvolvida pensando na coletividade, ela é focada nos usos comerciais, visando o lucro. De acordo com O'Neil, não existe democracia no seu desenvolvimento.

A IA é frequentemente considerada mais objetiva do que os humanos. Na realidade, no entanto, os algoritmos de IA tomam decisões com base em dados anotados por humanos, que podem ser tendenciosos e excludentes.

Caroline Criado Perez alerta (PEREZ, 2021): *“em vez de melhorar o mundo extraíndo vieses históricos de gênero, o aprendizado de máquina, ao que parece, está entrenchando ainda mais a discriminação”*. Ela destaca, ainda, que *“seja um software de reconhecimento de voz incapaz de detectar a voz feminina, ou algoritmos que preferem currículos masculinos, o Big Data está introduzindo novas formas de viés de gênero”*.

Por falar em currículos, um caso que ficou bastante conhecido foi o do sistema de recrutamento da Amazon³, o qual - embora tenha sido posteriormente abandonado pela empresa ao serem evidenciados seus comportamentos tendenciosos - possuía um mecanismo de seleção de perfis de candidatos que terminava rejeitando currículos e candidaturas de mulheres qualificadas somente em razão de seu gênero.

³ Disponível em: <https://www.reuters.com/article/us-amazon-com-jobs-automation-insight-idUSKCN1MK08G>.

Outro caso digno de nota foi o da Apple Card - ocorrido em meados de novembro de 2019 e que gerou bastante repercussão no Twitter -, depois que o empresário de Tecnologia David Heinmeier Hansson escreveu que a Apple Card ofereceu a ele vinte vezes mais limite de crédito do que para a sua esposa, embora ela tivesse uma pontuação mais alta e uma saúde financeira muito mais estável que a dele. Após o relato, muitos usuários da rede social acabaram concordando com David e expuseram outras situações semelhantes pelas quais haviam passado⁴.

Em oposição a essa última afirmação, o Departamento de Serviços Financeiros do Estado de Nova Iorque conduziu uma análise estatística de quase 400.000 candidatos de NY que mostrou que os modelos e algoritmos utilizados para estipular os limites de crédito não consideravam as características 'proibidas' dos candidatos, como o gênero. O que determinava se um cônjuge recebia ou não de crédito, eram diretrizes normais, como o histórico de crédito⁵.

Em um artigo recentemente publicado, a autora Nina da Hora (DA HORA, 2022) afirma que o solucionismo tecnológico, como o que foi dado para o caso anteriormente mencionado, não cabe quando se trata de questões éticas e sociais e que *“Big Data e machine learning são áreas da computação com muitos evangelistas”*.

Nina ainda afirma que a crença de que aparatos tecnológicos possam solucionar todos os problemas é um obstáculo, principalmente por causa da opacidade algorítmica. A opacidade se revela de diversas formas, seja pelo segredo empresarial, seja por desconhecimento de questões técnicas relacionadas às aplicações de IA. A pesquisadora expõe que a opacidade se apresenta em três níveis:

Opacidade intencional: o funcionamento do sistema interno é propositalmente escondido daqueles afetados por sua operação;

Opacidade não-alfabetizada: o funcionamento interno do sistema é opaco porque somente aqueles com conhecimento técnico especializado pode entender como ele funciona;

Opacidade intrínseca: o funcionamento do sistema é opaco devido a uma incompatibilidade entre o entendimento dos algoritmos que são usados para modelar a sociedade no universo digital.

⁴ Disponível em: <https://edition.cnn.com/2019/11/10/business/goldman-sachs-apple-card-discrimination/index.html>.

⁵ Disponível em: <https://www.theverge.com/2021/3/23/22347127/goldman-sachs-apple-card-no-gender-discrimination>.

A opacidade reflete um problema maior: a falta de controle que os agentes - e, em segunda camada, os próprios cidadãos - têm sobre os resultados que são produzidos. Qual a garantia de que o processo automatizado esteja a favor dos interesses dos indivíduos?

Os algoritmos são “caixas pretas” (PERONGINI, 2019), porque os usuários não têm a menor ideia sobre a forma de como as suas tomadas de decisões são feitas, e com isso aumenta o abismo epistêmico entre os usuários e os detentores da tecnologia.

Os dados que são utilizados por essas tecnologias são dados históricos. As decisões são matemáticas e não éticas. Os modelos de aprendizagem de máquina replicam o mundo como ele é. Decisões automatizadas por meio da análise massiva de dados podem, portanto, apresentar efeitos discriminatórios ou criar “bolhas de filtro”, ainda que de forma não intencional, impondo tratamento menos favorável a grupos já desfavorecidos!

Atualmente, os algoritmos determinam quem recebe uma casa, ou quem será contratado. Os algoritmos podem propagar, de forma ampla, a discriminação que tanta gente arriscou e arrisca a vida para combater ao longo dos anos, quiçá dos séculos.

Dado que a tecnologia não é neutra, o discurso largamente difundido de neutralidade falha em considerar a falta de participação das mulheres no processo de construção da tecnologia. A neutralidade não existe, é uma utopia, portanto não existe nas decisões automatizadas.

Para que se possa combater essa pretensa neutralidade, deve-se falar sobre ela, e nesse contexto o accountability surge como um dos principais elementos de governança para assegurar que haja responsabilização e prestação de contas no contexto da IA. A possibilidade de auditoria, quando necessário, e a transparência são essenciais (WIMMER, 2019), senão vejamos:

Por outro lado, a dificuldade em oferecer explicações para os resultados – e, conseqüentemente, para justificar decisões a serem tomadas com base nas inferências identificadas por sistemas dessa natureza – suscita um outro conjunto de preocupações relacionadas ao papel do livre arbítrio individual em contraposição à chamada “ditadura dos dados”. Ao mesmo tempo em que a crescente automatização de tarefas oferece promessas de maior eficiência, objetividade e produtividade, a opacidade dos sistemas de IA, da qual decorre também a dificuldade de rastrear os critérios que conduziram a determinada resposta, tende a suscitar questões difíceis, à medida que aumenta a capacidade de extrair inferências imprevistas e cresce a dificuldade de concretizar ideias ligadas à transparência, à compreensibilidade e à auditabilidade.

A melhor forma para lidar com vieses de discriminação algorítmica é que as equipes de desenvolvimento dos algoritmos sejam plurais e multi e interdisciplinares, com mulheres, negros, pessoas LGBTQIAP+ e profissionais de diversas áreas e campos de conhecimento, tais como profissionais do direito, da psicologia, da sociologia, entre outros.

Nesse sentido, a Conferência Geral das Nações Unidas (UNESCO, 2021) fez um estudo sobre a aplicação ética da inteligência artificial, material que reconheceu os impactos positivos e negativos da IA nas sociedades, no meio ambiente, nos ecossistemas e sobretudo na mente humana, já que as interações das pessoas com essa tecnologia influenciam a cognição dos indivíduos, podendo modificar a forma com que tomam decisões.

A recomendação apresentada pela UNESCO é que a abordagem seja baseada no direito internacional, com foco na dignidade e direitos humanos, bem como a igualdade de gênero, e justiça econômica e desenvolvimento, física e mental, bem-estar, diversidade, interconexão, inclusão, proteção ambiental e do ecossistema para orientar a IA.

Ainda, considera-se que as que as tecnologias de IA podem ser de grande utilidade para a humanidade e todos os países podem se beneficiar, mas ao mesmo tempo levantam preocupações éticas, por exemplo, em relação aos preconceitos que eles podem incorporar e replicar, resultando em discriminação, desigualdade, exclusão digital, significando uma ameaça à diversidade cultural, social e biológica e divisões sociais ou econômicas; necessidade de transparência e compreensão do funcionamento dos algoritmos; e seu impacto potencial sobre os direitos humanos, igualdade de gênero e democracia.

No Brasil, a inteligência artificial ainda é um universo novo - o famoso “oceano azul” -, o que desperta a necessidade de que cada vez mais profissionais sejam capacitado(a)s tecnicamente para atuar nesta seara e, de maneira plural e inclusiva, efetivamente colaborar com o desenvolvimento dos algoritmos e das novas tecnologias, combatendo a discriminação e eventual viés propagado pela inteligência artificial. Por isso, é extremamente importante que mais mulheres ocupem espaços na área de tecnologia, contudo, ainda que não haja essa possibilidade, há que se contratar equipes capacitadas para avaliar a IA, seus riscos e impactos na sociedade e, mais especificamente, em grupos subrepresentados.

O desenvolvimento do algoritmo carece de uma equipe multidisciplinar, com representatividade para conseguir colaborar desde o início, sendo a ética fator indispensável nessa receita, bem como leis claras que determinem como a empresa detentora da IA deve se portar, se adequar e atender as regras básicas e princípios

inafastáveis como o da não discriminação, da dignidade e da equidade, observando os direitos humanos, além de fatores que devem ser vinculantes, tais como a transparência e a boa-fé, não esquecendo a privacidade por design e a proteção de dados.

Ao não se contrapor aos vieses - sejam eles explícitos ou implícitos -, acaba-se por reforçá-los, por isso, deve-se combater a discriminação e impedir que a IA perpetue preconceitos seculares, como o de gênero.

REFERÊNCIAS

CHU, Charlene; LESLIE, Kathleen; NYRUP, Rune; KHAN, Shehroz. **Artificial intelligence can discriminate on the basis of race and gender, and also age**. The Conversation, 2022. Disponível em: <https://theconversation.com/artificial-intelligence-can-discriminate-on-the-basis-of-race-and-gender-and-also-age-173617>.

DA HORA, Nina. **Solucionismo tecnológico não cabe para questões éticas e sociais**. Disponível em: <https://mittechreview.com.br/solucionismo-tecnologico-nao-cabe-para-questoes-eticas-e-sociais/>.

LIMA, Ana Paula Canto de; NÓBREGA, Juliana Targino. **Inteligência artificial: diretrizes, estratégias e verificação nos Tribunais brasileiros**. In: **Direito Exponencial**. Thomson Reuters Brasil, 2019. p. 67-83.

LINDOSO, Maria Cristine Branco. **Discriminação de Gênero no Tratamento Automatizado de Dados Pessoais**. Rio de Janeiro: Editora Processo, 2021.

MCCARTHY, John., et al. **A proposal for the dartmouth summer research project on artificial intelligence**, 1955. Disponível em: <http://www-formal.stanford.edu/jmc/history/dartmouth/dartmouth.html>.

MEDINA, Marco; FERTIG, Cristina. **Algoritmos e Programação, Teoria e Prática**. São Paulo: Novatec, 2005.

O'NEIL, Cathy. **Algoritmos de Destruição em Massa**. São Paulo: Editora Rua do Sabão, 2020.

PASQUALE, Frank. **The Black Box Society: The Secret Algorithms That Control Money and Information**. Harvard University Press, 2015.

PEREZ, Caroline Criado. **Invisible Women: Data Bias in a World Designed for Men**. Abrams Press, 2021.

PERONGINI, Maria Fernanda Hosken de Souza. **"Male by design": um ensaio sobre equidade, discriminação algorítmica por viés de gênero e proteção de dados pessoais**. In: **Direito Exponencial**. Thomson Reuters Brasil, 2019. p. 403-426.

UNESCO, 2021. **Recomendação sobre a Ética da Inteligência Artificial**. Disponível em: https://unesdoc.unesco.org/ark:/48223/pf0000381137_por.

WIMMER, Miriam. **Inteligência artificial, algoritmos e o direito. Um panorama dos principais desafios**. In: **Direito Digital: debates contemporâneos**. 1. ed. São Paulo: Thomson Reuters Brasil, 2019. p. 15-30.

4. Potencial Discriminação Algorítmica Regional

Beatriz Pistarini

Como reforçado nos demais tópicos desta contribuição, a inteligência artificial tem o potencial de ajudar os seres humanos a tomarem decisões isentas e até mais equânimes, porém, isso somente acontecerá com um trabalho cuidadoso para garantir que seus sistemas alcancem a imparcialidade desejada.

O uso crescente da inteligência artificial em áreas sensíveis, como contratação de pessoas, justiça criminal e saúde, gerou muita discussão sobre viés (tendência) e imparcialidade. No entanto, as decisões tomadas por seres humanos nesses e em outros campos também podem conter falhas por serem influenciadas por preconceitos individuais e sociais, e muitas vezes inconscientes. O grande entrave dessa questão é se seria a IA responsável por agravar cenários “analógicos” de opressão.

Quando se fala sobre viés ou tendência em algoritmos de IA, trata-se de uma distorção do julgamento do observador, podendo se manifestar como uma inclinação irracional a atribuir um julgamento mais favorável ou desfavorável a alguma coisa, pessoa ou grupo.

No que tange a discriminação algorítmica regional, já existem pesquisas referentes, levando a acreditar que se trata de uma realidade problemática. A partir do estudo realizado pelo *National Bureau of Economic Research*, constatou-se que, no mercado de hipotecas dos Estados Unidos, ao solicitar crédito para comprar um imóvel, os clientes latinos e afro-americanos geralmente acabam pagando 7,9 pontos-base (0,079 ponto percentual) a mais do que os brancos para pagar a hipoteca, e 3,6 pontos-base a mais quando refinanciam suas dívidas.

Ainda, outro caso que vale a menção é o projeto *PredPol* (abreviação do termo em inglês “policiamento preditivo”), desenvolvido por uma startup norte-americana e que tem como objetivo a atuação na prevenção de crimes, com uso de inteligência artificial na aplicação de patrulhamento. Verificou-se que o *PredPol* estabelecia critérios regionais, classificando preemptivamente determinadas áreas e grupos como “sociedade marginalizada” previamente demarcada, ignorando outras atenuantes e elementos relevantes e passando a estabelecer uma relação direta de potenciais cometimentos de crimes a partir do uso de dados presentes em redes sociais, ou de qualquer outra forma, gerados pelos indivíduos.

No Brasil, ainda não há pesquisas robusta e largamente difundidas sobre discriminação algorítmica regional, o que desperta preocupação elevada, especialmente ao considerarmos as dimensões continentais do País. Se fizermos um paralelo, inclusive, ampliando a reflexão para o ambiente digital como um todo, há estudos e evidências que denotam a existência de discriminação em razão de particularidades regionais, principalmente em relação às regiões Norte e Nordeste, o que se evidencia em discursos de ódio difundido nas mídias sociais. Estes discursos de ódio propagam a discriminação decorrente de ideias infundadas, que restam por prestar um desserviço ao fomentar uma visão negativa e pejorativa dessas regiões, além de não serem devidamente punidos à altura de suas consequências, o que constitui um problema no âmbito jurídico.

Desta maneira, se situações de teor discriminatório regional no ambiente virtual já ocorrem com certa frequência, não há dúvidas de que tais perspectivas também serão refletidas nos algoritmos de IA, visto que sua precisão depende, em grande parte, da qualidade dos dados e dos parâmetros codificados nos algoritmos de treinamento, além do fato de que, como já dito acima, a IA reflete os padrões de pensamento de quem a criou.

Portanto, a depender de quem desenvolveu determinado algoritmo de IA, as segregações já praticadas na sociedade poderão aumentar, tornando um ambiente excludente e baseado em discursos de ódio, marginalizando (inclusive no meio digital) grupos sociais.

Sendo assim, é fundamental que as equipes de desenvolvimento de sistemas que utilizem IA sejam multidisciplinares, contendo profissionais não somente da área de tecnologia, mas também sociólogos, juristas, bem como experts em áreas diversas, tais como etarismo, psicologia, sexualidade, saúde e outros temas. Quanto mais diversidade houver nas equipes de desenvolvimento dos sistemas que utilizam IA - espera-se -, menos erros e vieses elas terão.

REFERÊNCIAS:

DETRIXHE, John. **Fintech algorithms discriminate 40% less than traditional lenders**. Disponível em: <https://qz.com/1647701/fintech-algorithms-discriminate-40-less-than-traditional-lenders/>.

LUCENA, Pedro Arthur Capelari de. **Policimento preditivo, discriminação algorítmica e racismo: potencialidades e reflexos no Brasil**. Disponível em: <https://lavits.org/wp-content/uploads/2019/12/Lucena-2019-LAVITSS.pdf>.

PINTO, Danielle Jacón Ayres; SOUZA, Elany Almeida de. **A Discriminação em relação às regiões Norte e Nordeste do Brasil, presente no discurso de ódio difundido nas mídias sociais**. Disponível em: <http://www.publicadireito.com.br/artigos/?cod=37e9b839eeb8b2d3>.

PREDPOL, Aim Samuels. **The Myth of Crime Displacement**. Disponível em: <https://www.predpol.com/crime-displacement-predpol/>.

SILBERG, Jake; MANYIKA, James. **Tackling bias in artificial intelligence (and in humans)**. Disponível em: <https://www.mckinsey.com/featured-insights/artificial-intelligence/tackling-bias-in-artificial-intelligence-and-in-humans>.

5. Potencial Discriminação Algorítmica em razão da Idade

Gisele Truzzi
Josemara Meireles

Conforme pesquisa anual coordenada pela PwC e realizada com diretores e tomadores de decisão, foi demonstrado que, à medida que enfrentamos a crise atual, o otimismo em relação à IA aumentou significativamente de 72% para 92%, e a taxa de adoção de IA aumentou de 62% para 70%. Além disso, 94% dos entrevistados afirmam que implementaram ou planejam implementar a IA em suas organizações.

Levando em consideração os dados estatísticos apontados acima, que certamente evidenciam o quão presente está a inteligência artificial no mundo corporativo, fica claro que processos de recrutamento e seleção também se incluem nessa dinâmica. Dessa forma, atividades de recursos humanos também se utilizam de aplicações da inteligência artificial, podendo ocorrer na forma de assistentes virtuais, que desempenham diversas tarefas desde o início — por exemplo, quando o candidato precisa responder em um site as primeiras questões para inscrição no certame —, chegando ao emprego de robôs responsáveis por avaliar comportamentos, habilidades e competências das pessoas em situações reais.

Em tais processos, a discriminação sempre foi um tema de significativa relevância, mesmo porque, ainda que de forma involuntária, práticas excludentes estão impregnadas na nossa cultura, ainda que não toleradas pelo nosso ordenamento jurídico.

Uma pesquisa feita pelo InfoJobs — empresa de tecnologias para recrutamento — mostra que 70,4% dos profissionais com mais de 40 anos entrevistados já sofreram preconceito no mercado de trabalho por conta de sua idade. A pesquisa revela ainda que, na percepção de 78,5% dos entrevistados, o mercado não dá as mesmas chances para profissionais com mais de 40 anos quando comparado aos mais jovens.

É certo que um sistema com tecnologia de IA pode facilmente cruzar informações com banco de dados de contratações anteriores e criar um padrão, por exemplo, escolhendo candidatos com idades iguais aos indivíduos contratados em outros processos, abrindo margem, assim, para a tomada de decisões repletas de vieses, excluindo candidatos com idade acima das idades do “padrão” encontrado nos bancos de dados.

De acordo com as informações aqui trazidas, podemos chegar à conclusão de que a discriminação algorítmica por idade possui um grande recorte cultural já naturalmente inserido no País, que não tem por hábito a valorização dos indivíduos mais velhos.

Logo, ressaltando o que já foi mencionado anteriormente, ao se realizar o tratamento de dados pessoais que abrem margem para o etarismo, faz-se com que o algoritmo selecione previamente candidatos com determinada faixa etária, excluindo sumariamente inúmeras pessoas com habilidades excelentes que poderiam ser capacitadas para ocupar determinada vaga.

De acordo com Carla Reita Faria Leal e Antonio Raul Alencar, inclusive:

Diante desse cenário, e considerando que as ferramentas de seleção automatizadas já representam, segundo estudiosos, um negócio de 500 milhões de dólares anuais, tendente a crescer de 10 a 15% ao ano, a discriminação algorítmica é uma questão que tende a se tornar mais frequente, demandando vigilância e aprimoramento dos mecanismos de fiscalização, orientação e monitoramento no que se refere ao tratamento de dados do candidato ao emprego.

Conclui-se, portanto, que o risco potencial discriminatório em razão da idade é muito elevado, seja pelo exponencial uso da inteligência artificial em processos de recrutamento e seleção, seja pela falta de regras claras no desenvolvimento dessas tecnologias, seja pela falta de diversidade dos atores envolvidos nos fluxos de deliberação existentes.

Nesse contexto, e diante da conclusão de que o uso da inteligência artificial em processos seletivos atrai o risco de potencial discriminatório em razão da idade, o que se propõe é a utilização de uma abordagem de hierarquização dos riscos, classificando o risco potencial discriminatório em razão da idade como sendo inaceitável pelos agentes, com a consequente criação de regras robustas de governança que permitam o gerenciamento de riscos de maneira preventiva.

Sabendo do elevado potencial de discriminação pelos sistemas desenvolvidos com base em inteligência artificial, o legislador deverá aprofundar a análise de riscos e, em seu exercício regulatório, permitir a criação de regras bem definidas para que as empresas que desenvolvem novas tecnologias possam incorporar, desde a concepção, reflexões com o objetivo de evitar ou mitigar o risco potencial discriminatório em razão da idade.

REFERÊNCIAS:

ALVES, Walter. **Algoritmo do preconceito**. Blog Maturi, 2022. Disponível em: <https://www.maturi.com.br/diversidade/algoritmo-do-preconceito/>.

BELMONTE, Camile. **O Etarismo no mercado de trabalho: preconceito de idade existe?** PlaceRH, 2022. Disponível em: <http://blog.placerh.com.br/2021/07/22/etarismo/>.

BELCHIOR, Sales Wilson. **Inteligência Artificial, princípios e recomendações da OCDE**. Migalhas, 2022. Disponível em: <https://www.migalhas.com.br/depeso/330983/inteligencia-artificial--principios-e-recomendacoes-da-ocde>.

CÓPPOLA, Giovana. **Inteligência artificial no recrutamento: como funciona esta tendência?** Inquietaria, 2022. Disponível em: <https://inquietaria.99jobs.com/intelig%C3%A2ncia-artificial-no-recrutamento-como-funciona-esta-tend%C3%A2ncia-d1d33ecd5b1c>.

OPAS. **Discriminação por idade é um desafio global, afirma relatório da Organização das Nações Unidas**. Organização Pan-Americana da Saúde, 2022. Disponível em:

<https://www.paho.org/pt/noticias/18-3-2021-discriminacao-por-idade-e-um-desafio-global-afirma-relatorio-da-organizacao-das>.

Discriminação etária na Inteligência Artificial ameaça saúde dos idosos, segundo a OMS. IstoÉ dinheiro, 2022. Disponível em: <https://www.istoedinheiro.com.br/discriminacao-etaria-na-inteligencia-artificial-ameaca-saude-dos-idosos-segundo-a-oms/>.

FILHO, Demócrito Reinaldo. **A Proposta Regulatória da União Europeia para a Inteligência Artificial (1ª parte) – A hierarquização dos riscos.** Juristas, 2022. Disponível em: <https://juristas.com.br/2021/05/21/a-proposta-regulatoria-da-uniao-europeia-para-a-inteligencia-artificial-1a-parte-a-hierarquizacao-dos-riscos/>.

MARCONDES, Leticia. **Etarismo: 6 sinais de discriminação por idade no trabalho.** Safespace, 2022. Disponível em: <https://safe.space/conteudo/etarismo-sinais-discriminacao-por-idade-trabalho-ageismo>.

PASCUAL, Manuel G. **Quem vigia os algoritmos para que não sejam racistas ou sexistas?.** Elpaís, 2022. Disponível em: [Quem vigia os algoritmos para que não sejam racistas ou sexistas? | Tecnologia | EL PAÍS Brasil \(elpais.com\)](https://elpais.com/tecnologia/2022/05/19/quien-vigila-a-los-algoritmos-para-que-no-sean-racistas-o-sexistas-20220519.html).

TUNES, Suzel. **A parcialidade dos algoritmos.** Nexo, 2022. Disponível em: <https://www.nexojournal.com.br/externo/2019/11/24/A-parcialidade-dos-algoritmos>.

70% dos profissionais com mais de 40 anos já sofreram preconceito no mercado de trabalho, mostra pesquisa. Portal G1-Economia, 2022. Disponível em: <https://g1.globo.com/economia/concursos-e-emprego/noticia/2021/05/19/70percent-dos-profissionais-com-mais-de-40-anos-ja-sofreram-preconceito-no-mercado-de-trabalho-mostra-pesquisa.ghtml>.

IA: uma oportunidade em meio à crise. Pwc, 2022. Disponível em: <https://www.pwc.com.br/pt/estudos/preocupacoes-ceos/mais-temas/2021/reinventando-o-futuro/ia-uma-oportunidade-em-meio-a-crise.html>.

LEAL, Carla Reita Faria. ALENCAR, Antônio Raul. **A discriminação algorítmica na seleção do emprego.** Disponível em: <https://oliveira.com.br/a-discriminacao-algoritmica-na-selecao-ao-emprego>.

6. Potencial Discriminação Algorítmica Religiosa

Andressa Almeida
Caroline Vilas Boas
Júlyanne de Bulhões

A necessidade de encontrar formas pacíficas de convivência entre as pessoas, especialmente com relação às diferentes crenças - considerando como marco a ruptura da unidade cristã - impulsionou o debate sobre a tolerância religiosa. Como consequência da reforma protestante, grupos religiosos minoritários na sociedade (minorias religiosas) passaram a defender o direito da sua própria interpretação da fé, desenvolvendo a ideia de tolerância religiosa. Esse é o primeiro direito inalienável do homem, inaugurando a concepção contemporânea dos direitos fundamentais (CANOTILHO, 2003).

A Declaração Mundial de Princípios sobre a Tolerância, aprovada pela Conferência Geral da UNESCO em 1995, está inserida em um contexto de busca de aperfeiçoamento humano,

com impacto no cotidiano dos indivíduos. No texto final, a Declaração assim define tolerância (UNESCO, 1995):

A tolerância é o respeito, a aceitação e o apreço da riqueza e da diversidade das culturas de nosso mundo, de nossos modos de expressão e de nossas maneiras de exprimir nossa qualidade de seres humanos. É fomentada pelo conhecimento, a abertura de espírito, a comunicação e a liberdade de pensamento, de consciência e de crença. A tolerância é a harmonia na diferença. Não é só um dever de ordem ética; é igualmente uma necessidade política e jurídica. A tolerância é uma virtude que torna a paz possível e contribui para substituir uma cultura de guerra por uma cultura de paz.

Chama-se liberdade religiosa a proibição ao Estado de impor ou impedir a confissão de uma fé religiosa, além de permitir ou propiciar a quem deseja seguir determinada religião com o cumprimento dos deveres que dela decorrem (BRANCO; JACOBINA, 2017). Trata-se do direito de ter ou não uma religião declarada, mudar de confissão religiosa e exercê-la, **sem interferência, mas com a garantia estatal do livre exercício pelo cidadão**. A liberdade religiosa é, portanto, um direito humano fundamental e expressão do fundamento constitucional da dignidade da pessoa humana (SANTOS JUNIOR, 2007).

Ainda que, à primeira vista, pareça desnecessário rememorarmos, hoje, as definições de tolerância, liberdade religiosa e como estes temas se encaixam dentro do ordenamento jurídico, dado seu indiscutível status de direito fundamental, é crucial evidenciarmos esses institutos também para garantir que os avanços tecnológicos os reflitam, com o apoio do Estado.

As informações pessoais relacionadas à convicção religiosa e/ou filiação a organização de caráter religioso são dados pessoais sensíveis, conforme ordena a Lei Geral de Proteção de Dados Pessoais, marco regulatório que representou um passo importante para a defesa de direitos personalíssimos. Dados pessoais sensíveis - como as informações relacionadas à religião, crença e fé - requerem uma maior e necessária proteção pelos agentes de tratamento, porque essas informações podem ser utilizadas, ainda que de forma não intencional, para práticas discriminatórias direcionadas a indivíduos.

Some-se ao debate, as inúmeras possibilidades tecnológicas: sistemas inteligentes são hoje capazes de processar inúmeras informações, inclusive em um cenário em que a fé esteja relacionada.

A inteligência artificial, se se concretizarem as expectativas, tende a confrontar o antropocentrismo de forma mais radical do que as tecnologias digitais, deslocando o centro

do poder na sociedade: “As novas tecnologias do século XXI podem, assim, reverter a revolução humanista, destituindo humanos de sua autoridade e passando o poder a algoritmos não humanos” (HARARI, 2016, p. 347).

As soluções aplicadas à inteligência artificial podem, mesmo que em maior raridade, se voltar para religião ou igrejas. O *ChurchIX*, aplicação com a funcionalidade de monitorar a assiduidade dos fiéis, é capaz de gerar relatórios de frequência e horário de chegadas aos cultos, inclusive considerando critérios de sexo e idade de cada participante do culto (BRAVO, 2019). O *BlessU-2*, por sua vez, é um pastor robô implementado na Alemanha que, através de um painel *touchscreen*, oferece bênçãos em cinco idiomas, em voz feminina ou masculina, de acordo com as escolhas do fiel (GALILEU, 2017).

Apesar de os exemplos acima serem todos pensados para impulsionar - e, portanto, “enaltecer” - alguma crença, não se pode perder de vista que o viés discriminatório também está presente no cenário tecnológico. As empresas, e seus desenvolvedores, adotam determinados padrões para orientar as ações executadas por estes equipamentos que, por sua vez, reproduzem valores e atitudes discriminatórias aprendidos e reproduzidos ao longo da vida pelas pessoas.

A inteligência artificial, assim como muitas das novas tecnologias desenvolvidas pela humanidade, perpetua o preconceito porque foi desenvolvida com essa base, o que pode se estender também na problemática do preconceito religioso a partir do processo de análise de dados em larga escala. Os algoritmos enviesados dependem da configuração desenvolvida pelos humanos que os constroem, por isso é fundamental considerar o papel dos algoritmos na tomada de decisões, especialmente nos casos em que os algoritmos desempenham um papel fundamental na formação de um processo de decisão, mesmo quando a decisão final é feita por humanos.

Em 2020, o lançamento do GPT-3, software considerado um dos mais sofisticados em aprendizado de máquinas (machine learning) desenvolvido por uma empresa chamada OpenAI, demonstrou a capacidade de processar e reproduzir linguagens para geração de textos, como livros, roteiros, artigos, dentre outras produções, como se fosse um ser humano. O feito por si só soa incrível, se não fosse pelo fato de que além de reproduzir a linguagem humana, o sistema de geração de textos reproduz também os vieses presentes em diversos campos na sociedade, e a religião é um deles.

Em um artigo publicado na revista científica *Nature Machine Intelligence*, pesquisadores de Stanford, liderados pelo pesquisador muçumano Abubakar Abid, evidenciaram as associações discriminatórias feitas pela ferramenta de construção textual, especialmente,

quando ligado à religião mulçumana de forma desproporcional, em comparação às demais religiões. Através dos testes realizados, fica perceptível a constatação de que, ao imputar palavras relacionadas à religião mulçumana, a ferramenta automaticamente realiza a associação com “violência” e “terrorismo”. É possível ainda perceber associações discriminatórias com outras religiões, como por exemplo, o judaísmo, em que o software associa a religião a “dinheiro”.

Outro exemplo prático de preconceito e discriminação envolvendo religião consiste no fato de que até mesmo em leitura de nomes que poderiam fazer parte da religião mulçumana foi identificada rejeição algorítmica, como é possível verificar no estudo detalhado no artigo científico *Racial Discrimination in the Sharing Economy: Evidence from a Field Experiment* (EDELMAN *et al*, 2017).

Ainda que para máquinas seja difícil a incorporação de todas as variantes de filosofias éticas e morais, visto que há uma série de divergências e crenças que os humanos carregam através de contato com novas experiências do mundo exterior que mudaram a sua forma de pensar, há de se pontuar a importância da diversidade que deve estar presente na configuração dessas tecnologias.

A inteligência artificial é reflexo da visão de mundo do homem, o algoritmo não conhece o argumento moral. (...) Se os dados são de baixa qualidade, o resultado também será ruim. (NOOMIS, 2020).

Na seara dos dados, cuidados especiais devem ser tomados nas fases de desenvolvimento e implantação, especialmente quando o processamento é direto ou indiretamente baseado em dados "sensíveis", como é o caso do dado religioso. Esses dados podem incluir origem racial ou étnica, formação socioeconômica, opiniões políticas, crenças religiosas ou filosóficas, associação sindical, dados genéticos, dados biométricos, dados relacionados à saúde, entre outros.

Em que pese o entendimento de que o algoritmo de visão computacional não tem absolutamente nenhuma visão subjetiva, é preciso pontuar que essa visão é baseada em dados e em variantes inseridas por seres humanos, o que não torna tais dispositivos isentos de julgamento e de outputs gerados com base em preconceito religioso ou de qualquer outra natureza.

É, portanto, crucialmente importante a discussão sobre como garantir que soluções tecnológicas também considerem cenários para garantir liberdade religiosa, o direito de divulgação da fé e a não-discriminação em razão dessas informações. A garantia desse

cenário está diretamente relacionada ao exercício da dignidade humana, essencial ao exercício do ser humano por sua própria condição.

REFERÊNCIAS:

BRANCO, Paulo Gustavo Gonet; JACOBINA, Paulo Vasconcelos. **Liberdade de Gueto? Religião e espaço público**. In: DIP, Ricardo; FERNANDES, André Gonçalves (coord.). *Laicismo e Laicidade no Direito*. São Paulo: QuartierLatin, 2017. p. 153-163

BRAVO, Luiza. **Reconhecimento facial: câmeras inteligentes nas igrejas**. 2019. Disponível em: <https://www.whow.com.br/global-trends/reconhecimento-facial-cameras-inteligentes-chegaram-igrejas/>.

CANOTILHO, José Joaquim Gomes. **Direito Constitucional e Teoria da Constituição**. 7. ed. Coimbra: Almedina, 2003

EDELMAN, Benjamin; LUCA, Michael, SVIRSKY, Dan. **Racial Discrimination in the Sharing Economy: Evidence from a Field Experiment**. Disponível em: <https://pubs.aeaweb.org/doi/pdfplus/10.1257/app.20160213>.

GALILEU. **Bênção do futuro: alemães criam robô pastor**. 2017. Disponível em: <https://revistagalileu.globo.com/Tecnologia/noticia/2017/05/bencao-do-futuro-alemaes-criam-robo-pastor.html>.

KAUFMAN, Dora. **Inteligência Artificial: Questões Éticas A Serem Enfrentadas**. Disponível em: https://abciber.org.br/analseletronicos/wp-content/uploads/2016/trabalhos/inteligencia_artificial_questoes_eticas_a_serem_enfrentadas_dora_kaufman.pdf.

NOOMIS. **Treine bem seus algoritmos para evitar decisões tendenciosas**. 2020. Disponível em: <https://noomis.febraban.org.br/temas/inteligencia-artificial/treine-bem-seus-algoritmos-para-evitar-decisoes-tendenciosas>.

SANTOS JUNIOR, Aloisio Cristovam dos. **A liberdade de organização religiosa e o Estado laico brasileiro**. São Paulo: Mackenzie, 2007.

SILVA, Tarcizio. **Racismo Algorítmico Em Plataformas Digitais: Microagressões E Discriminação Em Código**. Disponível em: https://d1wqtxts1xzle7.cloudfront.net/65462972/Comunidades_Algoritmos_e_Ativismos_olhar-with-cover-page-v2.pdf.

TAURION, Cezar. **Inteligência artificial: até os algoritmos têm “preconceito”**. Disponível em: <https://neofeed.com.br/blog/home/inteligencia-artificial-ate-os-algoritmos-tem-preconceito/>.

UNESCO. **Declaração de Princípios sobre a Tolerância**. Paris, 1995. Disponível em: <http://www.dhnet.org.br/direitos/sip/onu/paz/dec95.htm>.

7. Conclusão: essencialidade de análise preventiva do potencial discriminatório

Melisa Curilla Rapelli
Nathália Criscito

Esse tópico conclusivo visa amarrar todos os pontos trazidos anteriormente para fins de reforço da importância de uma análise preventiva do potencial discriminatório de inteligências artificiais, eis que robustos e recentes estudos evidenciam que o caráter discriminatório presente no comportamento humano pode ser reproduzido por algoritmos que, por sua vez, acabam refletindo vícios do comportamento social.

Estudos apontam que o reconhecimento facial já foi utilizado para justificar a prisão de homens negros e beneficiaram homens brancos, mesmo o primeiro grupo tendo deixado de cometer qualquer infração ou cometido infrações mais leves que o segundo grupo citado, como é o caso ocorrido em Detroit, em julho de 2019⁶, em que a polícia local deteve um jovem negro de 26 anos sob acusação de furto, equivocadamente.

O algoritmo era utilizado pelos Tribunais de Justiça dos EUA para tomadas de decisões nos julgamentos, como ferramenta para medir o grau de periculosidade apresentado por cada réu, sendo que o diagnóstico gerado servia como parâmetro para a concessão ou não de fiança, liberdade condicional e outros instrumentos de relaxamento de pena. O algoritmo gerou polêmicas e ficou comprovado, via pesquisas realizadas pela “ProPublica”⁷, que o sistema de inteligência artificial empregado era tendencioso no sentido de gerar *outputs* muito mais severos em suas decisões contra réus negros.

Em suma, a inteligência estabelecia uma nota de 1 a 10 para medir a probabilidade de reincidência de um réu com base em mais de 100 fatores, incluindo idade, sexo e história criminal. Ainda que, a priori, a informação de raça não fosse utilizada como insumo primário na análise, constatou-se que os padrões de classificação que eram conferidos para grupos de homens brancos e grupos de homens negros eram completamente destoantes. Entre os réus que não reincidiram, os negros tinham duas vezes mais chances do que os brancos de serem classificados como de médio ou alto risco (42% contra 22%), bem como de terem sua fiança negada, além de outros impactos severos em suas vidas.

Ainda em 2013, os pesquisadores Sarah Desmarais e Jay Singh examinaram 19 metodologias de risco diferentes usadas nos Estados Unidos e descobriram que *“na maioria dos casos, a validade foi examinada apenas em um ou dois estudos”* e que *“frequentemente, essas investigações foram concluídas pelas mesmas pessoas que desenvolveram o instrumento”*.

Não distante dali, no ano de 2020, outro jovem negro foi preso depois de um software de reconhecimento facial indicá-lo como o suspeito de um crime. Nesse segundo caso, o jovem foi “confundido” com o autor do crime, que era consideravelmente diferente dele, caso que foi pauta de matéria em jornal televisivo de grande notoriedade no Brasil.

O que se observa em ambos os casos - sem mencionar incontáveis contextos em que este tipo de abordagem discriminatória acontece - é que não existe neutralidade nas soluções

⁶ Disponível em: <https://www.uol.com.br/tilt/noticias/redacao/2020/09/07/pela-2-vez-homem-negro-e-presos-nos-eua-apos-erro-de-reconhecimento-facial.htm>.

⁷ Disponível em: <https://www.propublica.org/article/machine-bias-risk-assessments-in-criminal-sentencing>.

tecnológicas utilizadas em nossa sociedade, pois elas reproduzem o contexto social daqueles que as desenvolvem, daí a importância da diversidade nas equipes de construção e produção tecnológica.

Fica evidente também a necessidade de se analisar preventivamente o potencial discriminatório dos dados inseridos em sistemas de inteligência artificial, considerando as inúmeras limitações humanas, evitando-se que as mesmas sejam reproduzidas em ambiente tecnológico e, por consequência, perpetuando a discriminação, o racismo, o sexismo, o machismo, dentre tantas outras problemáticas de caráter social.

Diante das considerações acima, sugere-se que, em uma primeira camada, seja obrigatória a composição de equipes que possuam diversidade de representatividade nos times de desenvolvimento de inteligência artificial, o que viabiliza a construção de tecnologias sobre olhares e realidades diversas.

Tanto é assim que um dos 17 Objetivos de Desenvolvimento Sustentável da Agenda 2030 das Nações Unidas⁸ deixa clara a importância da redução de desigualdades a partir da promoção de inclusão social, econômica e política de todos os indivíduos, independentemente da idade, gênero, deficiência, raça, etnia, origem, religião, condição econômica, entre outros. Tal aspecto, por si só, deveria ser considerado em todos os espaços de discussão e dinâmicas sociais, mas se mostra ainda mais necessário quando se está falando da implementação de novas tecnologias que repercutam na vida humana.

Em uma segunda camada, sugere-se que existam avaliações de risco e impacto, bem como auditorias algorítmicas de resultados também compostas por olhares diversos e multidisciplinares, principalmente quando tratamos de tecnologias que visam promover a segurança pública e assegurar os interesses sociais. O processo de auditagem, nesse contexto, deve ocorrer para que se questionem e validem as formas de desenvolvimento, coleta e processamento desses dados a fim de evitar arbitrariedades em razão de eventual subjetividade inerente ao comportamento humano, evitando-se o impacto negativo dessas tecnologias.

Em uma terceira e última camada, sugere-se a implementação de processos periódicos e regulares de monitoramento de resultados em tecnologias em funcionamento real, de modo a monitorar e mitigar riscos de impactos discriminatórios que não tenham sido identificados na primeira e segunda camada aqui indicadas.

⁸ Disponível em: <https://brasil.un.org/pt-br/sdgs>.

A título de referência, inclusive, vale mencionar as seguintes estratégias discutidas no cenário internacional - as quais, embora não enderecem suficientemente todas as nuances e especificidades do contexto brasileiro, podem servir como ponto de partida para as discussões travadas no País:

- AI and data protection risk toolkit, elaborado pelo Information Commissioner's Office (autoridade de proteção de dados do Reino Unido e disponível em: <https://ico.org.uk/for-organisations/guide-to-data-protection/key-dp-themes/guidance-on-ai-and-data-protection/ai-and-data-protection-risk-toolkit/>);
- Framework da OCDE para “Classification of AI systems”, disponível em https://www.oecd-ilibrary.org/science-and-technology/oecd-framework-for-the-classification-of-ai-systems_cb6d9eca-en;jsessionid=j5z3QhUiR4zRm4Bu2dglp4cg.ip-10-240-5-155 ;
- Report on artificial intelligence in a digital age do Parlamento Europeu, disponível em: https://www.europarl.europa.eu/doceo/document/A-9-2022-0088_EN.html;
- Materiais desenvolvidos (em progresso) pelo Council of Europe na discussão de um tratado internacional sobre regulação de IA e direitos humanos, disponíveis em: <https://www.coe.int/en/web/artificial-intelligence/work-in-progress> .

Nosso posicionamento vai na linha de que, qualquer que seja a estratégia regulatória adotada no cenário brasileiro, esta esteja comprometida em incentivar e assegurar que o desenvolvimento de tecnologias continue sendo fomentado, mas sempre de forma a permitir que sistemas de inteligência artificial, uma vez postos para uso da sociedade, contribuam para o alcance de resultados sociais, éticos e transparentes, assegurando a equidade no tratamento dos cidadãos.

Com o apoio de equipes técnicas compostas por pessoas diversas, conceitos básicos em ciência da computação - como abstração e design modular -, caminharemos de maneira mais próxima ao objetivo de alcance concreto das noções de justiça e não discriminação, de modo a produzir algoritmos de aprendizagem mais éticos, plurais e assertivos, contribuindo para o efetivo desenvolvimento da vida em sociedade.

Nas palavras de Danilo Doneda e Luca Belli:

Enfim, existe um último desafio crucial para que qualquer sistema regulatório seja verdadeiramente sustentável: como implementar a regulação que queremos? E como levar a regulação até o desenvolvedor, para que seja eficaz, em um ecossistema que, em boa parte, inclusive extrapola as fronteiras nacionais?

A experiência da proteção de dados em geral nos ensina que criar uma lei e uma autoridade reguladora é um esforço enorme, mas na verdade, é somente a primeira etapa.

Continuemos, portanto, na luta por uma sociedade mais justa, ética e inclusiva, que prestigie o desenvolvimento tecnológico, a economia criativa e a inovação sem se distanciar do respeito e garantia da dignidade da pessoa humana em todas as suas multiplicidades de manifestação.

REFERÊNCIAS:

DONEDA, Danilo; BELLI, Luca. O que a regulação da inteligência artificial pode aprender da proteção de dados? JOTA Info, 2021. Disponível em: <https://www.jota.info/opiniao-e-analise/artigos/inteligencia-artificial-protecao-dados-governanca-internet-05112021>.

SINGH, J. P.; DESMARAIS, S.; VAN DORN, R. A. **Measurement of predictive validity in violence risk assessment studies: A second-order systematic review**. Behavioral Sciences & the Law, Volume 31, Issue 1, 2013. Disponível em: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2473262.

SELBST, Andrew D. et.al. **Fairness and Abstraction in Sociotechnical Systems**. FAT* '19: Proceedings of the Conference on Fairness, Accountability, and Transparency, 2019. Disponível em: <https://dl.acm.org/doi/10.1145/3287560.3287598>.

Junho de 2022.

Subscrevem o presente documento:

Ana Paula Canto de Lima: Advogada com atuação em Direito Digital Privacidade e Proteção de dados, especializada em Direito Digital, mestra pela UFRPE em Ciências do Consumo, fundadora do escritório Canto de Lima Advocacia. Cofundadora do curso LGPD Learning, e da plataforma “Cadê meu dado?”. Professora de cursos e pós-graduações e extensão de Direito Digital, Privacidade e Proteção de Dados em diversos Estados, como a Damásio Educacional, Escola da Magistratura - Ajuris, e a Escola da Magistratura Federal - ESMAFE. Presidente da Comissão de Privacidade e Proteção de Dados da OAB/PE, preside a Comissão de Crimes Cibernéticos da Academia Brasileira de Ciências Criminais (ABCCRIM), é membro do Instituto de Juristas Brasileiros (IJB), entre outros trabalhos voluntários, destaca-se o dedicado às escolas públicas, onde leva informações sobre segurança na Internet, que é o projeto pelo qual tem um carinho especial.

Ana Cristina Oliveira Mahle: Advogada, especializada em Proteção de Dados e Direito Digital, Membro Regional da Comissão de Privacidade e Proteção de Dados da OAB/SP; Mestre em Ciência Tecnologia e Sociedade.

Andressa Martins Andrade Almeida: Advogada atuante na área de privacidade e proteção de dados pessoais. Membro do grupo Mulheres na Privacidade e pesquisadora voluntária no Instituto Goiano de Direito Digital.

Beatriz Pistarini: Advogada em Truzzi Advogados. Atuante nas esferas consultiva e contenciosa de Direito Digital, Privacidade e Proteção de Dados. Graduada em Direito pela PUC-Campinas. Pós-graduada em Direito Digital e Compliance pela Faculdade IBMEC São Paulo. Cofundadora do Grupo de Estudos de Direito Digital (GEDD) da PUC-Campinas. Membro efetivo da “Comissão Especial de Tecnologia e Inovação” da OAB São Paulo. Membro efetivo da “Comissão de Estudos em Inovação e Startups” e da “Comissão Especial de Privacidade e Proteção de Dados”, ambas da OAB/SP – 3ª Subseção de Campinas/SP. Membro do Comitê Público da Associação Nacional dos Profissionais de Privacidade de Dados – ANPPD. Coautora de obras jurídicas.

Caroline Vilas Boas: Advogada, especialista com pós-graduação em privacidade e proteção de dados pessoais pela PUC/MG, atualmente trabalhando em empresa irlandesa como gerente de segurança da informação, lidando diretamente com os desafios da GDPR no mercado europeu. Membro do grupo Mulheres na Privacidade.

Eloá de Azevedo Caixeta: Advogada e consultora, graduada pela Universidade de Uberaba e especialista em Direito Digital e Compliance pela Damásio Educacional. Possui curso de extensão em Proteção de Dados pelo Data Privacy Brasil. É coautora da revista Direito Digital, 1ª e 2ª edição, lançada pela Enlaw, e também do livro Direito Digital Aplicado 5.0, lançado em 2022. Facilitadora do grupo Mulheres na Privacidade, atuando a frente de estudos voltados a vieses algorítmicos. É membro do Comitê de Governança para o Desenvolvimento Social e Humano da Rede Governança Brasil. Membro da Comissão de Proteção de Dados da OAB de Uberaba/MG e da OAB de São Paulo/SP. Consultora atuante em Privacidade e Proteção de Dados.

Gisele Truzzi: Advogada especialista em Direito Digital e Segurança da Informação. CEO e sócia-fundadora de Truzzi Advogados. Atua nas esferas consultiva e contenciosa do Direito Digital desde 2005. Professora convidada de cursos de pós-graduação e MBA no Estado de SP (PUC-Campinas, EPD, Proordem). Autora de diversos artigos sobre temáticas relacionadas a direito e tecnologia, publicados em vários portais, tais como IstoÉ Dinheiro e Conjur, entre outros. Coautora das obras: "Direito Digital: Debates Contemporâneos" (Org.: Ana Paula Canto de Lima e outras. RT, 2019) e "Manual de educação digital, cidadania e prevenção de crimes cibernéticos" (Org.: Higor Vinícius Nogueira Jorge. Juspodium; 2019).

Josemara Meireles: Advogada e Analista de Compliance Sênior, Consultora em Projetos de Privacidade e Proteção de Dados Pessoais, possui Formação Profissional em Privacidade de Dados (LGPD) pela TI Exames. Realizou o Curso oficial IAPP CDPO/BR Certified Data Protection Officer, membro da Associação Internacional de Profissionais de Privacidade (IAPP), membro do Grupo Mulheres na Privacidade, membro da Comissão de Compliance da OAB/BA, em fase de Conclusão do Certified Expert in Compliance pelo Instituto Auditoria, Risco e Compliance (ARC), Certificação em Analista Sênior de Compliance pelo Grupo JG Compliance. Especialista em Direito Tributário, Previdenciário, Responsabilidade e Contabilidade Fiscal pela Universidade Cândido Mendes.

Julyanne de Bulhões: Advogada, especialista em privacidade e proteção de dados pessoais, mestranda em propriedade intelectual e transferência de tecnologia para inovação pela Universidade Federal de Pernambuco. Facilitadora no grupo Mulheres na Privacidade e pesquisadora no grupo de estudos do Centro de Direito, Internet e Sociedade - CEDIS, do IDP Privacy Lab.

Karolyne Utomi: Advogada especialista e atuante desde 2014 em Privacidade, Proteção de Dados Pessoais, Direito Digital, Compliance e Contratos formada pelo Insper e FGV. Mestranda em Ciências Jurídicas. Empresária. Sócia fundadora da KR Advogados, escritório especialista em Privacidade e Proteção de Dados Pessoais. Sócia fundadora da Consultoria Kaosu (com foco em Educação Digital). Coordenadora do Comitê de Governança para o Desenvolvimento Social e Humano (RGB). Facilitadora do Grupo de Pesquisa - Racismo, Dados e Tecnologia (Mulheres na Privacidade) e do Grupo de Pesquisa - Dados Pessoais, Crianças e Adolescentes (Mulheres na Privacidade). Presidente da Comissão de Privacidade, Proteção de Dados Pessoais e LGPD da OAB – Guarulhos. Presidente da Coordenação e Comissão Temática de LGPD e Compliance da CamCMR. Parecerista, Palestrante e Ativista em diversas ações relacionadas a Privacidade, Proteção de Dados Pessoais, Compliance, Cidadania Digital e Desenvolvimento Social e Humano pautado nas novas tecnologias e advocacy.

Maria Beatriz Saboya: Graduada em Direito pela Universidade Federal de Pernambuco – UFPE. Advogada especializada em Direito Digital, Privacidade e Proteção de Dados. Membro certificada com as credenciais CIPP/E, CIPM e CDPO/BR, todas emitidas pela International Association of Privacy Professionals – IAPP, da qual é membro ativa e contribui com produção de artigos, participação em grupos de trabalho e grupos de discussão. Fundadora da rede Mulheres Na Privacidade. Coautora dos livros “Direito Exponencial: O Papel das Novas Tecnologias no Jurídico do Futuro” (Ed. Revista dos Tribunais) e “LGPD Aplicada” (Ed. Atlas e Gen Jurídico), este último tendo sido indicado na lista de Bibliografias Seleccionadas sobre Lei Geral de Proteção de Dados Pessoais (LGPD) do Superior Tribunal de Justiça (STJ), no ano de 2022.

Melisa Curilla Rapelli: Advogada, graduada pela Universidade São Judas Tadeu e Pós-Graduada em Direitos Humanos, Responsabilidade Social e Cidadania Global pela PUC/RS. Possui curso de extensão em Políticas Públicas pela PUC/SP e Sustentabilidade pela FGV. É membro do grupo Mulheres na Privacidade, atuando em

pesquisas e estudos voltados à vieses algorítmicos e Multiplicadora da rede Politize! de educação política e cidadania para todos.

Nathália Criscito: Advogada, graduada pela Universidade Presbiteriana Mackenzie. Possui curso de extensão em Privacidade e Proteção de Dados pela Data Privacy Brasil. É membra do Grupo Mulheres na Privacidade atuando em pesquisas e estudos relacionados a temas de privacidade, com foco especial em vieses algorítmicos e atuante em governança em privacidade e proteção de dados.