

DOCUMENTO

PRE TO



Contribuições do Aqualtune Lab
para o debate sobre regulação
de **Inteligência Artificial no Brasil**



Documento Preto I

Contribuições do Aqualtune Lab
para o debate sobre regulação de
Inteligência Artificial no Brasil



DOCUMENTO

PRETO

Contribuições do Aqualtune Lab
para o debate sobre regulação de
Inteligência Artificial no Brasil

AUTORIA

Arthur Almeida Meneses Barbosa
Celso Eduardo Lins de Oliveira
Paulo Rená da Silva Santarém
Natane da Silva Santos
Fernanda dos Santos Rodrigues Silva
Valdinei Freire da Silva

PROJETO GRÁFICO, CAPA, DIAGRAMAÇÃO E FINALIZAÇÃO

Isabel Cristina Coronel Xavier

GRUPO DE TRABALHO DE INTELIGÊNCIA ARTIFICIAL

Ana Carolina Sousa Silva
Arthur Almeida
Bianca Kremer
Cassia Isac
Celso Oliveira
Clarissa França
Ernane Costa
Fernanda Rodrigues
Natane Santos
Paulo Rená
Valdinei Freire



DIRETORIA

Ana Carolina Lima
Celso Oliveira
Clarissa França
Paulo Rená

SUGESTÃO DE CITAÇÃO

BARBOSA, Arthur Almeida Meneses;
OLIVEIRA, Celso Eduardo Lins de;
SANTARÉM, Paulo Rená da Silva;
SANTOS, Natane; SILVA, Fernanda
dos Santos Rodrigues; SILVA, Valdinei
Freire da. Documento Preto I. Brasil:
Aqualtune Lab, maio de 2022.

Diante da discussão existente no âmbito do Congresso Nacional sobre a regulação do desenvolvimento e uso da Inteligência Artificial no Brasil, o coletivo Aqualtune Lab se posiciona com os princípios a seguir.

O racismo é parte da construção histórica da sociedade brasileira e seus efeitos são cotidianos na população negra, manifestando-se das mais diversas formas diretas e indiretas sendo perpetuado na forma da organização das estruturas públicas e privadas¹.

Esta perpetuação só é possível pois, a cada ponto de suposto desenvolvimento, a cada avanço econômico, social ou tecnológico, o racismo se adapta, se transforma e é aplicado até o ponto de ser considerado normal criando a invisibilidade histórica e material do povo negro que estatisticamente corresponde a mais de 55% de nossa população.

O entendimento deste processo evidencia que o uso e o domínio de técnicas de inteligência artificial representa um enorme potencial para que este fenômeno social ocorra e se aprimore provocando segregação em várias formas da vivência da negritude, desde as comunidades tradicionais até os mais altos níveis de ascensão social.

¹ “Podemos dizer que o racismo é uma forma sistemática de discriminação que tem a raça como fundamento, e que se manifesta por meio de práticas conscientes ou inconscientes que culminam em desvantagens ou privilégios para indivíduos, a depender do grupo racial ao qual pertençam.” ALMEIDA, Sílvio Luiz de. *O que é racismo estrutural?* Belo Horizonte, MG: **Letramento**, 2018. p. 32.

Por exemplo, ressaltamos o conhecido exemplo da aplicação das tecnologias de reconhecimento facial que atingem de forma cruel a população preta e pobre, sendo a nova forma da teoria do suspeito natural. Nesse sentido, no documento “A Inteligência Artificial na era digital”, de 3 de maio de 2022, o Parlamento Europeu observa que, embora muitos receios sejam hipotéticos, são manifestos os efeitos nocivos de racismo, entre outras discriminações, decorrentes da adoção de decisões baseadas em IA sem as correspondentes salvaguardas.²

Tecnologias como esta devem ser banidas até que seu uso possa ser considerado seguro para toda a população independentemente da cor da sua pele, nível social, sexo ou idade. Neste sentido, já encontramos dificuldade da própria definição da inteligência artificial e de seus aspectos definidores que devem preservar as garantias dos arts. 4º, VIII, e 5º, XLII Constituição da República e 20, §§ 2º, 3º, III, da Lei nº 7.716/1989. O Brasil hoje está comprometido, no plano nacional e internacional, em enfrentar o racismo, a discriminação racial e formas correlatas de intolerância, nos termos da Convenção Interamericana firmada pela República

² “O Parlamento Europeu (...) 16. Observa que muitos dos receios associados à IA assentam em conceitos hipotéticos, como a IA geral, a superinteligência artificial e a singularidade, que poderiam, em teoria, conduzir a uma inteligência artificial capaz de superar a inteligência humana em muitos domínios; realça que existem dúvidas quanto à possibilidade de esta IA especulativa ser alcançada com as nossas tecnologias e leis científicas; considera, no entanto, que os riscos que atualmente se colocam relativamente às decisões tomadas com base na IA têm de ser tratados pelos legisladores, uma vez que é manifesto que os efeitos nocivos, como a discriminação racial e sexual, já são atribuíveis a casos específicos em que a IA foi posta em prática sem salvaguardas” (destaque nosso). Ver UNIÃO EUROPEIA. *A Inteligência Artificial na era digital. Resolução do Parlamento Europeu*, de 3 de maio de 2022, sobre a inteligência artificial na era digital (2020/2266(INI)). Estrasburgo (FRA): Parlamento Europeu, 3 mai. 2022. https://www.europarl.europa.eu/doceo/document/TA-9-2022-0140_PT.html.

Federativa do Brasil, na Guatemala, e promulgada em 5 de junho de 2013 por meio do Decreto nº 10.932, de 10 de janeiro de 2022.

Desta norma, destacam-se especificamente os arts. 1, 3, 4, e 7, sobre as definições de discriminação racial, discriminação racial indireta, racismo, e intolerância; a proteção de direitos humanos individuais e coletivos; os deveres dos Estados; e o compromisso em adotar legislação antirracista.

Especialmente pertinente às aplicações de inteligência artificial, o conceito de *discriminação racial indireta* expõe dispositivos, práticas ou critérios aparentemente neutros que acarretam desvantagem particular. Isto ocorre com base em raça, cor, ascendência ou origem nacional ou étnica, para pessoas pertencentes a um grupo específico, ou as prejudicam, sem um objetivo ou justificativa razoável e legítima à luz do Direito Internacional dos Direitos Humanos.

Assim, é necessário afirmar em lei o dever de que os sistemas de inteligência artificial serem antirracistas: ativamente contrárias à produção de desigualdades raciais mediante a adoção de qualquer vínculo causal entre as características fenotípicas ou genotípicas de pessoas físicas ou grupos e seus traços intelectuais, culturais, comportamentais e de personalidade.

Tampouco se pode admitir, na definição de princípios legais, o uso de termos tão amenos como “*busca*”, “*recomendação*” e “*mitigação*”, em detrimento da assunção de uma incumbência e engajamento, em um efetivo dever jurídico expresso em palavras como “*realização*”, “*obrigação*” e “*eliminação*”.

Uma norma legal abstrata não tem efetiva capacidade de pronta aplicação direta no dia a dia do trabalho tecnológico de quem programa, revisa ou aprimora códigos na área de inteligência artificial. Não obstante, as previsões principiológicas do PL 21/2020 não podem ser aceitas sem um compromisso expressamente assumido com o antirracismo como um critério legal de validade para atividades de fomento, desenvolvimento e uso de inteligência artificial no Brasil.

Nesse sentido, destaca-se também a importância da existência de instrumentos de ação preventiva, como o relatório de impacto de inteligência artificial, a fim de que seja possível mensurar os possíveis riscos envolvidos no desenvolvimento e uso de novas tecnologias. Isso poderá ajudar a identificar previamente danos potenciais à sociedade, em especial, no que tange a minorias, como também a população negra, e tem o condão de trazer maior responsabilidade para os agentes que atuam nesse processo, que deverão participar ativamente para sanar eventuais problemas encontrados.

Tanto a União Europeia quanto o Canadá já possuem iniciativas nessa direção, o que ressalta a importância da adoção de medidas preventivas. Inclusive, a proposta de regulamento da inteligência artificial da UE traz situações específicas em que determinadas tecnologias são proibidas, como aquelas que possibilitam a classificação social de pessoas para uso pelo Poder Público, por exemplo, em razão do elevado risco de produzir resultados discriminatórios.³

3 COMISSÃO EUROPEIA. *Avaliação de impacto da IA e código de conduta*. Disponível em: <https://futurium.ec.europa.eu/en/european-ai-alliance/best-practices/ai-impact-assessment-code-conduct>. Acesso em 27 abr. 2022.

No Brasil, a campanha pelo banimento do reconhecimento facial tem sido uma das pautas mais latentes a esse respeito, voltada para luta antirracista na tecnologia e contra vieses discriminatórios. De fato, somente em 2019, mais de 90% dos presos por reconhecimento facial no Brasil eram negros.⁴ No entanto, além de haver estudos que já apontam que sistemas que utilizam essa tecnologia possuem maior índice de falibilidade sobre rostos de pessoas negras,⁵ a conhecida seletividade do sistema penal no país faz com que essa ferramenta seja potencialmente violadora de direitos fundamentais dessa parte da população.

Assim, faz-se importante pensar também em que tipo de sistemas de inteligência artificial podem ser efetivamente relevantes para o desenvolvimento e evolução social e quais, pelos seus riscos e problemas, possuem mais desvantagens do que benefícios.

No cenário mundial, há distintos países concentrados na concepção de legislações e diretrizes que visem o estabelecimento e defesa do direito do titular de dados pessoais para que este não esteja sujeito a decisões pautadas exclusivamente em processos

4 NUNES, Pablo. *Exclusivo: levantamento revela que 90,5% dos presos por monitoramento facial no Brasil são negros*. **The Intercept**, 21 novembro 2019. Disponível em: <https://theintercept.com/2019/11/21/presos-monitoramento-facial-brasil-negros/>. Acesso em: 27 abr. 2022.

5 BUOLAMWINI, Joy; BEGRU, Timnit. *Gender shades: intersectional accuracy disparities in commercial gender classification*. **Proceedings of machine learning research**, v. 81, 2018, pp. 1-15.

automatizados⁶. Esta preocupação ocorre porque, sistemas de IA são suscetíveis a falhas que acarretam prejuízos sobre direitos fundamentais e riscos discriminatórios inadmissíveis. Além disso, os valores éticos e democráticos devem ser levados em consideração, especialmente, quando utiliza-se sistemas de IA como apoio ou substituição de decisões humanas verificando se há a necessidade de intervenção humana conforme a análise do contexto histórico, socioeconômico e político.

A preponderância da obrigatoriedade de revisão humana de decisões automatizadas pressupõe a garantia efetiva dos direitos: (i) autodeterminação informativa; (ii) não discriminação e transparência; (iii) direito de informação sobre critérios e parâmetros de decisões, revisão, explicação e oposição as decisões automatizadas⁷.

Entre as disposições normativas com a finalidade de proteger os interesses, direitos e garantias dos usuários cujos dados são analisados por distintos modelos de inteligências artificiais, ou seja, regulamentações que defendem direitos relacionados à participação humana em decisões automatizadas: Regulamento

6 DONEDA, Danilo; WIMMER, Miriam. “Falhas de IA” e a intervenção humana em decisões automatizadas: parâmetros para a legitimação pela humanização. *Revista Direito Público*. V.18 n. 100 (2021). Disponível em: <https://www.portaldeperiodicos.idp.edu.br/direitopublico/article/view/6119>. Acesso em: 25 abr. de 2022

7 SANTOS, Natane da Silva. Monografia Jurídica *Lei Geral de Proteção de Dados e os possíveis impactos da não obrigatoriedade de revisão humana de decisões automatizadas*. 2021.2. Faculdade Nacional de Direito da Universidade Federal do Estado do Rio de Janeiro - FND/UFRJ.

Geral de Proteção de Dados da Europa (*GDPR General Data Protection Regulation*)⁸; Convenção para a Proteção das Pessoas relativamente ao Tratamento Automatizado de Dados de Caráter Pessoal (conhecida como Convenção 108+); Recomendação do Conselho de Inteligência Artificial da Organização para a Cooperação e Desenvolvimento Econômico (OCDE); Proposta de Regulamento de IA da União Europeia e Recomendação da UNESCO sobre Ética da Inteligência Artificial. Convém enfatizar a relevância do respeito a Convenção Interamericana contra o Racismo, a Discriminação Racial e Formas Correlatas de Intolerância, visto que há uma obrigatoriedade da promoção e defesa de direitos relativos à proteção contra formas de racismo, logo, isto impõe deveres aos agentes de sistemas de inteligência artificial, dentre eles, pode-se apontar, a identificação e eliminação de vieses racistas.

Então, é salutar uma cautela na edição de leis e regulamentos sobre a aplicação de inteligência artificial em diferentes nações. Especialmente, no que tange à elaboração e execução ética de sistemas de IA oportunizando a centralidade na pessoa humana, regime de responsabilidade adequado, crescimento socioeconômico global inclusivo e responsável, bem como, a observância de alguns princípios relevantes para promoção efetiva dos direitos do titular de dados pessoais, dentre eles: equidade, proteção contra o racismo, transparência, explicabilidade e *accountability*.

| **8** Art. 22º e Considerando nº 71 RGPD (GDPR).

Vale ressaltar que a intervenção humana é fundamental⁹ para proporcionar maior segurança jurídica, através do estabelecimento de normas, boas práticas, fiscalização e sanções, de fato privilegiando o respeito ao tratamento igualitário e a diversidade nos setores público e privado.

Assim, manifestamos os seguintes princípios que devem ser observados em qualquer proposta de marco legal sobre a IA:

1) Qualquer legislação deve ser explicitamente antirracista, evitando termos ambíguos ;

2) Qualquer legislação deve ter evidente o princípio da transparência, efetivando questões de culpabilidade e consequências, evitando portanto a consecução de erro sem a devida medida reparadora pelo responsável.

3) Deve haver uma agência independente e autônoma que coordene a fiscalização descentralizada, na qual haja representação diversa da sociedade, incluindo obrigatoriamente os grupos sujeitos a riscos;

4) Deve haver classificação de riscos com critérios transparentes, auditáveis e que atendam ao item 1;

9 MEDON, Filipe. *Decisões automatizadas: o necessário diálogo entre a Inteligência Artificial e a proteção de dados pessoais para a tutela de direitos humanos*. Livro *O Direito Civil na era da Inteligência Artificial* Editora Revista dos Tribunais. p. 370.

5) As ações de IA consideradas de risco alto devem ser banidas até que a tecnologia possa apresentar resultados que eliminem tal margem de risco.

6) As tecnologias de IA de risco médio devem em seu processo final passar por revisão humana, com admissão de responsabilidade pelo agente de tal revisão;

7) Sanções e responsabilização no caso de descumprimento dos princípios elencados.



Quem somos:

AqaltuneLab é um coletivo jurídico com suporte multidisciplinar, pautado no estudo e elaboração de propostas que comportam a análise das inter-relações entre Direito, Tecnologia e Raça. Nosso grupo é composto por pessoas de diversas regiões do país, e temos como objetivo racializar discussões em temas como uso de tecnologias no sistema jurídico a exemplo – mas não nos limitando – das vigilâncias pública e privada, políticas de proteção de dados, identificação biométrica, segurança na internet, aplicativos móveis e mídias sociais. Buscamos compreender através de um olhar jurídico e multidisciplinar qualificados e orientados na comunicação, TICs e áreas correlatas, compreender as intersecções entre tecnologia e raça, descortinando as armadilhas do racismo como estruturantes das relações de saber e poder no Brasil.





www.aqualtunelab.com.br



contato@aqualtunelab.com.br



Instagram: @aqualtunelab



Twitter: @aqualtunelab



Facebook: @aqualtunelab



Linkedin: Aqualtune Lab

