

Contribuição da Coding Rights à Consulta Pública da CJSUBIA

joana@codingrights.org

qui 30/06/2022 21:06

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

Cc: Vanessa <vanessa@codingrights.org>;

 1 anexo

Contribuicao_CodingRights.pdf;

[Geralmente, você não obtém emails de joana@codingrights.org. Saiba por que isso é importante em <https://aka.ms/LearnAboutSenderIdentification>]

À Sua Excelência Ministro Ricardo Villas Bôas Cueva,

Exmo. Sr. Presidente da Comissão de Juristas do Senado Federal responsável por debater e elaborar minuta de substitutivo aos projetos de lei sobre inteligência artificial, em atendimento aos termos da convocatória da referida Consulta Pública e no intuito de contribuir com as ações desta Comissão, mais especificamente com os eixos de trabalho definidos como "2. Impactos da Inteligência Artificial" e "3. Direitos e Deveres", a Coding Rights vem, respeitosamente, compartilhar no documento em anexo os acúmulos obtidos por esta organização no decorrer de seus trabalhos e estudos sobre Inteligência Artificial.

Desde já, agradecemos a atenção despendida. E ficamos à disposição para eventuais consultas que visem melhor informar a Comissão, bem como explicitar melhor qualquer ponto que careça de maior nitidez.

Quando possível, solicito, encarecidamente, confirmação do recebimento do anexo.

Atenciosamente,

Joana Varon

--

Joana Varon
Diretora Executiva
Coding Rights

Technology and Human Rights Fellow
Carr Center for Human Rights Policy
Harvard Kennedy School

Affiliate
Berkman Klein Center for Internet and Society
At Harvard University

@joana_varon
@codingrights
codingrights.org

encrypting is caring, so:

Fingerprint A240 141B 5D29 7C63 5D58 CE97 1FC4 586A B078 CD8A

São Paulo, 30 de junho de 2022

À Sua Excelência Ministro Ricardo Villas Bôas Cueva

Exmo. Sr. Presidente da Comissão de Juristas do Senado Federal responsável por debater e elaborar minuta de substitutivo aos projetos de lei sobre inteligência artificial

Ref. Contribuição da Coding Rights focada nos eixos temáticos: II - Impactos da inteligência artificial e III - Direitos e Deveres, de acordo com o plano de trabalho da Consulta Pública da Comissão de Juristas do Senado Federal - CJUSBIA

Excelentíssimo Senhor Presidente,

A Coding Rights, pessoa jurídica devidamente inscrita no CNPJ/MF sob o nº 11.758.057/0001-70, é uma organização criada em 2015, um *think tank* que tem como parte de suas missões realizar pesquisa aplicada à políticas públicas na área de tecnologia e direitos humanos, buscando influenciar implementações de tecnologias rumo à equidade de gênero e suas interseccionalidades de raça, classe, sexualidade e território.

Diante disso, em atendimento aos termos da convocatória da referida Consulta Pública e no intuito de contribuir com as ações desta Comissão e especificamente com os eixos de trabalho definidos como "2. Impactos da Inteligência Artificial" e "3. Direitos e Deveres", esta organização vem, respeitosamente, compartilhar acúmulos obtidos no decorrer de seus trabalhos.

Desde 2018, a Coding Rights vem desempenhando um extenso trabalho de pesquisa, em parceria com outras especialistas em inteligência artificial e direitos humanos, sobre os impactos da inteligência artificial na América Latina e no Brasil. A pesquisa segue em execução e parte de seus resultados estão consolidados no website "Que inteligência?" (<https://queinteligencia.org>), disponível também em inglês sob o nome "Not my AI"

Coding Rights

<https://codingrights.org>
contato@codingrights.org

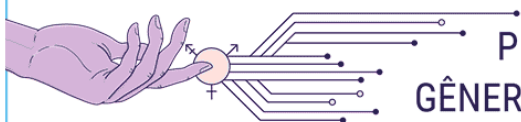
(<https://notmy.ai/>), parte desta contribuição que apresentamos aqui foi escrita com base nos resultados deste estudo. O intuito é difundir os acúmulos já existentes acerca dos impactos do aprendizado de máquina e estimular os debates entre organizações engajadas na promoção da justiça social, particularmente em temas que visam assegurar e avançar em direitos de equidade de gênero e suas interseccionalidade de raça, classe e sexualidade. O que se espera é que novas tecnologias não sejam implementadas no sentido de automatizar desigualdades históricas.

Infelizmente, isso já está ocorrendo. Governos de toda a América Latina estão implementando ou testando uma ampla variedade de sistemas de Inteligência Artificial na prestação de serviços públicos. Algumas vezes com o apoio de empresas dos EUA que usam a região como um laboratório de ideias e de aplicação de tecnologias que não ousam testar em seus países de origem.

Para ter uma visão geral de tendências na região, no projeto *Queinteligencia.org* realizamos pesquisa documental e aplicamos um [questionário](#), distribuído em redes de direitos digitais na região, para mapear projetos nos quais sistemas algorítmicos de tomada de decisões estão sendo implantados por governos e que possuam prováveis implicações negativas em matéria de igualdade de gênero e todas as suas interseccionalidades. Como resultado, até abril de 2021, mapeamos 24 casos no Chile, Brasil, Argentina, Colômbia e Uruguai, que pudemos classificar em cinco categorias: Educação; Sistema Judicial; Policiamento; Saúde Pública e Benefícios Sociais. Vários deles estavam em estágio inicial de implantação ou são desenvolvidos como pilotos.

O mapa abaixo ilustra projetos levantado. Embora este seja um levantamento contínuo e não exaustivo para mapear principais áreas de implementação e ilustrar narrativas vigentes em torno desses projetos, ele já revela tendências preocupantes que deveriam receber a atenção.

PROJETOS DE I.A. DO SETOR PÚBLICO NA AMÉRICA LATINA



PRECONCEITO E DISCRIMINAÇÃO DE
GÊNERO E SUAS INTERSECCIONALIDADES

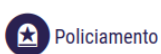
PROJETOS POR ÁREA DE APLICAÇÃO:



Benefícios sociais



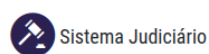
Educação



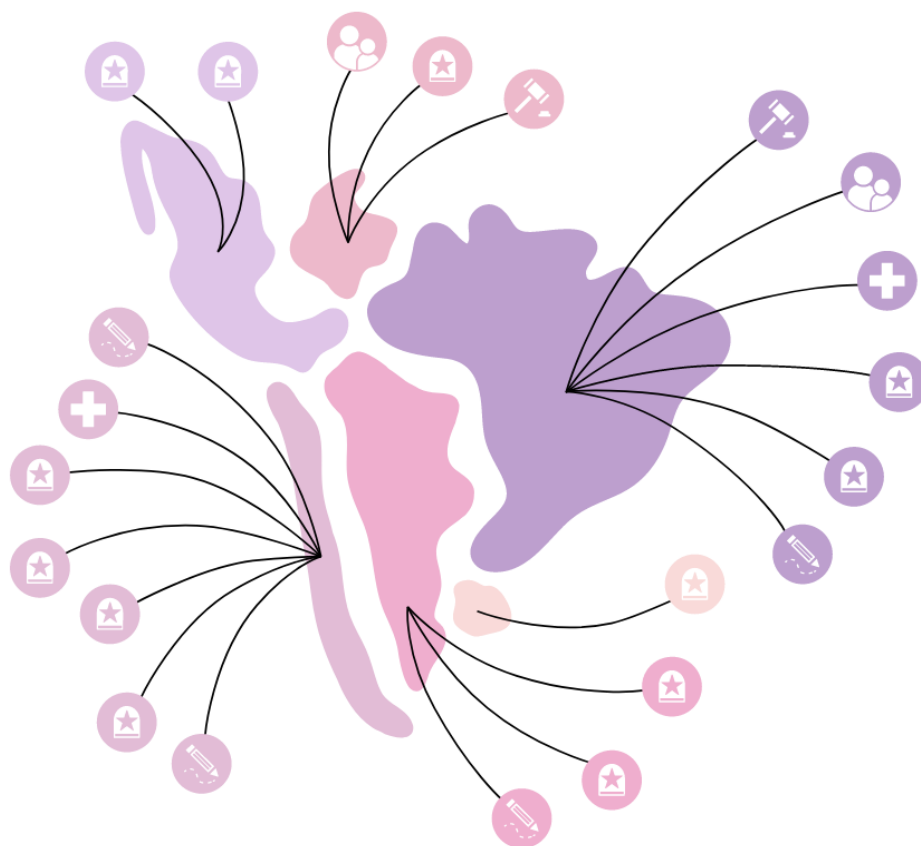
Policiamento



Saúde pública



Sistema Judiciário



Mapeamento feito pela plataforma “Que Inteligência?”, disponível em
<https://notmy.ai/pt/mapeamento-de-projetos/>

Coding Rights Projetos

<https://codingrights.org>
contact@codingrights.org

Ademais, recentemente, utilizando-nos do sistema presente na Lei de Acesso à Informação (Lei nº 12.527 de 2011), realizamos outro levantamento mais extenso sobre o uso de inteligência artificial pelo setor público federal brasileiro. Para tanto, foram efetuados pedidos de acesso à informação a 45 (quarenta e cinco) entes públicos com as seguintes perguntas:

"Esta solicitação visa compreender se este órgão utiliza ou faz testagem, ainda que em fase piloto, de sistemas de inteligência artificial e/ou aprendizado de máquina (machine learning) para desenvolvimento de seus trabalhos e funções, incluindo a implementação das políticas públicas.

Se sim, favor:

- a. Listar o nome de todos os projetos que estão utilizando ou testando inteligência artificial e/ou aprendizado de máquina (machine learning).
- b. Explicar para que fins e objetivos são utilizados e em quais programas ou áreas de atuação deste órgão."

Os resultados iniciais, consolidados em Outubro de 2021, indicam que a maioria dos entes do setor público brasileiro federal perguntados afirmaram utilizar algum tipo de sistema de Inteligência Artificial para desempenho de suas funções, incluindo a implementação de políticas públicas. Ou seja, são 23 entes públicos que autodeclararam fazer uso de sistemas de inteligência artificial sem a presença de um panorama regulatório em vigor, e pouco ou nada se sabe se existem análises sobre a necessidade real desses sistemas, se representam alguma eficiência na sua implementação, nem se existem análises de impacto, risco e dano realizadas antes de que esses sistemas sejam implementados, além de ser comum que as partes envolvidas se eximam de eventuais responsabilidades. O quadro a seguir ilustra o resultado deste mapeamento:

Uso de Inteligência Artificial pelo Setor Público no Brasil | Ref. Out.2021

Este órgão utiliza ou faz testagem, ainda que em fase piloto, de sistemas de inteligência artificial e/ou aprendizado de máquina (machine learning) para desenvolvimento de seus trabalhos e funções, incluindo a implementação das políticas públicas?

Legenda:

Sim

Não

Indefinido

Advocacia Geral da União AGU	Banco Central do Brasil BCB	Banco do Brasil BB	Banco do Nordeste do Brasil BNB
Banco Nacional de Desenvolvimento Econômico e Social BNDES	Caixa Econômica Federal CEF	Conselho Administrativo de Defesa Econômica CADE	Controladoria Geral da União CGU
Coordenação de Aperfeiçoamento de Pessoal de Nível Superior CAPES	Departamento de Polícia Federal PF	Empresa Brasil de Comunicação EBC	Empresa de Tecnologia e Informações da Previdência DATAPREV
Fundação Instituto Brasileiro de Geografia e Estatística IBGE	Fundação Oswaldo Cruz FIOCRUZ	Instituto Brasileiro do Meio Ambiente e dos Recursos Naturais Renováveis IBAMA	Instituto Nacional de Colonização e Reforma Agrária INCRA
Instituto Nacional de Estudos e Pesquisas Educacionais Anísio Teixeira INEP	Instituto Nacional de Seguro Social INSS	Ministério da Mulher, da Família e dos Direitos Humanos	Ministério da Defesa
Ministério da Economia	SERPO Serviço Federal de Processamento de Dados	Receita Federal do Brasil RFB	
Agência Nacional de Águas ANA	Agência Nacional de Saúde Suplementar ANS	Autoridade de Proteção de Dados ANPD	Banco da Amazônia BASA
Conselho Nacional de Desenvolvimento Científico e Tecnológico CNPQ	Gabinete de Segurança Institucional da Presidência da República GSI-PR	Financiadora de Estudos e Projetos FINEP	Fundação Nacional do Índio FUNAI
Ministério da Ciência, Tecnologia e Inovações	Ministério da Agricultura, Pecuária e Abastecimento	Ministério da Cidadania	Ministérios das Relações Exteriores
Ministério da Educação	Ministério da Justiça e Segurança Pública	Ministro de Minas e Energia	Ministério da Saúde
Ministério das Comunicações	Ministério do Turismo	Ministério de Desenvolvimento Regional	Ministério do Meio Ambiente
Agência Nacional de Energia Elétrica ANEEL	Ministério da Infraestrutura		

Coding Rights Projetos

<https://codingrights.org>
contact@codingrights.org

Estes dados preliminares, por si, já indicam que o uso de algum tipo de inteligência artificial, ainda que estágios piloto, pelo Estado brasileiro é uma realidade. Acreditamos que seja necessário tomar um passo atrás, e repensar a ideologia dominante do “move fast and break things”, que foi imposta pelas empresas do Vale do Silício e tem tido consequências nefastas tanto para a justiça sócio-ambiental, como para democracias de todo o mundo. Andar com cautela e prever riscos para evitar danos tem se mostrado uma visão mais coerente com a defesa de nossos direitos fundamentais. Ou seja, enquanto não se provar que não está causando danos, principalmente se os alvos dos sistemas propostos são comunidades marginalizadas, muito provavelmente danos estão sendo causados e desigualdades históricas estão sendo automatizadas.

Sistemas de I.A. baseiam-se em modelos que são representações, universalizações e simplificações abstratas de realidades complexas nos quais muitas informações são deixadas de fora por decisão de seus criadores. Como aponta Cathy O’Neil em seu livro [“Weapons of Math Destruction” \[Armas de Destruição Matemática, em tradução livre\]: “\[M\]odelos, apesar de sua reputação de imparcialidade, refletem objetivos e ideologia. \[...\] Nossos próprios valores e desejos influenciam nossas escolhas, desde os dados que escolhemos coletar até as perguntas que fazemos. Modelos são opiniões incorporadas na matemática”](#). (O’Neil, 2016, tradução livre).

Consequentemente, algoritmos são criações humanas sujeitas a falhas. Os seres humanos estão sempre presentes na construção de sistemas automatizados de tomada de decisão: são eles que determinam os objetivos e usos dos sistemas, que definem quais dados devem ser coletados para tais objetivos e usos, que coletam esses dados, que decidem como capacitar as pessoas que usam esses sistemas, avaliam o desempenho do software e, em última análise, agem segundo as decisões e avaliações feitas pelos sistemas.

Mais especificamente, como afirma Tendayi Achiume, Relator Especial sobre formas contemporâneas de racismo, discriminação racial, xenofobia e intolerância relacionadas, no

relatório [“Racial discrimination and emerging digital technologies”](#) [Discriminação racial e tecnologias digitais emergentes, em tradução livre], os bancos de dados usados nesses sistemas são resultado do design humano e podem ser tendenciosos de várias formas, potencialmente levando – intencionalmente ou não – à discriminação ou à exclusão de certas populações, especialmente minorias em questões de identidade de raça, etnia, religião e gênero (Tendayi, 2020).

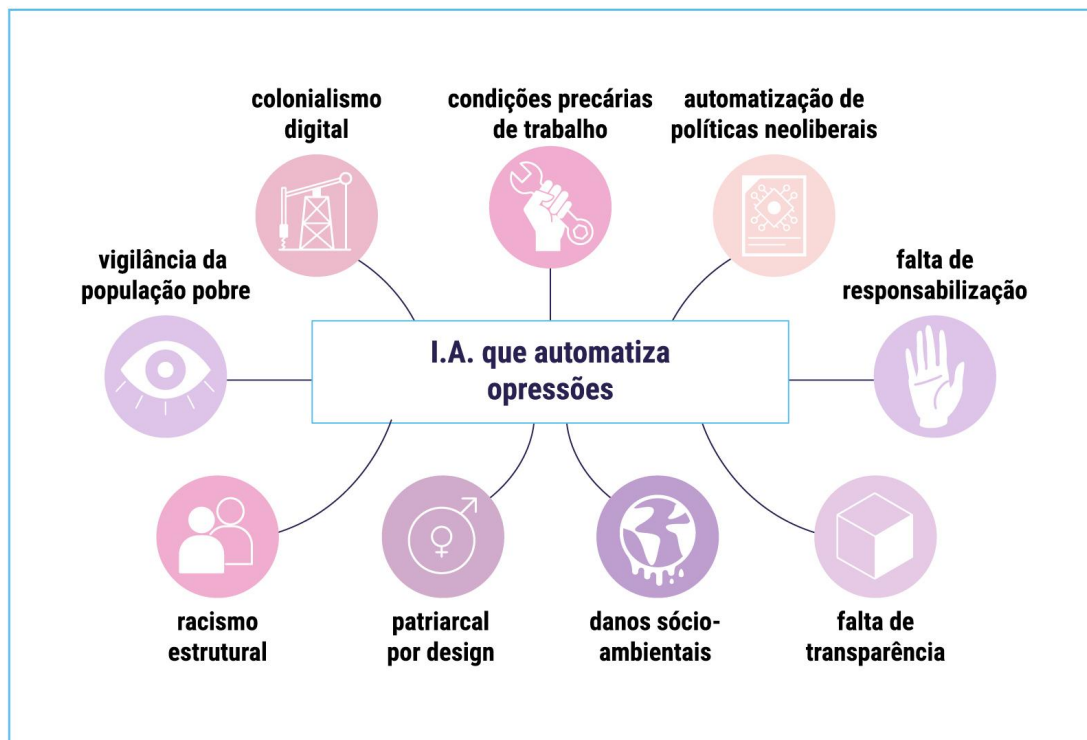
Diante desses problemas, deve-se reconhecer que parte da comunidade da tecnologia tem feito várias tentativas para definir matematicamente “justiça” e, assim, atender a um padrão demonstrável sobre o tema. Da mesma forma, várias organizações, privadas e públicas, têm empreendido esforços para estabelecer padrões éticos para a I.A. A visualização de dados [“Principled Artificial Intelligence”](#) [Inteligência Artificial com Princípios, em tradução livre] (Berkman Klein, 2020) mostra a variedade de quadros éticos e baseados em direitos humanos que surgiram em diferentes setores a partir de 2016 com o objetivo de orientar o desenvolvimento e o uso de sistemas de I.A. O estudo revela [“um consenso crescente em torno de oito tendências temáticas principais: privacidade, responsabilização, proteção e segurança, transparência e explicabilidade, justiça e não discriminação, controle humano da tecnologia, responsabilidade profissional e promoção dos valores humanos”](#) (tradução livre). No entanto, como podemos ver dessa lista, nenhum desses consensos é impulsionado por princípios de justiça social. Em vez de perguntar como desenvolver e implantar um sistema de I.A., não deveríamos antes perguntar “por que o construir?”, “é realmente necessário?”, “a pedido de quem?”, “quem se beneficia dele?”, “quem é prejudicado?” com a implantação de um determinado sistema de I.A.? Deveria tal sistema sequer ser desenvolvido e implantado?

Com base em uma extensão revisão da bibliografia existente e em análise baseada em casos mapeados na América Latina, o que se segue é uma *framework* analítica que tem sido testada em uma série de reuniões e oficinas realizadas junto a organizações preocupadas com a justiça social e equidade de gênero em suas interseccionalidades.

Possíveis danos causados pela implantação de tomada de decisão algorítmica em políticas públicas

A seguir, apresentamos um resumo das críticas que têm sido feitas a sistemas de I.A. implantados pelo setor público, que possuem pontos de intersecção e se retroalimentam. Com base em uma ampla revisão bibliográfica e em uma análise fundamentada em casos mapeados na América Latina, o que se segue é uma framework analítica que tem sido testada em uma série de reuniões e oficinas realizadas com organizações preocupada com a justiça social e equidade de gênero em suas interseccionalidades. O quadro analítico está em constante evolução (*work-in-progress*) e visa ultrapassar os discursos rasos de uma I.A. ética ou centrada em pessoas (que pessoas?) e busque ser uma estrutura holística que considere as relações de poder a fim de questionar a ideia de implantar sistemas de I.A. em vários campos do setor público. Esperamos que esse framework analítico também sirva de inspiração para a confecção do projeto de lei em questão, particularmente ao se pensar impactos da inteligência artificial, direitos e deveres. Este documento também termina com dois exemplos práticos: um exemplos de aplicação do quadro analítico em um caso de I.A. de predição de gravidez na adolescência e outro exemplo regulação que proíbe determinados usos da tecnologia de reconhecimento facial.

FRAMEWORK ANALÍTICO



Esquema desenvolvido por Joana Varon e Paz Peña e Clarote para a plataforma queinteligencia.org

A seguir, explicamos cada sessão do esquema analítico:

A. Vigilância dos pobres: transformando pobreza e vulnerabilidade em dados legíveis automaticamente

O ex-Relator das Nações Unidas sobre Pobreza Extrema e Direitos Humanos, Philip Alston, criticou o fenômeno segundo o qual [“os sistemas de proteção e assistência social são cada vez mais impulsionados por dados e tecnologias digitais que são usados para automatizar, prever, identificar, vigiar, detectar, mirar e punir.”](#) (A/74/48037 2019, tradução livre). Essas fontes de dados detalhados permitem às autoridades inferir movimentos, atividades e comportamento das pessoas, não sem ter implicações éticas, políticas e práticas a respeito de como os setores público e privado veem e tratam as pessoas. De acordo com

Coding Rights Projetos

<https://codingrights.org>
contact@codingrights.org

Linnet Tylor, em seu artigo [“What is data justice?” \(TYLOR, 2017\)](#) [O que é justiça de dados?, em tradução livre], isso é ainda mais desafiador nos casos de estratos da população de baixa renda, uma vez que a capacidade das autoridades de coletar dados estatísticos precisos sobre eles, que no passado era limitada, agora é o alvo de sistemas de classificação por regressão que perfilam, julgam, punem e vigiam.

A maioria desses programas se aproveita da [tradição de vigilância estatal sobre populações vulneráveis \(Eubanks, 2018, tradução livre\)](#), transforma sua existência em dados e agora usa algoritmos para definir a concessão de benefícios sociais pelos Governos. Analisando o caso dos EUA, Eubanks mostra como o uso de sistemas de I.A. está sujeito a uma longa tradição de instituições que gerem a pobreza e que buscam, por meio dessas inovações, se adaptar e manter seu desejo de conter, monitorar e punir os pobres. Trata-se, portanto, [de transformar pobreza e vulnerabilidade em dados legíveis automaticamente, com consequências reais na vida e na subsistência dos cidadãos envolvidos. \(Masiero & Das, 2019\). No mesmo sentido,](#) Cathy O’Neil (2016), analisando os usos da I.A. nos Estados Unidos (EUA), afirma que muitos sistemas de I.A. “tendem a punir os pobres”, o que significa que é cada vez mais comum que pessoas ricas sejam beneficiadas por interações pessoais, enquanto os dados dos pobres são tratados por máquinas que tomam decisões a respeito de seus direitos.

Isso se torna ainda mais relevante quando consideramos que o sistema de classes sociais possui um poderoso componente de gênero. É comum que as políticas públicas mencionem a “feminização da pobreza”. De fato, na IV Conferência das Nações Unidas sobre a Mulher, realizada em Pequim, em 1995, concluiu-se que 70% das pessoas pobres do mundo eram mulheres. Os motivos pelos quais a pobreza afeta mais as mulheres não têm a ver com motivos biológicos, mas com as estruturas de desigualdade social que dificultam a superação da pobreza pelas mulheres, tais como o acesso à educação e ao emprego (Aguilar, 2011).

B. Racismo estrutural integrado no algoritmo

Para o Relator Especial da ONU, E. Tandayi (2020), as tecnologias digitais emergentes também devem ser entendidas como capazes de criar e manter a exclusão racial e étnica em termos sistêmicos ou estruturais. Isso é também o que pesquisadoras de tecnologia dedicadas a estudo de raça e I.A. nos Estados Unidos, tais como [Ruha Benjamin](#), [Joy Buolamwini](#), [Timnit Gebru](#) e [Safiya Noble](#) destacam em seus estudos de caso. No contexto brasileira, [Nina da Hora](#), [Tarcísio Silva](#), [Pablo Nunes](#), Bianca Kremer, também apontam o que Tarcísio Silva chamou de racismo algorítmico ao investigarem desde tecnologias de reconhecimento facial a algoritmos de ferramentas de busca tamb. Ruha Benjamin (2019) discute especificamente como o uso de novas tecnologias reflete e reproduz as injustiças raciais existentes na sociedade estadunidense, ainda que elas sejam promovidas e percebidas como mais objetivas ou progressistas do que sistemas discriminatórios de épocas passadas. Nesse sentido, para a autora, quando a I.A. busca determinar o quanto pessoas de todas as classes merecem oportunidades, os desenvolvedores dessas tecnologias constroem um sistema de castas digital estruturado sobre a discriminação racial existente.

A partir do próprio desenvolvimento da tecnologia, em sua pesquisa, Noble (2018) demonstra como ferramentas de busca comerciais como o Google não apenas conciliam, como também são afetadas por uma série de imperativos comerciais que, por sua vez, são sustentados por políticas econômicas e de informação que acabam por endossar a mercantilização das identidades femininas. Nesse caso, a autora expõe o tema ao analisar uma série de pesquisas no Google em que mulheres negras acabam sendo sexualizadas pelas informações contextuais que o mecanismo de pesquisa exhibe (por exemplo, vinculando-as a mulheres selvagens e sexualizadas).

Outro estudo notável é o de Buolamwini & Gebru (2018), que analisaram três sistemas comerciais de reconhecimento facial que permitem classificar rostos por gênero. Elas descobriram que os sistemas apresentam taxas de erro mais altas para mulheres de pele mais

escura do que para qualquer outro grupo, com as taxas de erro mais baixas sendo as referentes aos homens de pele clara. As autoras atribuem esses preconceitos de raça e gênero à composição dos conjuntos de dados usados para treinar esses sistemas, esmagadoramente compostos de indivíduos de pele mais clara com aparência masculina.

C. Patriarcal desde o Projeto: sexismo, heteronormatividade compulsória e binaridade de gênero

Muitos [sistemas de I.A. funcionam classificando as pessoas segundo uma visão binária de gênero](#), bem como reforçando antiquados estereótipos de gênero e orientação sexual. Um [estudo recente de coautoria de Shakir Mohamed, cientista sênior da DeepMind](#), expõe como a discussão sobre justiça algorítmica ignorou a orientação sexual e a identidade de gênero, com impactos concretos na “censura, linguagem, segurança online, saúde e emprego”, levando à discriminação e à exclusão de pessoas LGBTQ+.

O gênero foi analisado de várias maneiras na Inteligência Artificial. West, Whittaker e Crawford (2019) argumentam que a crise de diversidade na indústria e as questões de sistemas tendenciosos de I.A. (especialmente no que se refere a raça e gênero) são faces inter-relacionadas do mesmo problema. No passado, as pesquisas costumavam examinar essas questões isoladamente, mas evidências crescentes mostram que elas estão intimamente interligadas. No entanto, alertam as autoras, apesar de todas as evidências sobre a necessidade de diversidade nos campos da tecnologia, tanto na academia quanto na indústria, esses indicadores estagnaram.

Inspiradas por Buolamwini & Gebru (2018), [Silva e Varon \(2021\) pesquisaram como tecnologias de reconhecimento facial afetam pessoas trans](#) e concluíram que, embora os principais órgãos públicos no Brasil já utilizem esses tipos de tecnologia para confirmar a identidade para se ter acesso a serviços públicos, há pouca transparência sobre a precisão dos mesmos (apontando falsos positivos ou falsos negativos), bem como sobre privacidade e

proteção de dados frente às práticas de compartilhamento de dados entre órgãos da administração pública e mesmo entre entidades privadas.

No caso da Venezuela, em meio a uma longa crise humanitária, o Estado implementou sistemas biométricos para controlar a aquisição de bens de primeira necessidade, resultando em várias denúncias de discriminação contra estrangeiros e pessoas trans. Segundo Díaz Hernández (2021), a legislação para proteger pessoas trans é praticamente inexistente. Não lhes é concedido o reconhecimento de sua identidade, o que faz com que essa tecnologia ressignifique o valor de seus corpos “e os transforme em corpos inválidos, que ficam, portanto, à margem do sistema e à margem da sociedade” (p.12, tradução livre).

No caso dos programas de gestão da pobreza por meio de big data e sistemas de Inteligência Artificial, é fundamental observar como as mulheres pobres estão particularmente sujeitas à vigilância estatal e como isso leva à reprodução das desigualdades econômicas e de gênero (Castro & López, 2021).

D. Colonialismo Digital

Autores como [Couldry e Mejias \(2018\)](#) e [Shoshana Zuboff \(2019\)](#) analisam o estágio atual do capitalismo, em que a produção e a extração de dados pessoais naturalizam a apropriação colonial da vida em geral. Para tanto, opera uma série de processos ideológicos nos quais, por um lado, os dados pessoais são tratados como matéria-prima, naturalmente disponível para a expropriação de capital e, por outro, as empresas são consideradas as únicas capazes de tratar os dados e, portanto, apropriar-se deles.

Com referência a colonialismo e Inteligência Artificial, [Mohamed et al. \(2020\)](#) examinam como a colonialidade se apresenta em sistemas algorítmicos por meio de opressão algorítmica institucionalizada (a subordinação injusta de um grupo social às custas do privilégio de outro), exploração algorítmica (formas pelas quais os atores institucionais e empresas tiram proveito de pessoas já frequentemente marginalizadas em benefício

desequilibrado de tais setores) e desapropriação algorítmica (centralização do poder em poucos e desapropriação do poder de muitos), numa análise que busca evidenciar as permanências históricas das relações de poder.

E. Danos sócio-ambientais

A pesquisadora Kate Crawford (2021) preconiza uma visão mais abrangente da Inteligência Artificial como forma crítica de entender que esses sistemas dependem da exploração: por um lado, de recursos energéticos e minerais, de mão de obra barata e, além disso, de nossos dados em escala. Em outras palavras, a I.A. é uma indústria extrativista.

Todos esses sistemas fazem uso intensivo de energia e dependem fortemente de recursos minerais, às vezes extraídos de áreas onde estão disponíveis. Só na América Latina, temos o triângulo do lítio na Argentina, Bolívia e Chile, bem como vários depósitos de minerais conhecidos como 3TGs, pelas letras iniciais de seu nome em inglês (estanho, tungstênio, tântalo e ouro) na região amazônica, todos eles usados em dispositivos eletrônicos de ponta. Como afirmam Danae Tapia e Paz Peña, [as comunicações digitais são construídas sobre a exploração, embora “análises sociotécnicas do impacto ambiental das tecnologias digitais sejam quase inexistentes na comunidade hegemônica de direitos humanos que trabalha no contexto digital”.](#) (Tapia & Peña, 2021, tradução livre). E, para além do impacto ambiental, Camila Nobrega e Joana Varon também afirmam que [as narrativas da economia verde, em conjunto com os tecnossolucionismos, estão “ameaçando diversas formas de existência, de usos históricos e de gestão coletiva de territórios”,](#) não sendo por acaso que as autoras tenham descoberto que a empresa [Alphabet Inc., controladora do Google, está explorando minerais 3TGs em regiões da Amazônia onde existe conflito fundiário com povos indígenas](#) (Nobrega & Varon, 2021).

F. Automatização das políticas neoliberais

Segundo [Payal Arora \(2016\)](#), os discursos a respeito do *big data* têm uma conotação esmagadoramente positiva graças à ideia neoliberal de que a exploração com fins lucrativos dos dados dos pobres por empresas privadas apenas beneficiará a população. Do ponto de vista econômico, os estados de bem-estar digital estão profundamente entrelaçados com a lógica do mercado capitalista e, em particular, com as doutrinas neoliberais que buscam reduzir drasticamente o orçamento geral do sistema de bem-estar social, incluindo o número de beneficiários, a eliminação de alguns serviços, a introdução de exigências e formas intrusivas de condições à concessão de benefícios, a tal ponto que – como Alston afirmou (2019) – os cidadãos não mais se veem como sujeitos de direitos, mas como requerentes de serviços ([Alston, 2019](#), [Masiero & Das, 2019](#)). Nesse sentido, é interessante verificar que os sistemas de I.A., em seus esforços neoliberais de destinar os recursos públicos, também classificam quem é o sujeito pobre por meio de mecanismos automatizados de exclusão e inclusão (López, 2020).

G. Precarização do Trabalho

Com foco específico na inteligência artificial e nos algoritmos das empresas do Big Tech, a antropóloga Mary Gray e o cientista da computação Siddharth Suri apontam o “[trabalho fantasma](#)” ou o trabalho invisível que impulsiona as tecnologias digitais. Rotulagem de imagens e limpeza de bancos de dados são trabalhos manuais, muitas vezes realizados em condições de trabalho insalubres “para fazer a Internet parecer inteligente”. As comunidades dedicadas a esses empregos têm condições de trabalho muito precárias, normalmente marcadas por trabalho em excesso, mal remunerado, sem benefícios sociais ou estabilidade, muito diferentes das condições de trabalho dos criadores de tais sistemas (Crawford, 2021). Quem cuida do seu banco de dados? Como sempre, as profissões do cuidado não são reconhecidas como trabalho valioso.

H. Falta de transparência

[De acordo com o AINOW \(2018\)](#), quando órgãos públicos adotam ferramentas algorítmicas sem transparência, responsabilização e supervisão externa adequadas, seu uso pode ameaçar as liberdades civis e exacerbar problemas existentes nos órgãos públicos. Na mesma linha, a OCDE [\(Berryhill et al., 2019\)](#) postula que a transparência [por parte de] é estratégica para fomentar a confiança do público na ferramenta.

Opiniões mais críticas comentam a abordagem neoliberal quando a transparência depende da responsabilidade dos indivíduos, uma vez que eles não têm o tempo necessário ou o desejo de se comprometer com formas mais significativas de transparência e de consentimento online [\(Annany & Crawford, 2018\)](#). Assim, intermediários do governo com entendimento e independência específicos deveriam desempenhar seu papel aqui [\(Brevini & Pasquale, 2020\)](#). Além disso, Annany & Crawford (2018) sugerem que o que a visão atual de transparência em I.A. faz é fetichizar o objeto da tecnologia, sem entender que a tecnologia é um conjunto de atores humanos e não humanos, pelo que, para entender o funcionamento da I.A., é preciso ir além de simplesmente olhar para o objeto.

Aplicação do framework analítico em um caso prático

Um exemplo evidente de I.A que incorre em todas essas práticas apareceu entre os mapeamentos realizados pela Coding Rights no queinteligencia.org. O Projeto Horus um sistema que se vende como sendo capaz de prever gravidez na adolescência, mas que se tornou um exemplo bastante eloquente de como aparentes soluções tecnológicas de inteligência artificial podem apoiar políticas públicas discriminatórias, quando adotadas pelo poder público sem uma efetiva aferição de riscos, sejam eles individuais ou coletivos.

Apuramos que tal sistema começou a ser testado em 2015 pela Microsoft em uma parceria firmada entre a empresa e a província de Salta, na Argentina, oferecendo um sistema

piloto que prometia ser capaz de ajudar o poder público a combater a evasão escolar e a gravidez na adolescência por meio de uma ferramenta preditiva. A plataforma continha um conjunto de dados de mais de 12 mil mulheres, de 10 a 19 anos, que incluíam informações como idade, bairro, etnia, nível de escolaridade da pessoa chefe de família, presença de deficiências físicas e mentais, número de pessoas que dividem a mesma casa e, até mesmo, a disponibilidade, ou não, de água quente.

Nesse sistema, os algoritmos identificavam certas características nas pessoas que as tornassem, potencialmente, sujeitas à gravidez precoce, alertando o governo para que ele trabalhasse na prevenção. A tecnologia prometia, portanto, a previsão de quais meninas (ainda crianças) estariam predestinadas a ter uma gravidez na adolescência cinco ou seis anos depois, com uma acurácia prometida de 86%.

Surgiram muitas críticas a essa plataforma na época, com destaque para a análise técnica desenvolvida pelo Laboratório de Inteligência Artificial Aplicada da Universidade de Buenos Aires, que analisou a metodologia empregada pelos engenheiros da Microsoft. Eles concluíram pela existência de resultados superdimensionados, erros estatísticos grosseiros, bancos de dados tendenciosos, e coleta inadequada de dados pela estigmatização de mulheres pobres.

Mesmo diante das críticas, a iniciativa continuou a ser testada no país, inclusive em outras províncias argentinas (como La Rioja, Chaco e Terra do Fogo). Também foi exportado para a Colômbia (no município de La Guajira) e para o Brasil. Em 2019, a Microsoft impulsionou a exportação deste sistema para o Brasil. Assim, fomos o 5º país da América Latina a receber o projeto Horus, e a primeira cidade a testar o programa foi Campina Grande, no Estado da Paraíba. A princípio, para supostamente subsidiar a melhoria das ações do chamado Programa Criança Feliz/ Primeira Infância.

Nós da Coding Rights questionamos o Ministério da Cidadania por maiores informações sobre o acordo de cooperação técnica firmado com a empresa Microsoft, através

Coding Rights Projetos

<https://codingrights.org>
contact@codingrights.org

de pedidos de acesso à informação. Como resposta, recebemos as informações de que: (i) não houve repasses financeiros à empresa Microsoft; (ii) os bancos de dados utilizados para a construção de ferramentas analíticas e de inteligência artificial foram três: Sistema Único de Assistência Social, Cadastro Único e CADSUAS, do Ministério de Desenvolvimento Social.

Ou seja, toda a base de dados que compunha a plataforma era oriunda de programas sociais no Brasil. A coleta e formação de bancos de dados, portanto, se deu sobre informações de crianças pobres, do sexo feminino e em situação de vulnerabilidade social. Questionado mais uma vez por nós o Ministério da Cidadania, informaram que o acordo de cooperação vigorou por 6 meses (entre setembro de 2019 e março de 2020), e não havia informações referentes à presença, ou não, de margens de erro nas tecnologias envolvidas. O Ministério também alegou que como tratou-se de um piloto “não poderiam atender a solicitação de dados estatísticos sobre o seu uso e sua efetividade”. Ou seja, os resultados do piloto ficaram apenas como insumo para a Microsoft.

Esses sistemas de I.A. constituem um novo estágio na tecnocratização das políticas públicas. Os níveis de participação na sua concepção e a transparência do processo são questionáveis; no caso do Brasil, o convênio de prova de conceito previa um plano de trabalho de seis meses, o que deixa claro que não havia a intenção de se conduzir um processo inclusivo mais amplo com a população-alvo, mas tão somente a de implementar uma ferramenta tecnológica. Assim, não apenas os cidadãos brasileiros não tiveram acesso a qualquer dado de avaliação do piloto, mas também as pessoas afetadas por esses sistemas, ou seja, crianças, adolescentes e famílias pobres, sequer foram objeto de consulta, por não serem reconhecidas como partes interessadas. Da mesma forma, o consentimento para usar os dados com o objetivo de receber ou não benefícios sociais abre toda uma discussão sobre a ética desses sistemas, discussão que segue sem solução.

Da mesma forma, há evidências de que o uso de I.A. para prever possíveis vulnerabilidades não só não funciona bem no atendimento social de crianças e adolescentes

([Clayton et al., 2020](#)), como também acaba sendo bastante oneroso para os Governos, ao menos nos estágios iniciais, o que parece ser contrário à doutrina neoliberal ([Bright et al. 2019](#)). Quantas horas de recursos humanos de servidores públicos foram usadas em uma prova de conceito cujos resultados não estão documentados ou disponíveis ao público?

Em resumo, podemos dizer que a “Plataforma Tecnológica de Intervención Social” e o Projeto Horus são apenas um exemplo bastante eloquente de como a pretensa neutralidade da Inteligência Artificial tem sido cada vez mais implementada em alguns países da América Latina para apoiar políticas públicas potencialmente discriminatórias que poderiam prejudicar os direitos humanos das pessoas sem privilégios, bem como para monitorar e censurar as mulheres e seus direitos sexuais e reprodutivos. Analisando nosso quadro analítico, poderíamos dizer que preenche todos os quesitos.

Um exemplo de regulação protetiva restringindo usos específicos de sistemas de reconhecimento facial

Quando se trata de tecnologia, por vezes, prevalece uma visão de que tudo que é novo é bom e a implementação de determinada inovação é inevitável. Mas casos recentes de regulação de sistemas de reconhecimento facial demonstram o contrário.

No Brasil, a partir do dia 21 de junho de 2022, mais de 50 parlamentares de diferentes partidos apresentaram projetos de lei pelo banimento do reconhecimento facial em espaços públicos. A iniciativa #SaiDaMinhaCara¹ demonstrou um consenso multipartidário sobre o caráter invasivo e discriminatório dessa tecnologia, principalmente quando aplicada sob uma pretensa narrativa de segurança pública, mostra também que não tudo que é nova tecnologia deve ser utilizada para qualquer finalidade. Com esse protocolo de projetos de leis, o Brasil entra na tendência mundial de restringir determinados usos do reconhecimento facial.

¹ O texto que se segue é uma cópia da explicação coletiva redigida para a iniciativa.

Pólos de desenvolvimento dessa tecnologia já eliminaram seu uso para fins de policiamento. São Francisco, onde estão localizadas as grandes empresas de tecnologia do Vale do Silício, foi a primeira cidade dos Estados Unidos a aprovar, em 2019, a proibição do uso do reconhecimento facial pela polícia e outras agências públicas. Em 2020, foi a vez de Boston e Cambridge, onde estão localizados Harvard e o MIT, pólos universitários do debate tecnológico. Na Europa, em outubro de 2021, o Parlamento Europeu também votou a favor do banimento após manifestação da Autoridade Europeia para a Proteção de Dados (AEPD).

Agora, 12 estados e o Distrito Federal tiveram PLs apresentados no Brasil, seja no nível estadual ou municipal. São eles: Bahia, Ceará, Espírito Santo, Distrito Federal, Minas Gerais, Pará, Pernambuco, Paraná, Rio de Janeiro, Rio Grande do Sul, Santa Catarina, Sergipe e São Paulo.

As razões para o banimento são várias. Estudos comprovam que essa tecnologia é falha e cheia de vieses, com margem de erro particularmente gritantes quando se trata de rostos de pessoas negras, principalmente se forem mulheres ou pessoas trans. O resultado disso é que, se essa tecnologia é utilizada pela polícia e identifica erroneamente alguém, torna-se bem difícil argumentar contra uma máquina que você é você e não alguém procurado pela Justiça.

E isso já está acontecendo. No Brasil, desde 2019 houve uma expansão da implementação do reconhecimento facial sob pretexto da segurança pública. No segundo dia de teste dessa tecnologia no Rio de Janeiro, uma mulher foi detida ao ser confundida com uma mulher que já estava em privação de liberdade. No Piauí, um homem foi transferido para o DF e preso injustamente por três dias após um sistema de reconhecimento facial apontá-lo erroneamente como pessoa procurada. Em Salvador, um rapaz de 25 anos com necessidades especiais foi abordado por policiais após ser confundido com um homem procurado por assalto.

Com o temor de serem acusadas de alimentar racismo e mais violência policial, ainda mais depois do assassinato de George Floyd, grandes empresas como IBM, Microsoft e Amazon já não vendem esse tipo de tecnologia para autoridades estatais e para fins de

policiamento. Mas muitas outras empresas estão por aí, fazendo vendas milionárias e acordos com governos. Uma tecnologia falha e imprecisa que custa caro para o setor público. Na Bahia, anunciou-se a expansão do sistema de reconhecimento facial para mais de 70 municípios do interior, com o gasto de R\$ 665 milhões de reais. Em algumas cidades que “ganharão” as câmeras, faltam escolas, hospitais, serviços de acesso à Justiça, e outros. Mesmo quando uma empresa se oferece para “doar” câmeras, o barato sai caro: aumentam-se gastos com pessoal para usar uma tecnologia que não funciona; a empresa lucra testando seu algoritmo com dados da nossa cara e acaba criando laços com o ente público que podem facilitar ganhar licitações futuras sobre o tema.

Não se trata apenas de aprimorar o algoritmos falhos com tendências racistas. Mesmo que essas tecnologias evoluam para funcionar perfeitamente, outra crítica recorrente à implementação desses sistemas em espaços públicos é a vigilância em massa. Em uma perversa reversão da presunção de inocência, todas as pessoas no exercício de seu direito fundamental de ir e vir passam a ser tratadas como suspeitas, filmadas, vigiadas e potencialmente identificadas, sem consentimento. Essas são algumas das preocupações endereçadas pela iniciativa #SaiDaMinhaCara.

A iniciativa #SaiDaMinhaCara é estimulada pela Coding Rights, MediaLab-UFRJ/Rede Lavits, Instituto Brasileiro de Defesa do Consumidor (IDEC) e Centro de Estudos de Segurança e Cidadania – (CESeC), organizações especialistas em temas de tecnologia, segurança e direitos humanos que mobilizaram parlamentares em torno da pauta. Os primeiros projetos de lei foram protocolados em 08 de dezembro de 2021 pela Deputada Estadual Dani Monteiro - PSOL RJ (Presidente da Comissão de Direitos Humanos da Assembleia Legislativa do Estado do Rio de Janeiro) e em 27 de setembro de 2021 pelo Vereador Reimont - PT (Presidente da Comissão de Monitoramento no Tema Cidades Inteligentes, na Câmara Municipal do Rio de Janeiro).

Considerações finais

Estamos assistindo a um mundo em que os Governos estão cada vez mais adotando sistemas algorítmicos de tomada de decisão como uma varinha mágica para “resolver” problemas sociais, econômicos, ambientais e políticos, gastando recursos públicos de maneiras questionáveis, compartilhando dados pessoais sensíveis de cidadãos com empresas privadas e, em última instância, descartando qualquer tentativa de resposta coletiva, democrática e transparente aos principais desafios da sociedade. Um alerta evidente para a necessidade de uma regulamentação, mas essa regulação precisa promover uma visão crítica a essas soluções mágicas, principalmente quando se trata de sistemas de Inteligência Artificial que ganham escala por serem implementados na gestão de polícias do setor público.

Não acreditamos em uma I.A.. justa, ética e/ou inclusiva se os sistemas automatizados de decisão não reconhecerem as desigualdades e injustiças estruturais que afetam as pessoas cujas vidas constituem o alvo a ser gerenciado por esses sistemas. A transparência não basta se os desequilíbrios de poder históricos não forem levados em consideração. Além disso, somos críticas à ideia de sistemas de I.A. serem concebidos para gerenciar pessoas pobres ou quaisquer comunidades marginalizadas. Esses sistemas tendem a ser desenvolvidos por estratos da população privilegiados, contra o livre arbítrio e sem ouvir a opinião nem ter a participação desde o início daqueles que provavelmente serão alvos ou “ajudados”, resultando em opressão e discriminação automatizadas praticadas pelos Estados de Bem-Estar Digital, que recorrem à matemática como desculpa para se esquivar de qualquer responsabilidade política.

Diante disso, por tudo o exposto, a Coding Rights espera que essas considerações sejam levadas em conta no processo de elaboração de minuta de substitutivo dos Projetos de Lei nº 5.051, de 2019, 21, de 2020, e 872, de 2021, que tramitam no Congresso Nacional e têm como objetivo estabelecer princípios, regras, diretrizes e fundamentos para regular o desenvolvimento e a aplicação da inteligência artificial no Brasil. Também aproveitamos a

ocasião para disponibilizar a fala por escrito e na íntegra da participação da Coding Rights em Audiência Pública do Senado, no dia 12 de maio de 2022, sobre regulamentação do uso de IA:

<https://tinyurl.com/CodingAudienciaPublicaIA>

Desde já, agradecemos a atenção despendida. E ficamos à disposição para eventuais consultas específicas no texto de lei ou participações em reuniões ou eventos que visem melhor informar a Comissão de Juristas do Senado que trata de regulação sobre inteligência artificial, bem como explicitar melhor qualquer ponto que careça de maior nitidez.

Apresentamos à Vossa Excelência os protestos de elevada estima e mais distinta consideração.

Joana Varon Ferraz

Diretora Executiva da Coding Rights
Fellow de Direitos Humanos e Tecnologia
Carr Center da Harvard Kennedy School
joana@codingrights.org

Paz Peña

Consultora
paz@pazpena.com

Vanessa Koetz

Analista de Policy e Advocacy
vanessa@codingrights.org

Bianca Kremer

Analista de Policy e Advocacy
bianca@codingrights.org

ANEXO - Bibliografía

Aguilar, P. L. (2011). La feminización de la pobreza: conceptualizaciones actuales y potencialidades analíticas. *Revista Katálisis*, 14(1),126-133

AINOW. 2018. Algorithmic Accountability Policy Toolkit.
<https://ainowinstitute.org/aap-toolkit.pdf>

Alston, Philip. 2019. Report of the Special rapporteur on extreme poverty and human rights. Promotion and protection of human rights: Human rights questions, including alternative approaches for improving the effective enjoyment of human rights and fundamental freedoms. A/74/48037. Seventy-fourth session. Item 72(b) of the provisional agenda.

Ananny, M., & Crawford, K. (2018). Seeing without knowing: Limitations of the transparency ideal and its application to algorithmic accountability. *New Media & Society*, 20(3), 973–989.
<https://doi.org/10.1177/1461444816676645>

Arora, P. (2016). The Bottom of the Data Pyramid: Big Data and the Global South. *International Journal of Communication*, 10, 1681-1699.

Benjamin, R. (2019). Race after technology: Abolitionist tools for the new Jim code. *Polity*.

Berryhill, J., Kok Heang, K., Clogher, R. & McBride, K. (2019). Hello, World: Artificial Intelligence and its Use in the Public Sector. OECD.

Brevini, B., & Pasquale, F. (2020). Revisiting the Black Box Society by rethinking the political economy of big data. *Big Data & Society*. <https://doi.org/10.1177/2053951720935146>

Bridle, J. (2018). *New Dark Age: Technology, Knowledge and the End of the Future*. Verso Books.

Bright, J., Bharath, G., Seidelin, C., & Vogl, T. M. (2019). Data Science for Local Government (April 11, 2019). Available at SSRN: <https://ssrn.com/abstract=3370217> or <http://dx.doi.org/10.2139/ssrn.3370217>

Buolamwini, J. & Gebru, T. (2018). Gender Shades: Intersectional Accuracy Disparities in Commercial Gender Classification. Proceedings of the 1st Conference on Fairness, Accountability and Transparency, in PMLR 81:77-91
Castro, L. & López, J. (2021). Vigilancia a las “buenas madres”. Aportes desde una perspectiva feminista para la investigación sobre la datificación y la vigilancia en la política social desde Familias En Acción. Fundación Karisma. Colombia.

Couldry, N., & Mejias, U. (2019). Data colonialism: rethinking big data’s relation to the contemporary subject. *Television and New Media*, 20(4), 336-349.

Crawford, K. (2021). *Atlas of AI. Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press.

Díaz Hernández, M. (2020) Sistemas de protección social en Venezuela: vigilancia, género y derechos humanos. In *Sistemas de identificación y protección social en Venezuela y Bolivia. Impactos de género y otros tipos de discriminación*. Derechos Digitales.

Eubanks, V. (2018). *Automating inequality: how high-tech tools profile, police, and punish the poor*. First edition. New York, NY: St. Martin’s Press.

López, J. (2020). Experimentando con la pobreza: el SISBÉN y los proyectos de analítica de datos en Colombia. Fundación Karisma. Colombia.

Masiero, S., & Das, S. (2019). Datafying anti-poverty programmes: implications for data justice. *Information, Communication & Society*, 22(7), 916-933.

Mohamed, S., M. Png & W. Isaac. (2020). Decolonial AI: Decolonial Theory as Sociotechnical Foresight in Artificial Intelligence. *Philosophy & Technology*. <https://doi.org/10.1007/s13347-020-00405-8>

Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. New York University Press.

Nobrega, C; Varon, J. (2021). Big tech goes green(washing): feminist lenses to unveil new tools in the master's house. *Giswatch, APC*. Available at: <https://www.giswatch.org/node/6254>

O'Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown.

Silva M.R. & Varon, J. (2021). Reconhecimento facial no setor público e identidades trans: tecnopolíticas de controle e ameaça à diversidade de gênero em suas interseccionalidades de raça, classe e território. *Coding Rights*.

Taylor, L. (2017). What is data justice? The case for connecting digital rights and freedoms globally. *Big Data & Society*, July-December, 1-14. <https://journals.sagepub.com/doi/10.1177/2053951717736335>

Tapia, D; Peña. P. (2021). White gold, digital destruction: Research and awareness on the human rights implications of the extraction of lithium perpetrated by the tech industry in Latin American ecosystems. *Giswatch, APC*. Available at: <https://www.giswatch.org/node/6247>

Tendayi Achiume, E. (2020). Racial discrimination and emerging digital technologies: a human rights analysis. Report of the Special Rapporteur on contemporary forms of racism, racial discrimination, xenophobia and related intolerance. A/HRC/44/57. Human Rights Council. Forty-fourth session. 15 June–3 July 2020.

West, S.M., Whittaker, M. and Crawford, K. (2019). Discriminating Systems: Gender, Race and Power in AI. AI Now Institute.

Zuboff, S. (2019). The age of surveillance capitalism: the fight for a human future at the new frontier of power. Profile Books.