

Automated Ableism: A Systematic Review on AI and Discrimination against People with Disabilities

Janaina Nogueira de Souza Lopes¹, Valéria Quadros dos Reis¹, Amaury Antônio de Castro Junior¹, Anderson Corrêa de Lima¹

¹Faculdade de Computação – Universidade Federal de Mato Grosso do Sul (UFMS)
Campo Grande – Brasil

{janaina.nogueira, valeria.reis, amaury.junior, anderson.lima}@ufms.br

Abstract. Research Context: Artificial Intelligence (AI) systems have expanded in decision-making processes but raise concerns about algorithmic discrimination, with ableism remaining an underexplored dimension in Information Systems (IS). **Scientific and/or Practical Problem:** Despite advances in the literature on algorithmic bias, discrimination against people with disabilities (PwD) in AI systems remains underexplored. This gap limits the development of fairer and more inclusive AI systems. **Proposed Solution and/or Analysis:** This study presents a Systematic Literature Review (SLR) on how AI systems can reproduce, intensify, or mitigate ableist practices. **Related IS Theory:** It is grounded in algorithmic justice, fairness in machine learning, and socio-technical perspectives on inclusion. **Research Method:** Following PRISMA and PICOC, 26 articles published between 2020 and 2025 were analyzed and assessed against quality criteria. **Summary of Results:** Three main trends were identified: AI that reinforces ableism, mitigation initiatives, and ethical-regulatory debates. The problem has been discussed mainly by researchers from the United States and Europe, highlighting the need for studies that consider different languages as well as diverse social and cultural behaviors. **Contributions and Impact to IS area:** The study addresses a neglected dimension of algorithmic bias and provides insights for researchers, practitioners, and policymakers committed to developing inclusive AI.

1. Introdução

Os avanços em Inteligência Artificial (IA) transformaram a forma como decisões são automatizadas na sociedade. Sistemas baseados em aprendizado de máquina são capazes de executar tarefas complexas de maneira eficiente e em grande escala. No entanto, ao mesmo tempo em que esses sistemas promovem inovações e eficiência, crescem as preocupações quanto à reprodução e à amplificação de desigualdades sociais por meio da discriminação algorítmica. Segundo [Mehrabi et al. 2022], a discriminação algorítmica é uma fonte de injustiça originada de preconceitos humanos e estereótipos baseados em atributos sensíveis, que pode ocorrer de forma intencional ou não. Estudos demonstram que, quando treinados com dados históricos enviesados, os algoritmos tendem a perpetuar discriminações estruturais já existentes [Barocas and Selbst 2016, O’Neil 2016]. Dados enviesados, ao serem utilizados em modelos, incorporam padrões de desigualdade presentes na sociedade, influenciando as previsões e decisões produzidas e estando diretamente relacionados com a discriminação algorítmica.

A discussão sobre viés e discriminação em sistemas de IA tem se intensificado, especialmente no que diz respeito aos impactos sobre grupos socialmente marginalizados. Pode-se observar diversos estudos em diferentes áreas. Na questão racial, por exemplo, [Angwin et al. 2016] revelam um algoritmo que classificava pessoas negras como de maior risco de reincidência criminal, mesmo apresentando histórico semelhante ao de pessoas brancas. Quanto ao gênero, [Bolukbasi et al. 2016] demonstraram que vetores semânticos treinados em grandes corpora reproduzem estereótipos sexistas, como associar “homem” a “programador” e “mulher” a “doméstica”. Em relação à classe social, [Eubanks 2018] detalha como sistemas automatizados usados em políticas públicas de assistência social contribuíram para a exclusão sistemática de populações pobres de benefícios governamentais, baseando-se em critérios que reproduziam preconceitos históricos.

Apesar dos avanços nas discussões sobre justiça algorítmica, um eixo de exclusão ainda estigmatizado na literatura é o capacitismo, ou seja, a discriminação sistemática contra pessoas com deficiência. Segundo [Griffin et al. 2007], o capacitismo é um sistema de crenças e práticas que privilegia corpos e mentes considerados “normais”, marginalizando aqueles que se desviam desse padrão ideal de funcionamento físico, sensorial ou cognitivo. Essa lógica estabelece uma norma corporal e mental dominante, tratando qualquer diferença como inadequada ou indesejável.

A discussão sobre o capacitismo nas pesquisas em IA é especialmente preocupante diante do risco de que sistemas automatizados, ao se apoiarem em dados enviesados e não representativos, reforcem estereótipos históricos. Conforme apontam [Whittaker et al. 2019], a falta de dados acessíveis, inclusivos e sensíveis à diversidade corporal e cognitiva pode resultar em decisões algorítmicas injustas e imprecisas, aprofundando desigualdades sociais e estruturais já existentes.

No contexto brasileiro, embora tenham ocorrido avanços nas políticas de inclusão social nas últimas décadas, persistem barreiras significativas que dificultam a promoção da equidade para pessoas com deficiência, especialmente nos campos da educação e do trabalho. Dados do Instituto Brasileiro de Geografia e Estatística indicam que a taxa de analfabetismo entre pessoas com deficiência é de 19,5%, evidenciando um grande desafio educacional [IBGE 2023]. Além disso, 63,3% das pessoas com deficiência com 25 anos ou mais não possuem instrução ou o ensino fundamental completo, enquanto apenas 11,1% concluíram o ensino fundamental. No mercado de trabalho, a situação é igualmente desafiadora: apenas 26,6% das pessoas com deficiência têm acesso ao emprego formal e, entre essas, cerca de 55% atuam em condições de informalidade.

Esses indicadores revelam não apenas a exclusão e a desigualdade estrutural vivida pelas pessoas com deficiência, mas também os desafios em garantir o acesso a direitos fundamentais. Com o avanço das tecnologias digitais e o uso crescente da IA em processos decisórios, como na educação, no trabalho ou em políticas públicas, torna-se imprescindível refletir criticamente sobre os riscos de reprodução de práticas excludentes, bem como sobre o potencial dessas tecnologias para promover maior equidade e inclusão social.

As tecnologias digitais oferecem caminhos concretos para ampliar a inclusão social e reduzir desigualdades. Diversos estudos evidenciam esse potencial, seja por meio

de jogos digitais acessíveis voltados ao desenvolvimento do pensamento computacional em crianças com deficiência intelectual [Dutra et al. 2023], da educação a distância como estratégia inclusiva no ensino de programação [Santos et al. 2017], ou da adoção de metodologias de avaliação centradas no usuário para o desenvolvimento de tecnologias assistivas [Rodrigues et al. 2018]. Além disso, aplicações recentes de IA têm contribuído para ampliar a acessibilidade digital e favorecer a autonomia de pessoas com deficiência [Jacques et al. 2025, N. et al. 2025], reforçando a relevância dessas inovações no enfrentamento das desigualdades.

Diante desse cenário, esta Revisão Sistemática da Literatura (RSL) tem como objetivo identificar, analisar e categorizar evidências científicas sobre o papel da IA na relação com a discriminação contra pessoas com deficiência (PcD), com foco em práticas capacitistas e viés algorítmico. No contexto brasileiro, o capacitismo se expressa tanto no acesso restrito a uma educação básica de qualidade quanto na exclusão sistemática do mercado de trabalho formal. Esse quadro revela que a discriminação contra PcDs não decorre apenas da condição de deficiência, mas resulta também de barreiras sociais, culturais e institucionais que limitam sua participação plena. Assim, ao observar esse contexto com os avanços tecnológicos, busca-se compreender de que maneira a IA pode atuar não apenas como agente de perpetuação ou intensificação do capacitismo, mas também como ferramenta de mitigação dessa forma de discriminação, contribuindo para a construção de tecnologias e sistemas mais inclusivos e justos.

Para isso, este estudo investigará as diferentes abordagens e metodologias utilizadas em IA para lidar com o capacitismo, por meio de uma RSL, que pretende responder às seguintes questões de pesquisa:

- QP1: Qual é o papel atribuído à Inteligência Artificial nos estudos analisados: reprodução, amplificação ou mitigação de práticas capacitistas?
- QP2: Quais são as principais metodologias utilizadas na detecção de discurso capacitista?
- QP3: Quais são as perspectivas futuras para o uso de IA no combate ao capacitismo em tecnologias educacionais e em sistemas de IA?

A QP1 busca compreender qual é o papel atribuído à IA nos estudos analisados: se ela atua de maneira contrária à discriminação (mitigação) ou favorável a ela em uma escala similar à encontrada na sociedade (reprodução) ou aumentada (amplificação). Esta questão é fundamental para identificar como a literatura científica posiciona a IA frente às desigualdades enfrentadas por pessoas com deficiência, especialmente considerando seu impacto crescente em decisões sociais automatizadas.

A QP2 tem como foco identificar quais são as principais metodologias utilizadas na detecção de discurso capacitista. Essa investigação é essencial para mapear os métodos computacionais e abordagens teóricas adotadas nos estudos, avaliando sua robustez, limitações e aplicabilidade na identificação de viés algorítmico e linguagem discriminatória.

Por fim, a QP3 busca explorar as perspectivas futuras para o uso da IA no combate ao capacitismo, com ênfase em tecnologias educacionais e sistemas algorítmicos. Considerando os desafios enfrentados pelas pessoas com deficiência nos campos da educação e do trabalho, esta questão visa compreender o potencial da IA como aliada na promoção

da equidade e da inclusão social, bem como identificar direções promissoras apontadas pela literatura para o desenvolvimento de tecnologias mais justas.

Até onde vai nosso conhecimento, e considerando o escopo das fontes consultadas, não identificamos estudos que realizem uma revisão sistemática dedicada à análise da relação entre capacitismo e discriminação algorítmica em sistemas de IA. Tal lacuna na literatura evidencia a originalidade e a relevância do presente estudo.

2. Trabalhos relacionados

Os trabalhos relacionados abordam percepções de justiça em decisões algorítmicas, estratégias de identificação e mitigação de vieses em modelos de aprendizado de máquina, aplicações específicas como diagnósticos radiológicos e estudos sobre IA e deficiência, em geral restritos à visão médica e assistiva. Apesar desses avanços, ainda há uma lacuna importante no que se refere ao capacitismo e discriminação algorítmica da IA sobre PcD.

A revisão conduzida por [Starke et al. 2022] foca especificamente nas **percepções de justiça em decisões algorítmicas**, a partir de uma análise sistemática de 39 estudos. A revisão demonstra que há grande diversidade nas definições teóricas e metodológicas de justiça algorítmica, e aponta que os estudos se concentram majoritariamente em países ocidentais. O trabalho enfatiza a necessidade de abordagens interdisciplinares que considerem o contexto social e cultural na avaliação da justiça algorítmica, contribuindo assim para o desenvolvimento de sistemas mais inclusivos. Esse enfoque é particularmente útil para pesquisas que pretendem investigar como pessoas com deficiência percebem ou são afetadas por decisões automatizadas, ainda que esse grupo específico não tenha sido o foco da RSL.

Já [Pagano et al. 2022] conduzem uma RSL sobre viés e injustiça em modelos de aprendizado de máquina. A partir da análise de 45 artigos publicados entre 2017 e 2022, os autores investigam as principais abordagens para **identificação e mitigação de vieses, com ênfase em métricas de justiça, técnicas de processamento e ferramentas**. A revisão destaca a diversidade de métricas e estratégias empregadas, bem como a dificuldade de selecionar critérios adequados para contextos distintos. Além disso, são discutidos os desafios relacionados à transparência e explicabilidade dos modelos. Embora a RSL aborde vieses relacionados a atributos como gênero e raça, ela também evidencia a ausência de estudos voltados a outros grupos minorizados, como pessoas com deficiência, sinalizando uma lacuna na literatura que justifica novas investigações.

[Lima et al. 2024] realizaram uma RSL voltada à identificação de técnicas de justiça algorítmica aplicadas ao diagnóstico radiológico com o uso de IA. Os autores analisaram estudos publicados entre 2018 e 2023, investigando as estratégias utilizadas para **mitigar vieses relacionados a atributos sensíveis como sexo, idade e etnia**. Apesar do foco restrito à área médica, a RSL oferece contribuições metodológicas relevantes para investigações mais amplas sobre viés e justiça algorítmica.

Por fim, [Morr et al. 2024] realizaram uma revisão sistemática de escopo sobre **IA e deficiência**, analisando 45 estudos. Identificaram que a grande parte das pesquisas adota uma visão médica da deficiência, com foco em diagnóstico e tecnologias assistivas. Apenas alguns dos trabalhos consideram aspectos sociais ou justiça da deficiência. Os autores destacam a ausência de estratégias contra viés algorítmico e recomendam uma abordagem mais inclusiva e interdisciplinar no desenvolvimento de IA.

Em consonância com os trabalhos relacionados que abordam justiça algorítmica e viés em sistemas de IA, este estudo contribui ao avançar a discussão para o campo do capacitismo, propondo uma análise sistemática sobre como a IA pode reproduzir, intensificar ou mitigar práticas capacitistas. Diferentemente das revisões anteriores, que tratam majoritariamente de vieses ligados a gênero, raça ou idade, ou ainda restringem-se a perspectivas médicas e assistivas da deficiência, esta RSL busca preencher a lacuna existente ao focar especificamente na dimensão discriminatória associada à PcD no âmbito da IA.

3. Metodologia

A presente RSL foi conduzida com base em um protocolo estruturado, visando garantir a reprodutibilidade, transparência e abrangência da busca por estudos relevantes sobre o papel da IA na perpetuação, amplificação ou mitigação do capacitismo. A metodologia adotada seguiu os princípios das diretrizes PRISMA (Preferred Reporting Items for Systematic reviews and Meta-Analyses) descritas por [Page et al. 2021]. O protocolo tem sido amplamente empregado em pesquisas por promover transparência e rastreabilidade. A adoção do PRISMA oferece um arcabouço metodológico consolidado para a condução de RSL também no contexto da computação.

A seleção e organização dos estudos foram realizadas utilizando a plataforma Parsifal¹, que possibilitou a gestão centralizada das etapas da revisão, incluindo definição da string de busca, critérios de inclusão e exclusão, extração de dados e registro das decisões de triagem. Essa ferramenta contribuiu para manter o controle e a rastreabilidade de todo o processo, assegurando maior consistência e transparência metodológica².

3.1. Processo de seleção e fontes de informação

Inicialmente, foi definido o protocolo da revisão, incluindo o objetivo geral, as perguntas de pesquisa e a estrutura PICOC. A estrutura PICOC, de acordo com Carrera-Rivera et al. (2022) decompõe os objetivos da RSL em componentes (*Population, Intervention, Comparison, Outcome e Context*), que se tratam de palavras-chave e auxiliam na formulação de questões de pesquisa de forma clara e sistemática [Carrera-Rivera et al. 2022]. Com base nesse protocolo, elaborou-se a *string* de busca, composta por termos relacionados ao capacitismo e à IA, bem como seus sinônimos e variações linguísticas, combinados por operadores booleanos.

A *string* de busca foi aplicada nas seguintes bases de dados: ACL Anthology, ACM Digital Library, El Compendex, IEEE Xplore, ISI Web of Science, Scopus e SOL (SBC Open Library). Para acesso ao conteúdo completo nas bases: IEEE Xplore e Scopus foi utilizado o Café por meio do Portal de Periódicos da CAPES. Foram utilizados filtros por período de publicação (2020–2025) e por acesso aberto. O período de publicação entre 2020 e 2025 foi escolhido para contemplar os modelos de aprendizado mais recentes e relevantes, o que inclui não somente grandes modelos de linguagem, mas também modelos de inteligência artificial generativa. As buscas foram realizadas entre os dias 23 de maio e 6 de junho de 2025.

Na etapa inicial da triagem, foram removidos os registros duplicados e realizada a leitura dos títulos e resumos, com base nos critérios de inclusão e exclusão previamente

¹<https://parsif.al/>

²Dados da RSL: https://github.com/janaina-nogueira/RSL_IA_SBSI2026.git

estabelecidos. Os estudos potencialmente relevantes foram então avaliados na íntegra, considerando sua pertinência temática e metodológica em relação ao escopo da pesquisa.

Após a seleção final, foi realizada a Avaliação de Qualidade dos estudos incluídos, com o objetivo de assegurar a robustez e a confiabilidade das evidências analisadas. Por fim, os dados extraídos foram organizados, analisados e discutidos à luz das perguntas de pesquisa, evidenciando tendências, lacunas e oportunidades futuras para a ética referente à IA no contexto do capacitismo.

3.2. Critérios de elegibilidade

Foram definidos critérios de inclusão e exclusão com o objetivo de refinar os resultados e garantir a relevância dos estudos selecionados. Os critérios de inclusão consideraram publicações que abordassem diretamente a discriminação de PcD no contexto da IA, com foco em práticas capacitistas, viés algorítmico ou linguagem discriminatória.

Foram incluídos apenas artigos publicados entre 2020 e 2025, redigidos em português ou inglês, disponíveis em acesso aberto e que empregassem técnicas de IA, como modelos de linguagem, aprendizado de máquina, Processamento de Linguagem Natural ou análise de corpus automatizada. Foram excluídos estudos que tratassem exclusivamente de dados não textuais (como imagens), que não abordassem diretamente o capacitismo ou a discriminação contra PcD, bem como estudos secundários, como revisões sistemáticas e resumos expandidos. Também foram eliminados trabalhos duplicados, publicações sem acesso ao conteúdo completo ou que não se enquadrassem nos critérios metodológicos estabelecidos.

3.3. Estratégia de busca

Para a *string* de busca, foram utilizados termos relacionados ao capacitismo e às tecnologias de IA, com o intuito de capturar estudos que analisassem possíveis vieses discriminatórios automatizados. A seguir, apresenta-se a *string* de busca utilizada:

```
("Capacitismo" OR "Ableism" OR "Disability Bias" OR "Disability  
Discrimination" OR "Discriminação pessoa com deficiência")  
OR  
("Automated Ableism" OR "Capacitismo automatizado")  
AND  
("Inteligência artificial" OR "AI" OR "Artificial intelligence"  
OR "IA" OR "Aprendizado de máquina" OR "Machine learning" OR "ML"  
OR "Modelo de Linguagem" OR "Language model" OR "Processamento  
de Linguagem Natural" OR "PLN" OR "Natural Language Processing"  
OR "NLP" OR "Corpus Analysis" OR "Análise de Corpus" OR  
"Leitura distante" OR "Distant Reading" OR "Deep learning")
```

3.4. Seleção dos estudos

Na Tabela 1 é apresentada a distribuição do número de estudos identificados em cada base após a aplicação das estratégias de busca.

Foram encontradas inicialmente 599 publicações. Dentre essas, 37 foram removidas por se tratarem de duplicatas e 48 por entradas incompletas, como resumos estendidos ou apenas citações de conferências, resultando em um total de 514 estudos únicos.

A triagem inicial foi realizada por meio da leitura dos títulos e resumos, com base em critérios previamente estabelecidos de inclusão e exclusão. Nessa etapa, 488

Base de Dados	Número de Estudos
ACL Anthology	9
ACM Digital Library	436
Compendex (El Compendex)	22
IEEE Digital Library	28
ISI Web of Science	62
Scopus	42
SOL	0
Total	599

Tabela 1. Número de estudos identificados em cada base de dados

publicações foram rejeitados por não atenderem ao escopo da pesquisa. Com isso, 26 estudos foram selecionados para leitura na íntegra.

Na Figura 1 é apresentado o fluxo do processo de seleção dos estudos, elaborado conforme o modelo PRISMA [Page et al. 2021].

3.5. Avaliação de Qualidade

A Avaliação de Qualidade dos estudos foi realizada com base em uma lista personalizada, composta por cinco perguntas, com respostas em: “Sim” (1 ponto), “Parcialmente” (0,5 ponto) e “Não” (0 pontos). A pontuação máxima foi de 5 pontos por artigo. As perguntas avaliadas foram:

- O objetivo da pesquisa está claramente descrito?
- O estudo realizou um experimento bem descrito e passível de replicação, com dados e códigos disponibilizados?
- As conclusões estão claras?
- O artigo é bem estruturado e fácil de seguir?
- O artigo trata do tema objetivo da RSL?

Os critérios de qualidade foram definidos para verificar se os estudos eram claros, bem estruturados e relevantes para a revisão, sem a intenção de excluir ou penalizar trabalhos. A pergunta sobre experimento replicável foi criada pensando nos estudos empíricos de IA, que costumam disponibilizar dados e código, mas a revisão também incluiu artigos conceituais e jurídicos, nos quais esse tipo de avaliação não se aplica. Nesses casos, o critério foi interpretado de forma adequada ao tipo de pesquisa, garantindo que diferentes abordagens fossem consideradas de maneira justa na análise dos resultados.

Os critérios foram definidos de maneira ampla para permitir a avaliação de estudos com naturezas metodológicas distintas (empíricos, conceituais e jurídicos), evitando a exclusão indevida de contribuições relevantes. Os resultados dessa avaliação auxiliaram na avaliação dos estudos incluídos, considerando sua confiabilidade metodológica e relevância para os objetivos da RSL.

4. Resultados

Na Tabela 2 são apresentados os estudos selecionados para esta revisão sistemática, incluindo o identificador, autor e ano da publicação, o título e a pontuação de Avaliação de Qualidade atribuída a cada trabalho.

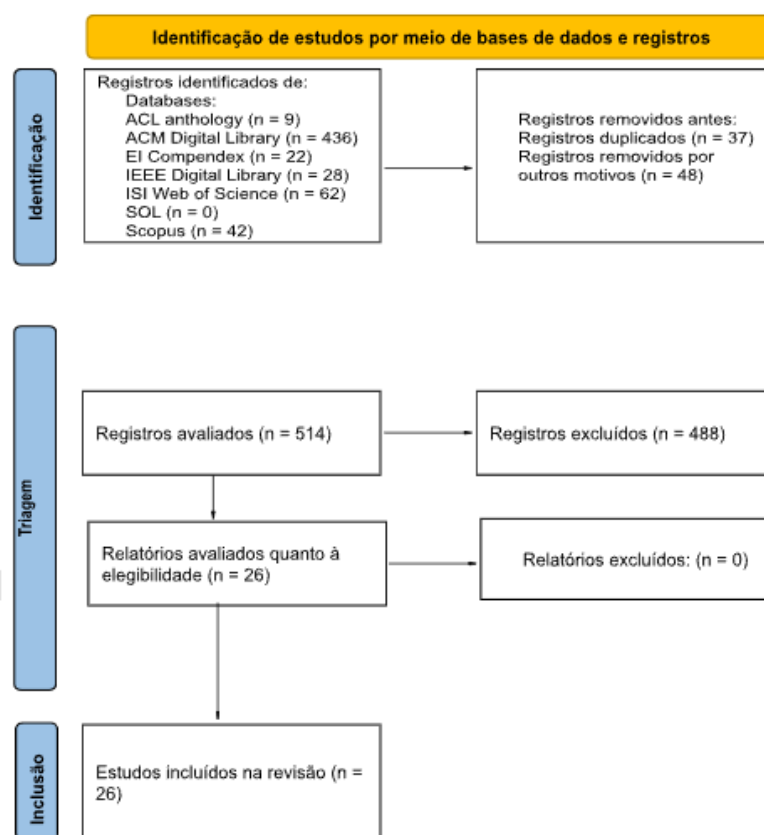


Figura 1. Diagrama de fluxo PRISMA 2020

Esses dados fornecem uma visão geral que guia as análises detalhadas acerca das abordagens metodológicas e das contribuições específicas dos trabalhos revisados para a compreensão da relação da IA e discriminação algorítmica no contexto do capacitismo.

Entre os 26 estudos selecionados, cinco se destacaram por alcançar pontuação máxima na avaliação de qualidade, evidenciando elevada relevância e rigor metodológico. Esses trabalhos investigam vieses capacitistas em modelos de linguagem, propõem abordagens interseccionais para sua análise, discutem implicações éticas em contextos educacionais, examinam barreiras impostas por preconceitos presentes em sistemas de PLN e contribuem com a criação de datasets voltados à representação de pessoas com deficiência. Em conjunto, esses estudos apresentaram objetivos bem definidos, metodologias claras e aderência direta ao tema da discriminação algorítmica contra PcD, servindo como referências centrais para a análise dos resultados desta RSL.

Ademais, os resultados estão organizados de acordo com as questões de pesquisa definidas: na Subseção 4.2 apresenta-se o papel atribuído à Inteligência Artificial nos estudos analisados, considerando se ela atua na reprodução, amplificação ou mitigação de práticas capacitistas; a Subseção 4.3 discute as principais metodologias utilizadas na detecção de discurso capacitista; e a Subseção 4.4 aborda as perspectivas futuras para o uso de IA no combate ao capacitismo em tecnologias educacionais e em sistemas de IA. Além dessas questões, outras análises complementares são apresentadas na Subseção 4.1.

4.1. Características dos estudos

Além da Tabela 2, que apresenta os estudos selecionados com base em autor, ano e pontuação da avaliação de qualidade, a Tabela 3 amplia essa análise ao oferecer uma caracterização mais abrangente dos artigos revisados. Essa tabela organiza os estudos segundo o tipo de metodologia empregada, o papel atribuído à IA (como agente de reprodução, mitigação ou abordagem conceitual do capacitismo), a localidade aproximada da publicação e o foco temático principal. Essa categorização contribui para a identificação de padrões recorrentes e lacunas na literatura, oferecendo oportunidade para uma análise mais profunda das contribuições dos estudos no contexto abordado.

A literatura revisada abrange o período de 2020 até o início de 2025, refletindo um interesse acadêmico crescente na interseção entre deficiência, preconceito e IA. Observa-se um aumento na quantidade de publicações ao longo dos anos, especialmente em 2023 e 2024, o que sugere um impulso nas pesquisas que abordam o capacitismo e a discriminação algorítmica em modelos de linguagem. Essa tendência está ilustrada na Figura 2.

Outro aspecto importante a se observar é a localidade de investigação ou publicação dos estudos. A maioria dos trabalhos revisados foi conduzida ou publicada em países do hemisfério norte, em especial nos Estados Unidos e Europa. Entretanto, essa centralização também revela uma lacuna importante na literatura produzida em contextos do Sul Global, incluindo América Latina, África e partes da Ásia. A escassez de publicações nesses locais fomenta uma sub-representação de realidades sociais, culturais e políticas distintas.

ID	Autor e Ano	Título	Índice de qualidade
A01	[Park et al. 2025]	<i>“As an Autistic Person Myself:” The Bias Paradox Around Autism in LLMs</i>	4,5
A02	[Venkit et al. 2022]	<i>A Study of Implicit Bias in Pretrained Language Models against People with Disabilities</i>	5,0
A03	[Shew 2020]	<i>Ableism, Technoableism, and Future AI</i>	3,5
A04	[Marks 2020]	<i>Algorithmic Disability Discrimination</i>	3,0
A05	[Glazko et al. 2023]	<i>An Autoethnographic Case Study of Generative Artificial Intelligence’s Utility for Accessibility</i>	4,5
A06	[Herold et al. 2022]	<i>Applying the Stereotype Content Model to assess disability bias in popular pre-trained NLP models underlying AI-based assistive technologies</i>	4,5
A07	[Glazko et al. 2025]	<i>Autoethnographic Insights from Neurodivergent GAI “Power Users”</i>	4,5
A08	[Narayanan Venkit et al. 2023]	<i>Automated Ableism: An Exploration of Explicit Disability Biases in Sentiment and Toxicity Analysis Models</i>	4,5
A09	[Newman et al. 2024]	<i>Crafting Disability Fairness Learning in Data Science: A Student-Centric Pedagogical Approach</i>	4,0
A10	[Li et al. 2024]	<i>Decoding Ableism in Large Language Models: An Intersectional Approach</i>	5,0
A11	[Krupiy and Scheinin 2023]	<i>Disability Discrimination in the Digital Realm: How the ICRPD Applies to Artificial Intelligence Decision-Making Processes and Helps in Determining the State of International Human Rights Law</i>	3,5
A12	[Urbina et al. 2025]	<i>Disability Ethics and Education in the Age of Artificial Intelligence: Identifying Ability Bias in ChatGPT and Gemini</i>	5,0
A13	[Packin 2021]	<i>Disability discrimination using artificial intelligence systems and social scoring: Can we disable digital bias?</i>	3,5
A14	[Mukherjee et al. 2023]	<i>Global Voices, Local Biases: Socio-Cultural Prejudices across Languages</i>	4,5
A15	[Binns and Kirkham 2021]	<i>How Could Equality and Data Protection Law Shape AI Fairness for People with Disabilities?</i>	4,0
A16	[Glazko et al. 2024]	<i>Identifying and Improving Disability Bias in GPT-Based Resume Screening</i>	4,5
A17	[Wu and Ebling 2024]	<i>Investigating Ableism in LLMs through Multi-turn Conversation</i>	4,5
A18	[Manzoor et al. 2024]	<i>Out of dataset, out of algorithm, out of mind: a critical evaluation of AI bias against disabled people</i>	4,5
A19	[Elglaly and Liu 2023]	<i>Promoting Machine Learning Fairness Education through Active Learning and Reflective Practices</i>	4,5
A20	[Moss 2021]	<i>Screened out onscreen: Disability discrimination, hiring bias, and artificial intelligence</i>	4,0
A21	[Hutchinson et al. 2020]	<i>Social Biases in NLP Models as Barriers for Persons with Disabilities</i>	5,0
A22	[Buyl et al. 2022]	<i>Tackling Algorithmic Disability Discrimination in the Hiring Process: An Ethical, Legal and Technical Analysis</i>	4,0
A23	[Narayanan Venkit 2023]	<i>Towards a Holistic Approach: Understanding Socio-demographic Biases in NLP Models using an Interdisciplinary Lens</i>	4,0
A24	[Hassan et al. 2021]	<i>Unpacking the Interdependent Systems of Discrimination: Ableist Bias in NLP Systems through an Intersectional Lens</i>	4,5
A25	[Tadimalla et al. 2024]	<i>WIP: Moving from Accessibility to Anti-Ableism through the Explication of Disability in the AI Ecosystem</i>	4,5
A26	[Mondal et al. 2022]	<i>“#DisabledOnIndianTwitter”: A Dataset towards Understanding the Expression of People with Disabilities on Indian Twitter</i>	5,0

Tabela 2. Artigos selecionados na RSL.

ID	Tipo de Metodologia	Papel da IA	Localidade	Foco Temático
A01	NLP Bias	Reproduz	EUA	LLMs e autismo
A02	Bias implícito	Reproduz	EUA	Viés implícito em LLMs
A03	Conceitual	Conceitual	Europa	Technoableism e futuro da IA
A04	Jurídica	Conceitual	EUA	Discriminação algorítmica
A05	Qualitativa Auto-etnográfica	Mitiga	EUA	Acessibilidade com GAI
A06	Modelo de estereótipos	Reproduz	EUA	Bias em modelos assistivos
A07	Qualitativa Auto-etnográfica	Mitiga	EUA	Neurodivergência e uso de GAI
A08	Toxicidade	Reproduz	EUA	Bias explícito em NLP
A09	Conceitual	Mitiga	EUA	Ensino de justiça algorítmica
A10	LLMs	Reproduz	EUA	Capacitismo interseccional em LLMs
A11	Jurídica	Conceitual	Internacional	Direitos humanos digitais internacionais
A12	Educacional	Mitiga	EUA	ChatGPT, Gemini e ética educacional
A13	Jurídica	Conceitual	Europa	Exclusão digital
A14	Análise de dataset	Reproduz	Internacional	Viés sociocultural nos dados
A15	Jurídica	Conceitual	Europa	Direito e equidade na IA
A16	NLP Bias	Reproduz	EUA	Triagem de currículos com GPT
A17	LLMs	Reproduz	EUA	Bias em interações multivolta
A18	Datasets	Reproduz	Europa	Exclusão algorítmica por dados ausentes
A19	Auditoria educacional	Mitiga	EUA	Educação crítica e fairness
A20	Jurídica	Conceitual	Europa	Discriminação no recrutamento
A21	Bias	Reproduz	EUA	Barreiras em modelos de PLN
A22	Jurídica	Conceitual	Europa	Ética e IA no processo seletivo
A23	NLP	Reproduz	EUA	Viés sociodemográfico em NLP
A24	NLP	Mitiga	EUA	Interseccionalidade em NLP
A25	Conceitual	Conceitual	EUA	Anticapacitismo no ecossistema de IA
A26	Dataset	Reproduz	Ásia	Expressão de PcD no Twitter

Tabela 3. Caracterização expandida dos artigos selecionados na RSL

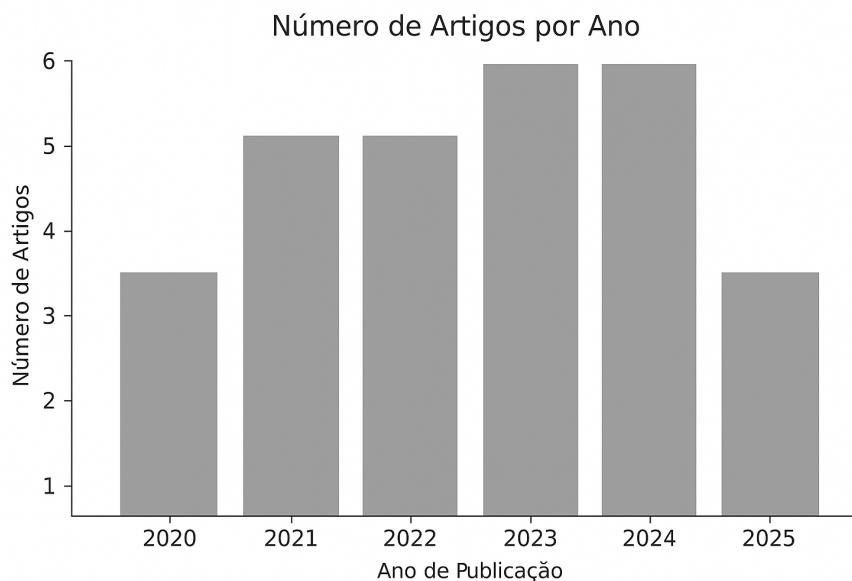


Figura 2. Número de artigos por ano de publicação (2020–2025).

4.2. Qual é o papel atribuído à Inteligência Artificial nos estudos analisados: reprodução, amplificação ou mitigação de práticas capacitistas?

A análise dos artigos selecionados revela que a maioria dos estudos considera a IA como um agente que **reproduz ou amplifica práticas capacitistas**. As publicações revisadas destacam a forma como modelos de linguagem natural e algoritmos de decisão automatizada incorporam ou reforçam estereótipos, preconceitos e exclusões estruturais já presentes na sociedade. Estudos como **A02, A06, A08, A10, A16, A17, A18, A21 e A23** evidenciam como esses sistemas podem reproduzir vieses capacitistas presentes nos dados de treinamento ou em decisões algorítmicas.

Por outro lado, uma parcela menor dos estudos propõe a IA como ferramenta de mitigação do capacitismo. Trabalhos como **A09, A12, A19 e A24** discutem abordagens e estratégias de aprendizado de máquina justo para detectar e corrigir práticas discriminatórias contra pessoas com deficiência.

Além disso, artigos como **A03, A04, A11, A13, A15, A20, A22 e A25** apresentam abordagens mais conceituais ou jurídicas, discutindo os fundamentos éticos, normativos e legais que envolvem o uso da IA em contextos relacionados à deficiência, com foco nos riscos e potencialidades do ponto de vista dos direitos humanos.

4.3. Quais são as principais metodologias utilizadas na detecção de discurso capacitista?

As metodologias identificadas nos artigos variam em complexidade e escopo, mas podem ser agrupadas em algumas categorias principais. A primeira delas é a análise de viés em modelos de linguagem natural, presente nos artigos **A01, A02, A06, A08, A10, A16, A17, A21 e A23**. Esses estudos utilizam modelos como BERT, GPT e suas variantes para identificar padrões discriminatórios ou enviesados nos *outputs* gerados por sistemas de IA, aplicando técnicas como análise de toxicidade, polaridade de sentimentos e analogias semânticas.

Outra abordagem recorrente é a avaliação de *datasets*, como observado nos artigos **A14**, **A18** e **A26**, que examinam os dados utilizados no treinamento de sistemas de IA com foco na sub-representação ou uso de descrições estigmatizantes de pessoas com deficiência.

Destacam-se também os estudos autoetnográficos e qualitativos, como os artigos **A05** e **A07**, que utilizam as experiências de pessoas neurodivergentes para avaliar a acessibilidade e os impactos subjetivos causados por sistemas de IA.

Além disso, alguns trabalhos propõem auditorias técnicas de equidade, como evidenciado nos artigos **A19** e **A24**, os quais desenvolvem métricas e indicadores técnicos para mensurar desigualdades de tratamento, erros de classificação e impactos discriminatórios nos sistemas analisados.

Por fim, uma abordagem de destaque envolve as análises jurídicas e conceituais, como nos artigos **A04**, **A11**, **A13**, **A15**, **A20** e **A22**, que discutem o discurso capacitista com base em legislações e marcos regulatórios, ressaltando a importância de fundamentos ético-legais no desenvolvimento de tecnologias.

4.4. Quais são as perspectivas futuras para o uso de IA no combate ao capacitismo em tecnologias educacionais e em sistemas de IA?

As perspectivas futuras delimitadas pelos artigos revisados apontam para uma crescente preocupação com o desenvolvimento de tecnologias mais inclusivas e com a formação ética e crítica de profissionais da área. Um dos caminhos mais enfatizados é a educação crítica em ciência de dados, como defendido pelos artigos **A09**, **A12**, **A19** e **A24**, que sugerem a inclusão de temas como justiça social, equidade e deficiência na de ciência de dados.

Outro ponto amplamente discutido é a construção de *datasets* mais diversos e representativos, conforme indicam os artigos **A14**, **A18**, **A21** e **A26**. Esses estudos ressaltam a importância de bases de dados que representem com dignidade e respeito a vivência de pessoas com deficiência, como forma de reduzir vieses estruturais.

A criação de tecnologias voltadas à acessibilidade também surge nos artigos **A05**, **A07** e **A25**, que destacam o papel que sistemas de IA podem desempenhar na promoção da autonomia e do empoderamento das pessoas com deficiência, desde que desenvolvidos com princípios de anti-capacitismo.

Adicionalmente, os artigos **A04**, **A11**, **A13**, **A15**, **A20** e **A22** reforçam a urgência de se estabelecer uma regulação ética e governança algorítmica eficaz, que garanta a não reprodução de exclusões por sistemas de IA e assegure que os impactos dessas tecnologias sejam monitorados e mediados por políticas públicas consistentes e inclusivas.

5. Discussão

Os resultados desta RSL evidenciam que o capacitismo em sistemas de IA é uma preocupação crescente, embora ainda relativamente estigmatizada no escopo das pesquisas sobre ética e justiça algorítmica. O capacitismo apresenta uma sub-representação tanto em número de publicações quanto na diversidade metodológica e de localidade dos trabalhos encontrados.

A análise dos estudos revela que o capacitismo em IA se manifesta de múltiplas formas: desde a ausência ou sub-representação de pessoas com deficiência em *datasets* amplamente utilizados [Mukherjee et al. 2023, Manzoor et al. 2024], até a reprodução de estereótipos e representações estigmatizantes nos *outputs* gerados por modelos de linguagem [Park et al. 2025, Venkit et al. 2022, Narayanan Venkit et al. 2023]. Tal cenário levanta questões éticas fundamentais sobre os processos de coleta, anotação e modelagem de dados.

A concentração dos estudos em contextos de países desenvolvidos, como Estados Unidos e Europa, destaca um viés geográfico que limita a aplicabilidade dos resultados a realidades mais diversas. Poucos trabalhos abordam contextos do Sul Global, principalmente em termos de idiomas, ou apresentam reflexões interseccionais que considerem, por exemplo, deficiência combinada a outras categorias como pobreza, etnia ou acesso limitado à educação e tecnologia [Glazko et al. 2023, Glazko et al. 2025]. Isso aponta para a urgência de ampliar o escopo de localidade das pesquisas.

Outro ponto importante diz respeito às abordagens metodológicas adotadas. Grande parte dos trabalhos recorre a auditorias técnicas e análises quantitativas do viés algorítmico, utilizando métricas de toxicidade, polaridade de sentimento e analogias semânticas [Li et al. 2024, Glazko et al. 2024, Narayanan Venkit 2023]. Estudos com abordagens qualitativas, como autoetnografias e relatos de experiência de pessoas neurodivergentes [Glazko et al. 2023, Glazko et al. 2025], têm se mostrado promissores ao expor os efeitos concretos das tecnologias na vida cotidiana de pessoas com deficiência.

Por fim, os estudos convergem na defesa de uma formação ética e crítica para profissionais da área de IA, com ênfase na construção de *datasets* mais diversos e na criação de tecnologias que não apenas evitem discriminação, mas também promovam ativamente a autonomia e a dignidade de pessoas com deficiência [Urbina et al. 2025, Hutchinson et al. 2020, Tadimalla et al. 2024].

Dessa forma, esta RSL contribui para preencher uma lacuna importante na literatura, ao sistematizar as principais abordagens e desafios envolvidos no enfrentamento do capacitismo em IA, além de apontar caminhos para pesquisas futuras que sejam interdisciplinares, éticas e inclusivas.

5.1. Limitações do trabalho

Primeiramente, há o viés de seleção, pois os artigos analisados foram filtrados a partir de critérios como acesso aberto e indexação em bases específicas, o que pode ter deixado de fora estudos relevantes, especialmente aqueles publicados em periódicos pagos. Além disso, a revisão abrangeu publicações entre 2020 e o primeiro semestre de 2025, o que pode ter causado um viés temporal, excluindo artigos relevantes publicados posteriormente ou ainda não indexados.

5.2. Estudos futuros

A partir das lacunas identificadas na literatura, diversos caminhos se abrem para pesquisas futuras. Primeiramente, há uma demanda por investigações que explorem o impacto direto de sistemas de IA capacitistas em contextos educacionais e profissionais, especialmente a partir da perspectiva de pessoas com deficiência. Além disso, futuros estudos podem se concentrar na construção e validação de *datasets* específicos que representem

adequadamente a diversidade de deficiências, promovendo maior equidade nos modelos treinados. Outra frente de trabalho promissora consiste em investigar como o capacitismo algorítmico pode se interseccionar com outras formas de discriminação, como classe social ou raça, que são particularmente prementes no contexto brasileiro.

Outro eixo relevante é o desenvolvimento de métricas mais refinadas para a detecção automática de discurso capacitista, incorporando não apenas aspectos linguísticos explícitos, mas também elementos contextuais, culturais e subjetivos. Também se mostra promissor o avanço de abordagens interdisciplinares que integrem fundamentos da justiça social, estudos da deficiência e computação, contribuindo para o desenvolvimento de tecnologias inclusivas.

6. Conclusão

Esta Revisão Sistemática da Literatura revelou que, embora ainda inicial, a discussão sobre capacitismo em sistemas de IA vem ganhando espaço nas pesquisas acadêmicas. Os estudos analisados apontam que os sistemas de IA não são neutros, podendo reproduzir e amplificar preconceitos historicamente direcionados a pessoas com deficiência, seja pela ausência de representação nos dados, seja por vieses incorporados nos próprios modelos.

As metodologias encontradas variam entre análises técnicas de viés, em sua maioria em modelos de PLN, avaliações de *datasets*, abordagens qualitativas e reflexões ético-jurídicas, evidenciando a necessidade de um olhar multidisciplinar sobre o tema. Além disso, as perspectivas futuras destacam a importância de iniciativas que promovam a construção de tecnologias mais inclusivas, como a educação crítica em ciência de dados, a criação de bases de dados representativas e a formulação de políticas públicas voltadas à equidade algorítmica.

Por fim, as ações educacionais assumem papel central no enfrentamento do capacitismo algorítmico. A incorporação de conteúdos sobre justiça social, equidade e deficiência nos currículos de ciência de dados e áreas afins pode sensibilizar futuros profissionais para os impactos sociais de suas decisões técnicas. Além disso, a participação de pessoas com deficiência na construção e avaliação de sistemas amplia a consciência ética e a capacidade crítica dos estudantes, incentivando práticas de desenvolvimento mais justas e representativas. Ao promover uma educação voltada à inclusão e ao pensamento interdisciplinar, cria-se uma base sólida para que tecnologias de IA não apenas evitem reproduzir preconceitos, mas também atuem como agentes efetivos de transformação social.

Referências

- Angwin, J., Larson, J., Mattu, S., and Kirchner, L. (2016). Machine bias. *ProPublica*.
- Barocas, S. and Selbst, A. D. (2016). Big data's disparate impact. *California Law Review*, 104(3):671–732.
- Binns, R. and Kirkham, R. (2021). How could equality and data protection law shape AI fairness for people with disabilities? *CoRR*, abs/2107.05704.
- Bolukbasi, T., Chang, K.-W., Zou, J. Y., Saligrama, V., and Kalai, A. T. (2016). Man is to Computer Programmer as Woman is to Homemaker? Debiasing Word Embeddings. In *Advances in Neural Information Processing Systems (NeurIPS)*.

- Buyl, M., Cociancig, C., Frattone, C., and Roekens, N. (2022). Tackling algorithmic disability discrimination in the hiring process: An ethical, legal and technical analysis. page 1071–1082.
- Carrera-Rivera, A., Ochoa, W., Larrinaga, F., and Lasa, G. (2022). How-to conduct a systematic literature review: A quick guide for computer science research. *MethodsX*, 9:101895.
- Dutra, T. C., Maschio, E., and Gasparini, I. (2023). Pensar e lavar: Processo de desenvolvimento e avaliação de um jogo digital educacional para promover o pensamento computacional para crianças neurotípicas e com deficiência intelectual. *Revista Brasileira de Informática na Educação*, 31:659–690.
- Elglaly, Y. N. and Liu, Y. (2023). Promoting machine learning fairness education through active learning and reflective practices. *SIGCSE Bull.*, 55(3):4–6.
- Eubanks, V. (2018). *Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor*. St. Martin’s Press.
- Glazko, K., Cha, J., Lewis, A., Kosa, B., Wimer, B. L., Zheng, A., Zheng, Y., and Mankoff, J. (2025). Autoethnographic Insights from Neurodivergent GAI “Power Users”. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, CHI ’25, New York, NY, USA. Association for Computing Machinery.
- Glazko, K., Mohammed, Y., Kosa, B., Potluri, V., and Mankoff, J. (2024). Identifying and Improving Disability Bias in GPT-Based Resume Screening. page 687–700.
- Glazko, K. S., Yamagami, M., Desai, A., Mack, K. A., Potluri, V., Xu, X., and Mankoff, J. (2023). An autoethnographic case study of generative artificial intelligence’s utility for accessibility. In *Proceedings of the 25th International ACM SIGACCESS Conference on Computers and Accessibility*, ASSETS ’23, New York, NY, USA. Association for Computing Machinery.
- Griffin, P., Peters, M. L., and Smith, R. M. (2007). Ableism curriculum design. In Adams, M., Bell, L. A., and Griffin, P., editors, *Teaching for Diversity and Social Justice*, page 24. Routledge, 2nd edition.
- Hassan, S., Huenerfauth, M., and Alm, C. O. (2021). Unpacking the interdependent systems of discrimination: Ableist bias in NLP systems through an intersectional lens. In Moens, M.-F., Huang, X., Specia, L., and Yih, S. W.-t., editors, *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 3116–3123, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Herold, B., Waller, J., and Kushalnagar, R. (2022). Applying the Stereotype Content Model to Assess Disability Bias in Popular Pre-Trained NLP Models Underlying AI-Based Assistive Technologies. *Proceedings of the AAAI Conference on Artificial Intelligence*, 36(10):10750–10757.
- Hutchinson, B., Prabhakaran, V., Denton, E., Webster, K., Zhong, Y., and Denuyl, S. (2020). Social biases in NLP models as barriers for persons with disabilities. pages 5491–5501.
- IBGE, M. (2023). Brasil tem 18,6 milhões de pessoas com deficiência, indica pesquisa divulgada pelo IBGE e MDHC. <https://www.gov.br/mdh/pt->

- br/assuntos/noticias/2023/julho/brasil-tem-18-6-milhoes-de-pessoas-com-deficiencia-indica-pesquisa-divulgada-pelo-ibge-e-mdhc. Acesso em: 13 jan. 2025.
- Jacques, E., Sacramento, C., Gouveia, Y., Silva, W., Barros, Y., and Ferreira, S. (2025). Preservação da memória com acessibilidade digital: um plugin para descrição de imagens com IA generativa. In *Anais Estendidos do XXI Simpósio Brasileiro de Sistemas de Informação*, pages 157–161, Porto Alegre, RS, Brasil. SBC.
- Krupiy, T. T. and Scheinin, M. (2023). Disability Discrimination in the Digital Realm: How the ICRPD Applies to Artificial Intelligence Decision-Making Processes and Helps in Determining the State of International Human Rights Law. *Human Rights Law Review*, 23(3):ngad019.
- Li, R., Kamaraj, A., Ma, J., and Ebling, S. (2024). Decoding ableism in large language models: An intersectional approach. pages 232–249.
- Lima, L., Lima, L., Riquelme, M., and Ricarte, D. (2024). Uma revisão sistemática das técnicas de justiça algorítmica para diagnóstico radiológico: Avanços, desafios e perspectivas futuras. In *Anais Estendidos do XXIV Simpósio Brasileiro de Computação Aplicada à Saúde*, pages 37–42, Porto Alegre, RS, Brasil. SBC.
- Manzoor, R., Hussain, W., and Anjum, M. L. (2024). Out of dataset, out of algorithm, out of mind: a critical evaluation of AI bias against disabled people. *AI Soc.*, 40(5):3941–3951.
- Marks, M. (2020). Algorithmic disability discrimination. *Journal of Ethics and Information Technology*, 22(3):345–355.
- Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K., and Galstyan, A. (2022). A survey on bias and fairness in machine learning.
- Mondal, I., Kaur, S., Bali, K., Vashistha, A., and Swaminathan, M. (2022). “#DisabledOnIndianTwitter” : A dataset towards understanding the expression of people with disabilities on Indian Twitter. In He, Y., Ji, H., Li, S., Liu, Y., and Chang, C.-H., editors, *Findings of the Association for Computational Linguistics: ACL-IJCNLP 2022*, pages 375–386, Online only. Association for Computational Linguistics.
- Morr, C. E., Kundi, B., Mobeen, F., Taleghani, S., El-Lahib, Y., and Gorman, R. (2024). AI and disability: A systematic scoping review. *Health Informatics Journal*, 30(3):14604582241285743. PMID: 39287175.
- Moss, H. (2021). Screened out onscreen: Disability discrimination, hiring bias, and artificial intelligence. *Denver Law Review*, 98(4):2021. Posted: 18 Aug 2021.
- Mukherjee, A., Raj, C., Zhu, Z., and Anastasopoulos, A. (2023). Global Voices, local biases: Socio-cultural prejudices across languages. In Bouamor, H., Pino, J., and Bali, K., editors, *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 15828–15845, Singapore. Association for Computational Linguistics.
- N., F., N., L., F., L., R., M., Budaruiche, R., and Cunha, V. (2025). Tecnologia assistiva: O papel da inteligência artificial em sistemas embarcados para acessibilidade. In *Anais do XVII Encontro Unificado de Computação do Piauí*, pages 205–210, Porto Alegre, RS, Brasil. SBC.

- Narayanan Venkit, P. (2023). Towards a Holistic Approach: Understanding Sociodemographic Biases in NLP Models using an Interdisciplinary Lens. page 1004–1005.
- Narayanan Venkit, P., Srinath, M., and Wilson, S. (2023). Automated ableism: An exploration of explicit disability biases in sentiment and toxicity analysis models. In Ovalle, A., Chang, K.-W., Mehrabi, N., Pruksachatkun, Y., Galystan, A., Dhamala, J., Verma, A., Cao, T., Kumar, A., and Gupta, R., editors, *Proceedings of the 3rd Workshop on Trustworthy Natural Language Processing (TrustNLP 2023)*, pages 26–34, Toronto, Canada. Association for Computational Linguistics.
- Newman, P., Opdahl, T., Liu, Y., Wehrwein, S., and Elglaly, Y. N. (2024). Crafting disability fairness learning in data science: A student-centric pedagogical approach. In *Proceedings of the 55th ACM Technical Symposium on Computer Science Education V. 1, SIGCSE 2024*, page 944–950, New York, NY, USA. Association for Computing Machinery.
- O’Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown Publishing Group.
- Packin, N. G. (2021). Disability Discrimination Using AI Systems, Social Media and Digital Platforms: Can We Disable Digital Bias? *Journal of International and Comparative Law*, 8(2):487. Posted: 11 Jan 2021; Last revised: 30 Jan 2023.
- Pagano, T. P., Loureiro, R. B., Lisboa, F. V. N., Cruz, G. O. R., Peixoto, R. M., de Sousa Guimarães, G. A., dos Santos, L. L., Araujo, M. M., Cruz, M., de Oliveira, E. L. S., Winkler, I., and Nascimento, E. G. S. (2022). Bias and unfairness in machine learning models: a systematic literature review.
- Page, M. J., McKenzie, J. E., Bossuyt, P. M., Boutron, I., Hoffmann, T. C., Mulrow, C. D., Shamseer, L., Tetzlaff, J. M., Akl, E. A., Brennan, S. E., Chou, R., Glanville, J., Grimshaw, J. M., Hróbjartsson, A., Lalu, M. M., Li, T., Loder, E. W., Mayo-Wilson, E., McDonald, S., McGuinness, L. A., Stewart, L. A., Thomas, J., Tricco, A. C., Welch, V. A., Whiting, P., and Moher, D. (2021). The PRISMA 2020 statement: an updated guideline for reporting systematic reviews. *BMJ*, 372.
- Park, S., Min, A., Beltran, J. A., and Hayes, G. R. (2025). ”As an Autistic Person Myself:”The Bias Paradox Around Autism in LLMs. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems, CHI ’25*, New York, NY, USA. Association for Computing Machinery.
- Rodrigues, A. S., Costa, V. K., Cardoso, R. C., Machado, M. B., and Tavares, T. A. (2018). A Systematic Mapping Study about the Evaluation in Human-Computer Interaction of Assistive Technology Focused on People with Motor Disability. *iSys - Brazilian Journal of Information Systems*, 11(3):90–126.
- Santos, F., Gomes, L. A., Barros, E., Pinheiro, F., Silva, L., Nascimento, M., and Maia, P. (2017). Desafios e Experiências no Ensino de Programação Java através de Educação a Distância para Pessoas com Deficiência. In *Anais do XXIII Workshop de Informática na Escola*, pages 1109–1118, Porto Alegre, RS, Brasil. SBC.
- Shew, A. (2020). Ableism, Technoableism, and Future AI. *IEEE Technology and Society Magazine*, 39(1):40–85.

- Starke, C., Baleis, J., Keller, B., and Marcinkowski, F. (2022). Fairness perceptions of algorithmic decision-making: A systematic review of the empirical literature. *Big Data & Society*, 9(2):20539517221115189.
- Tadimalla, S. Y., Figard, R., and Song, Y. (2024). WIP: Moving from Accessibility to Anti-Ableism through the Explication of Disability in the AI Ecosystem. In *2024 IEEE Frontiers in Education Conference (FIE)*, pages 1–5.
- Urbina, J. T., Vu, P. D., and Nguyen, M. V. (2025). Disability Ethics and Education in the Age of Artificial Intelligence: Identifying Ability Bias in ChatGPT and Gemini. *Archives of Physical Medicine and Rehabilitation*, 106(1):14–19.
- Venkit, P. N., Srinath, M., and Wilson, S. (2022). A study of implicit bias in pretrained language models against people with disabilities. In Calzolari, N., Huang, C.-R., Kim, H., Pustejovsky, J., Wanner, L., Choi, K.-S., Ryu, P.-M., Chen, H.-H., Donatelli, L., Ji, H., Kurohashi, S., Paggio, P., Xue, N., Kim, S., Hahm, Y., He, Z., Lee, T. K., Santus, E., Bond, F., and Na, S.-H., editors, *Proceedings of the 29th International Conference on Computational Linguistics*, pages 1324–1332, Gyeongju, Republic of Korea. International Committee on Computational Linguistics.
- Whittaker, M., Alper, M., Bennett, C. L., Hendren, S., Kaziunas, E., Mills, M., and West, S. M. (2019). Disability, Bias, and AI. Technical report, AI Now Institute.
- Wu, G. and Ebling, S. (2024). Investigating ableism in LLMs through multi-turn conversation. In Dementieva, D., Ignat, O., Jin, Z., Mihalcea, R., Piatti, G., Tetreault, J., Wilson, S., and Zhao, J., editors, *Proceedings of the Third Workshop on NLP for Positive Impact*, pages 202–210, Miami, Florida, USA. Association for Computational Linguistics.