

BORSODI Zsófia

Eötvös Loránd University, Faculty of Education and Psychology, Doctoral School of Education
zsolia.borsodi@gmail.com
<https://orcid.org/0009-0005-7504-1768>

SZENCZI Beáta

Eötvös Loránd University, Bárczi Gusztáv Faculty of Special Education
szenczi.beata@barczi.elte.hu
<https://orcid.org/0000-0003-4709-8810>

SZEKERES Ágota

Eötvös Loránd University, Bárczi Gusztáv Faculty of Special Education
Hungarian University of Agriculture and Life Sciences, Institute of Education
HAS-ELTE Autism and Intellectual Disability in Education Research Group
szekeres.agota@barczi.elte.hu
<https://orcid.org/0000-0001-7120-9096>

VIRÁNYI Anita

Eötvös Loránd University, Bárczi Gusztáv Faculty of Special Education
HAS-ELTE Autism and Intellectual Disability in Education Research Group
viranyianita@barczi.elte.hu
<https://orcid.org/0000-0002-4601-4876>

<https://doi.org/10.31287/FT.en.2026.1.1>

Article published: 15th June, 2026

“Nothing About Us – Without Algorithms?”

A Theoretical Synthesis of the Contradictions in the Relationship Between Artificial Intelligence and Disability

Abstract

Artificial intelligence (AI) is one of the most transformative forces of the 21st century; it reshapes industries, economies, governmental systems, and the functioning of societies at an unprecedented pace. The tools of AI, including large language models (LLMs), are reshaping the frameworks of human communication, knowledge transmission, and social participation. From the perspective of disability studies, these technologies exhibit a dual nature: they offer access and new instruments to support inclusion, while also carrying the risk of exclusion. These systems are therefore both technological and normative constructions:

the linguistic patterns and cultural assumptions embedded in them determine who appears in the digital space and how they are represented. Focusing on this duality, the study examines how the issues of fairness, autonomy, transparency, and accessibility emerge in the development and application of large language models, with a specific focus on people with disabilities.

Keywords: artificial intelligence, ethics, disability, fairness, large language models, LLM

„Semmit rólunk – algoritmusok nélkül?”

A mesterséges intelligencia és a fogyatékoság kapcsolatában megjelenő ellentmondások elméleti szintézise

Absztrakt

A mesterséges intelligencia (MI) a 21. század egyik legátalakítóbb ereje; példátlan ütemben formálja át az iparágakat, a gazdaságokat, a kormányzati rendszereket és a társadalmak működését. Az MI eszközei, köztük a nagy nyelvi modellek (Large Language Models, LLM-ek) újraformálják az emberi kommunikációt, a tudásközvetítést és a társadalmi részvétel kereteit. A fogyatékoságtudomány szemszögéből a technológiák kettős természetét mutatnak: egyszerre kínálnak hozzáférést és új eszközöket az inklúzió támogatására, ugyanakkor a kirekesztés kockázatát is magukban hordozzák. Ezek a rendszerek tehát egyszerre technológiai és normatív konstrukciók is: a bennük rejlő nyelvi mintázatok és kulturális előfeltevések meghatározzák, kik és hogyan jelennek meg a digitális térben. A tanulmány erre a kettősségre fókuszálva vizsgálja, hogy miként jelennek meg a méltányosság, autonómia, átláthatóság és hozzáférhetőség kérdései a nagy nyelvi modellek fejlesztésében és alkalmazásában, különös tekintettel a fogyatékosággal élő emberek csoportjára.

Kulcsszavak: mesterséges intelligencia, etika, fogyatékoság, méltányosság, nagy nyelvi modellek, LLM



“Nothing About Us – Without Algorithms?”

A Theoretical Synthesis of the Contradictions in the Relationship Between Artificial Intelligence and Disability

1. Introduction

Artificial intelligence, particularly generative systems, are increasingly redefining human interactions in both digital and physical spaces. Analyses by international policy bodies have noted that AI is increasingly embedded in public services and administrative practices, where it can enhance productivity, tailor services to citizen needs, and support decision-making, but also raises concerns about transparency, accountability and equitable access (OECD, 2025). Natural language models such as ChatGPT are capable of producing human-like responses, conducting natural dialogues, and generating coherent, creative texts (Dwivedi et al., 2023). Large language models are highly sensitive technologies that mediate the linguistic and

cognitive dimensions of human communication and, in doing so—often in hidden ways—also carry values, cultural assumptions, and social norms (Bender et al., 2021; Navigli et al., 2023). Linguistic and cultural biases embedded in datasets may render the experiences of people with disabilities invisible or reinforce stereotypes, thus reproducing social exclusion in digital form.

Floridi's (2013) "information ethics" asserts that the handling of data and models is always a moral decision: every act of data processing is a choice of representation, determining whose story enters the digital canon and who is left out. Information systems are consequently not value-neutral; the "infosphere" conveys the values and norms of human contexts.

The application of AI in education and special education further illustrates these dynamics. An increasing number of studies examine the opportunities and challenges that AI-based tools pose for inclusive pedagogy and learning innovation. AI in education is not entirely new, but its application has entered a new dimension with the emergence of next-generation generative models (e.g., ChatGPT, Synthesia, NotebookLM, DALL-E, Elicit). The integration of AI tools into educational practice may, even in the short term, lead to the formation of new learning settings supported by adaptive learning management systems, learning analytics procedures that analyse student performance, and systems that provide automated feedback (Horváth, 2023). These technological solutions initiate institutional-organisational transformations as well as methodological changes, which raise ethical and social issues at the level of educational practice (Bokor, 2023).

There is a growing interest in AI tools within the field of special education. However, the practical implementation is shaped by motivational, competency-related, and ethical considerations alike (Mező, 2025). At the same time, the application of AI in education is often described in an instrument-oriented manner in the literature and is not always accompanied by appropriate methodological or inclusive conceptual frameworks (Borsodi & Virányi, 2024).

The post-model of disability provides a critical lens for examining digital infrastructures as cultural environments rather than merely technical artefacts. If disability is understood as discursively constructed, shaped by social and environmental contexts, then AI systems—particularly language models—become active participants in constructing the meaning of disability in the digital sphere. Thus, AI ethics and disability theory intersect in their shared concern with power, representation, and the conditions of participation.

This paper offers a theoretical synthesis at the intersection of disability studies and AI ethics, drawing on interdisciplinary scholarship in disability theory, algorithmic fairness research, and AI governance. It aims to:

- 1) analyse how the post-model of disability reframes the ethical evaluation of large language models;
- 2) identify structural tensions between statistical AI systems and the heterogeneity of disability;
- 3) examine how algorithmic bias, representational homogenisation, and technoableist assumptions shape the digital construction of disability;
- 4) outline normative and technical approaches toward disability-inclusive AI development.

The evidentiary basis consists primarily of peer-reviewed academic publications, complemented by publicly available policy documents and recent scholarly preprints. The literature was selected according to its thematic relevance to these analytical objectives and identified through focused searches in major international academic databases (including Scopus and Web of Science Core Collection, as well as discipline-specific repositories such as the ACM Digital Library and IEEE Xplore), complemented by citation tracing of influential works in disability studies and AI ethics. The corpus integrates foundational theoretical contributions with contemporary research published between 2019 and 2025, reflecting the period during which large language models and generative AI systems became widely deployed and disability-focused critiques of these algorithms intensified. The analysis moves from key developments in disability theory to the structural characteristics of AI systems, before examining manifestations of algorithmic bias and discussing possible governance and design-level responses.

2. The Post-Model of Disability and the New Dimensions of the Technological Environment

The conceptualisation of disability shifted in the second half of the twentieth century from the individualising perspective of the medical model to the collective approach of the social model (Barnes, 1998; Oliver, 1990). This shift redirected attention from individual impairment to the role of social and environmental barriers and is reflected in the World Health Organization's International Classification of Functioning, Disability and Health (ICF), which adopts a biopsychosocial framework integrating functional, social and environmental dimensions (World Health Organization, 2001).

Building on the conceptual contributions of the social model, the human rights model of disability emerged through disability-led civil rights movements and international human rights advocacy and was institutionally articulated at the normative level in the UN Convention on the Rights of Persons with Disabilities (UNCRPD; United Nations, 2006). The UNCRPD declares that people with disabilities must be ensured access to all areas of social life, including information and communication technologies and systems. This logic may also be extended to the information environment: the informational sphere can itself be limiting if it excludes or misrepresents certain groups. Accordingly, the accessibility or biased operation of algorithmic systems constitutes a technological issue as well as a human rights concern.

In Hungarian disability studies, the social model of disability was expanded by Könczei and Hernádi (2011) toward the post-model, which understands disability as an open and discursive cultural construction. This conceptualisation is consistent with developments in international disability studies that have moved beyond a rigid opposition between medical and social models. Relational approaches have examined the cultural and representational dimensions of disability (Shakespeare, 2014; Garland-Thomson, 2002), while intersectional analyses have explored how



disability interacts with gender, race, and socio-economic position within broader social structures (Erevelles, 2011).

From this standpoint, the examination of the technological environment also reveals that the norms, categorisation logics, and representations that operate in AI systems reflect how the concept of “disability” is constructed in information society (Foley & Melese, 2025). The analytical shift from social space to digital infrastructure allows disability theory to inform the interpretation of machine learning architectures and their socio-cultural effects.

Contemporary AI ethics scholarship likewise emphasises that AI systems should be understood as socio-technical constructs embedded in institutional, cultural, and regulatory contexts rather than purely technical artefacts (Jobin et al., 2019; Stahl, 2021). From this perspective, disability theory and AI ethics intersect in their shared concern with fairness, representation, and participation. If disability emerges within structured environments, then algorithmic systems that organise communication may also influence how disability is represented, interpreted, and evaluated in digital spaces. Designing ethical and inclusive AI systems therefore requires approaches that foreground lived experience, user needs, and contextual understanding, principles often associated with human-centred design (Jain, 2025).

Automated assistive technologies for communication, speech recognition, and image-based support may create new forms of participation, while algorithms trained on incomplete datasets may reproduce social prejudices and the digital-informational forms of exclusion (Nazer et al., 2023; Soldatic et al., 2025). Recent reviews indicate that AI research concerning disability remains strongly influenced by the medical model and often reflects ableist assumptions (El Morr et al., 2024).

The concept of ableism (Campbell, 2009) is based on the implicit assumption that the “able” body and mind constitute the normative state, thus framing disability through the lens of deficit or lack. In the discourse on the relationship between technology and people with disabilities, the notion of technobleism has become increasingly prominent: a technology-based form of ableism that makes visible the often hidden prejudices and normative expectations behind the seeming inclusivity of technological advancement. Ashley Shew (2020) introduced the term to describe the widespread assumption that technology is inherently “good” for people with disabilities and is both capable and obliged to “fix” disability.

The phenomenon of technobleism can thus be described as dual in nature: while technology operates with the promise of “empowerment,” it often continues to rest on the implicit assumption that the normative body and mind are the desirable state to which people with disabilities must adapt. Shew (2020) presents technobleist thinking as more than a flawed application of technology, describing it as a structural ideology that ties disability to an individualised biological deficit, disregarding the role of social and environmental barriers.

The technological support of “independence” and “autonomous living” often further reinforces this ideology by obscuring the naturalness of mutually dependent human relations—that assistance is not deviance, but a fundamental part of human existence. According to Shew (2020), this does not mean that the critique of technobleism entails an anti-technology stance. Quite the contrary: it advocates the creation of a radically participatory technological vision in which people with disabilities go beyond being “users” or test subjects, and instead take part as active

creators and critics of technological innovations (Shew, 2020). This approach is likewise embodied in the Crip technoscience movement, which places the lived experience and knowledge of people with disabilities at the centre of technological development (Hamraie & Fritsch, 2019).

Taken together, these theoretical developments frame disability as relational, context-dependent, and shaped by cultural and institutional environments. The post-model thus shifts attention from individual deficit to the structures that organise interaction and meaning. This conceptual orientation provides the basis for examining how digital systems participate in the construction and portrayal of disability.

3. The Structural Contradictions of Artificial Intelligence

The structural characteristics of contemporary AI systems generate several tensions when examined from the perspective of disability studies. The following sections analyse these dynamics across four interrelated dimensions: statistical homogenisation in machine learning systems, representational bias in AI-generated narratives, intersectional forms of exclusion, and the cognitive and institutional dynamics that shape human–AI interaction.

3.1. Statistical homogenisation and algorithmic monoculture

Large language models are statistical technologies that model the probability distributions of linguistic patterns, visual depictions, and behavioural data. Their outputs therefore mirror the most frequently occurring (“average”) structures in the training data (Bender et al., 2021; Bommasani et al., 2022; Jurafsky & Martin, 2025). This gives rise to the problem of *homogenisation*: although AI technologies promise new forms of access and participation for people with disabilities, the logic of their operation tends to smooth differences out and reinforce typical patterns.

Among protected minority groups, people with disabilities form the most heterogeneous population (McGuire, 2020). The concept of disability is itself complex and dynamic, encompassing a broad spectrum of sensory, physical, cognitive, and psychosocial conditions (World Health Organization, 2001; United Nations, 2006; World Health Organization & World Bank, 2011). This complexity poses challenges for machine learning models, which often rely on uniform parameters and are therefore especially prone to disregarding differences between groups (Trewin et al., 2019; Mehrabi et al., 2021). Consequently, the diversity that the post-model recognises as a fundamental characteristic of human experience may gradually disappear through the statistical homogenisation inherent in machine learning systems.

The literature refers to this phenomenon as algorithmic monoculture, which occurs when widely used models lead to converging outputs (Kleinberg & Raghavan, 2021; Bommasani et al., 2022). The consequence of conformity to a statistical average is that rarer linguistic, bodily, or cognitive patterns become less visible in model outputs.



For people with disabilities, this homogenisation may appear as implicit, “silent” exclusion. Large language models do not necessarily generate overtly offensive content but reproduce subtle, paternalistic schemas that portray people with disabilities more as objects (sources of inspiration or moral lessons) than as autonomous agents (Trewin et al., 2019).

It is important to emphasise that the dangers of homogenisation are not only cultural but may also have clinical consequences. When AI is used to process medical or care documentation, standardisation aligned to an “average” may obscure individual differences and context-dependent details, ultimately leading to diagnostic bias and inequality in access to care (Obermeyer et al., 2019; Ghassemi et al., 2021).

3.2. Representational bias and ableist narratives

The statistical tendencies described above extend beyond internal model dynamics and shape the cultural and representational dimensions of AI-generated content. In the case of large language models, the question of how people with disabilities are represented in datasets and algorithms—and how AI shapes the social narrative about them—demands careful consideration. According to the digital interpretation of the “nothing about us without us” principle (Charlton, 1998), ethical technological development is inconceivable without the participation of those concerned.

The problem of homogenisation is further complicated by the fact that AI is not only a technological but also an interpretive medium. Language models such as ChatGPT or Gemini operate in the realm of communication and meaning-making and therefore become mediators of cultural norms. Since their learning processes rely on biased datasets and linguistic assumptions, the resulting outputs are prone to *reproducing social prejudices* (Hu et al., 2025; Nazer et al., 2023). In relation to disability, this may mean that instead of nuanced understandings of human differences, AI systems reinforce a simplified dichotomy of normality and deficiency.

According to the New York City Bar Association (2025), portrayals of disability in generative systems often follow stereotypical patterns (e.g., heroism, dependency, pity) and provide visually homogenised portrayals (most often white men with visible physical disabilities), which—through narrative homogenisation—contribute to the cultural marginalisation of people with disabilities.

Anglo-centrism is a practical manifestation of linguistic and cultural dominance in AI-driven systems. Most large language models learn predominantly from English-language sources, embedding Western—and more specifically Eurocentric—cultural values and discourses as defaults (Bender & Friedman, 2018; Joshi et al., 2020). This data dependency particularly disadvantages those communities—including people with disabilities—whose experiences and expressive forms rarely appear in English-language corpora (Blasi et al., 2022). Interestingly, models trained in non-English languages are not exempt from distortion either: in translation-based systems, the prejudices embedded in the source language are transferred and may further distort outputs in the target language (Prates et al., 2020).

The corpora used to train LLMs and other machine learning systems are constructed largely from non-accessible digital content, mirroring the structural absence of accessibility in algorithms as well (WebAIM, 2023). According to WebAIM's



(2023) annual report, the homepages of 97% of the world's one million most visited websites did not comply with the basic criteria of the Web Content Accessibility Guidelines (WCAG; World Wide Web Consortium [W3C], 2023). Structural exclusion in the digital sphere is thus reproduced at the level of model training.

3.3. Intersectionality and structural exclusion

Representational absences and statistical distortions do not remain at the level of internal model logic. Numerous studies (AI Now Institute, 2019; Glazko et al., 2024; Kamruzzaman & Kim, 2025; Tilmes, 2022) have demonstrated that AI—and large language models in particular—reproduce systemic biases against people with disabilities.

In AI-supported labour-market selection systems, applicants who openly disclosed their disability often received lower scores, even when they had identical qualifications, professional experience, gender, and ethnic backgrounds to non-disabled applicants (Kamruzzaman & Kim, 2025). This indicates that LLMs and other AI systems may favour “normative” characteristics aligned with ableist social norms, thereby implicitly disadvantaging people with disabilities.

Bias often intensifies intersectionally, when disability intersects with other marginalised identities—such as gender, ethnicity, or social class (Phutane et al., 2025). To identify disability-related biases in AI systems, Phutane and colleagues (2025) introduce the ABLEist framework, which identifies four main forms of harm associated with ableism: ableism, inspirational narrative, superhumanisation, and tokenism. Alongside ableism, inspirational narratives position people with disabilities as tools of motivation for others, often wrapping their biographies in emotionally charged praise (e.g., “if they can do it, you can too”) (Young, 2014). Related to this is the concept of superhumanisation which depicts people with disabilities as superheroes who overcome their bodily or cognitive “limitations” with extraordinary strength (Garland-Thomson, 2002). Although seemingly uplifting, this narrative imposes expectations and obscures structural barriers. Finally, tokenism (Ahmed, 2012) creates the appearance of inclusion without meaningful participation. Intersectional harms—especially tokenism and superhumanisation—disproportionately affect disabled women and socially disadvantaged disabled individuals (Phutane et al., 2025).

3.4. Cognitive, relational and governance dimensions

Beyond representational and structural biases, AI systems also influence how users perceive and interact with technological systems. One such contradiction lies between *independence and dependence*: while assistive technologies aim to enhance autonomy, they may simultaneously create new forms of technological dependence.

Research has demonstrated that although language models can improve human cognitive performance, they also weaken the accuracy of self-reflection and critical self-evaluation. Fernandes and colleagues (2025) examined how the use of LLMs influences logical reasoning performance and metacognitive accuracy. According to their findings, participants who used ChatGPT achieved significantly higher scores

on logical tasks from the Law School Admission Test, yet overestimated their own knowledge and accuracy, indicating metacognitive bias. This phenomenon relates to the Dunning–Kruger effect (Kruger & Dunning, 1999), whereby individuals with low performance tend to overestimate their abilities, while highly competent individuals tend to underestimate their knowledge. This metacognitive distortion stems from the fact that poor performance is often accompanied by the absence of the knowledge required for accurate self-assessment (Kruger & Dunning, 1999; Pennycook et al., 2017). According to Fernandes and colleagues (2025), the use of LLMs interestingly mitigated this distortion: the Dunning–Kruger effect weakened because AI support increased the performance of lower-skilled participants, while participants with higher AI literacy showed greater confidence but lower accuracy in self-evaluation. This indicates that AI-assisted systems may narrow performance gaps, yet at the same time increase overconfidence and reduce the capacity for self-reflection and metacognitive awareness (Buçinca et al., 2021).

Recent research highlights that the relationship between humans and AI may create a form of psychological dependence. Through constant interaction with AI systems—especially language models and voice-based services—users tend to attribute human qualities to technology, which enhances trust but reduces critical distance. According to Naseer, Ahmad and Chishti (2025), this process creates the “illusion of technological comfort,” while increasing cognitive dependence and decreasing reflective autonomy. AI-based decision support thus paradoxically both increases and limits autonomy. One of the main psychological drivers of the phenomenon is anthropomorphism: the attribution of human characteristics to machine systems. Analysing the structural model of trust in AI, Dang and Li (2026) concluded that the main components of user trust—competence, benevolence, and integrity—are constructed through anthropomorphic attributions. That is, the more human-like an AI system appears, the more likely users are to project human norms of reliability onto it, even when the decisions of the algorithm cannot be interpreted through such norms.

Marketing and consumer psychology research has revealed similar dynamics. According to Gomes, Lopes and Nogueira (2025), anthropomorphic AI models increase emotional attachment and brand loyalty, while also encouraging shifts in user behaviour toward psychological dependency, especially where the system creates the appearance of empathy or attentiveness.

Transparency is one of the greatest challenges for AI tools, including large language models, and has become, since the 2010s, a cornerstone of ethical AI discourse, together with the concept of accountability (Diakopoulos, 2016; Ananny & Crawford, 2018). The increasingly complex functioning of AI systems gradually undermines trust between developers, decision-makers, and users through the so-called “*black box problem*” (Burrell, 2016; Carabantes, 2020; Lipton, 2018). Diakopoulos (2016) defines the concept of algorithmic responsibility as maintaining a trace of human decisions throughout the entire process of development, deployment, and evaluation. Ananny and Crawford (2018) further nuance the concept: transparency alone does not ensure ethical operation if social context and power relations remain unacknowledged. The notion of human-centred transparency (Liao & Vaughan, 2023) builds on these insights. According to this, transparency does not mean the disclosure of code or data, but the support of user understanding:



ensuring that stakeholders—developers, users, decision-makers—comprehend how the model operates in ways aligned with their roles and competencies. Inclusive design approaches similarly emphasise that technological systems must account for human variability from the outset. In this framework, complexity and diversity are considered inherent features of system architecture and development processes (Treviranus, 2023).

In recent years, numerous technical and methodological initiatives have emerged to increase the transparency of AI systems. Examples include model cards (Mitchell et al., 2019) and datasheets for datasets (Gebru et al., 2021), which aim to make the development process and the origins, limitations, and purposes of the data used transparent. Additionally, the field of explainable AI (Doshi-Velez & Kim, 2017; Miller, 2019) seeks to ensure that AI decisions are accurate while remaining understandable and traceable from a human perspective.

From the standpoint of disability studies, all this is crucial because transparency and interpretability enable users with different abilities to understand and critically evaluate the functioning of AI. Mapping the terrain of algorithmic ethics, Mittelstadt and colleagues (2016) warn that transparency is not solely a technical matter: the lack of interpretability may reproduce social inequalities, as AI decisions often appear unquestionable to non-expert users.

Legal regulations also aim to reduce AI bias. For example, the European Union's regulation on AI establishes a common legal framework for the development and use of AI systems in Europe (European Parliament & Council of the European Union, 2024). In October 2025, Hungary adopted the first comprehensive piece of legislation aimed at implementing this regulation at the national level. Act LXXV of 2025, which entered into force on 1 December 2025, places emphasis on ensuring a balance between technological innovation and the protection of fundamental rights in the development and application of AI. The legislation stresses that the use of AI tools must not violate the principle of human dignity and requires special consideration for vulnerable groups. The Act also acknowledges the uncertainties arising from the rapid technological changes associated with AI applications: in the absence of practical experience with implementation, it deems necessary the issuance of unified guidance, recommendations, and opinions to support consistent legal practice (Act LXXV of 2025 on the Hungarian implementation of the European Union's regulation on artificial intelligence).

The structural characteristics of large-scale AI systems—statistical optimisation, dependence on training data, and the amplification of dominant patterns—create conditions under which minority perspectives risk marginalisation. These dynamics emerge from the operational logic of machine learning rather than from isolated technical errors. Understanding this structural dimension is essential for analysing how disability is represented in AI-mediated environments.

4. Technical Responses to Algorithmic Bias

For people with disabilities, AI thus offers opportunities yet carries risks. Language models often invisibly encode the social assumptions that interpret disability as deficit rather than identity.



Reducing AI bias is also an important issue in educational practice, especially when AI is used to assess learning outcomes or personalise learner support. According to empirical research by Demeter and Mező (2023), the majority of Hungarian special education teacher trainees are open to the application of AI; however, the management of algorithmic bias, data protection concerns, and ethical risks is a necessary condition for AI to create genuine pedagogical value for learners with special educational needs. Mező (2025) emphasises that, in addition to bias mitigation techniques, the creation of specialised databases and the involvement of users—including representatives of vulnerable groups—are indispensable for establishing fair AI systems in pedagogy.

The literature agrees that these systems require both value-based and technically fairness-aware approaches (Li et al., 2023; Gallegos et al., 2023). In recent years, increasing attention has been paid to technological models and algorithmic solutions designed to reduce bias, preserve diversity, and support inclusive participation.

Fairness is not a single algorithm or metric, but a set of design decisions and technical procedures layered together. The classical literature distinguishes three complementary approaches: preprocessing, in-processing, and post-processing.



1) Preprocessing

Interventions at the dataset level constitute the first line of defence against bias. The aim of preprocessing is to mitigate distortions in the distribution of input training data before the model learns from them. Classic methods include reweighting and resampling, which address disproportionalities between groups by adjusting their relative representation (Kamiran & Calders, 2012).

In educational AI systems, besides reweighting, synthetic data augmentation is often used to increase the representation of underrepresented learner groups (such as students with special educational needs or minority language speakers) in the model training space. The “counterfactual data augmentation” technique described by Li and colleagues (2023), for example, inserts artificially generated samples into training datasets that model the characteristics of minority groups, thereby balancing training distributions. Through such procedures, the model does not evaluate individuals against a “normative” learner profile but learns from a broader and more realistic diversity.

2) In-processing

At the level of model training, adversarial debiasing has become increasingly widespread in recent years. This approach creates a balance between two competing components of the model: while one component optimises the primary task, the other seeks to identify and neutralise sensitive attributes (e.g., gender, language, disability). The MinDiff algorithm developed by Google Research, for instance, applies a loss function that explicitly penalises the model if it produces different error rates for different groups (Google Developers, 2024). Such solutions are well suited to educational LLM systems, where equal evaluation is essential when processing student feedback, essays, or speech patterns.

3) Post-processing

When the model is already built, post-processing modifies its outputs. The latest developments in post-processing include the integration of fairness-checking metrics into AI systems. According to the taxonomy presented by Gallegos and colleagues (2023), probabilistic fine-tuning and output regulation are increasingly common in modern LLMs, applying real-time statistical correction to balance model outputs. This means that the system automatically checks whether it prefers certain vocabularies, social groups, or cultural perspectives before presenting the text to the user.

Table 1. below summarises the main risks and opportunities identified in the previous sections, together with corresponding ethical and technical responses discussed in the literature.

TABLE 1. RISKS, OPPORTUNITIES AND ETHICAL RESPONSES IN DISABILITY-INCLUSIVE AI

Dimension	Risks	Opportunities	Ethical / Technical Responses
Data and modelling	Statistical homogenisation and underrepresentation of disability-related experiences in training data	AI-enabled assistive technologies and accessibility tools	Diversified datasets, dataset documentation, bias evaluation methods
Representation	Reproduction of ableist narratives and stereotypical portrayals of disability in datasets and outputs	Greater visibility of disability perspectives in digital environments	Participatory design and disability-sensitive evaluation frameworks
Intersectionality	Amplification of inequalities where disability intersects with gender, ethnicity or socio-economic disadvantage	Potential improvement of access to services and information	Intersectional bias auditing and inclusive governance approaches
Human–AI interaction	Cognitive overreliance, anthropomorphism, and reduced critical evaluation of automated outputs	Support for reasoning, communication and learning processes	AI literacy, explainability, and human-centred transparency
Governance and regulation	Limited transparency and accountability in complex AI systems	Regulatory frameworks may promote responsible AI development	Transparency requirements, auditing mechanisms, and AI governance frameworks

In pedagogical contexts, these technological solutions play a central role as AI generates content while shaping learning experiences. The inclusive educational AI systems examined by Kooli and Chakraoui (2025), for example, create new forms of participation for students with special educational needs through assistive technologies based on speech recognition and image processing, while built-in bias-filtering algorithms ensure that the system does not privilege “average” language-use patterns.



5. Conclusion

The relationship between AI and disability is both a technological and a social issue, shaped by the deep structures of digital systems together with the theoretical frameworks of disability studies. Databases, learning patterns, and informational environments that influence the functioning of large language models are not neutral: they reflect the cultural assumptions, norms, and asymmetries that also condition social perceptions of disability.

One of the most important intersections between the post-model of disability and the analysis of AI systems is that both direct attention to the complex interactions between environment and individual. The post-model emphasises that disability is not an internal deficit, but a social-spatial construction that emerges in different interactional and communicational situations. Extending this logic to the digital sphere reveals that AI operates with the limitations and preferences embedded in datasets and algorithmic procedures.

The statistical functioning of large language models is based on the amplification of frequent patterns and the backgrounding of rare ones, a mechanism that creates structurally disadvantageous conditions for minority groups. Due to the heterogeneity of people with disabilities, this effect is further intensified: linguistic, cognitive, or cultural experiences that are sparsely documented often fail to appear in the training data, leaving the models without an adequate basis for producing nuanced representations. This study has highlighted that algorithmic homogenisation, ableist representations, and Anglo- and Eurocentric data dependency are not merely side-effects but phenomena inherent in the internal logic of machine learning systems.

While the continually evolving technical procedures of bias mitigation may offer some advancement, they cannot substitute for the involvement of people with disabilities throughout the entire development process. Participatory auditing, critical design methodologies, and approaches grounded in community knowledge are tools capable of alleviating the shortcomings that arise from exclusively technical perspectives.

In practical terms, disability-inclusive AI development requires:

- 1) institutionalised participation of persons with disabilities in AI design and auditing processes;
- 2) the systematic integration of disability-sensitive fairness metrics in LLM evaluation;
- 3) the expansion of multilingual and accessibility-compliant datasets;
- 4) the embedding of disability studies perspectives into AI ethics guidelines and regulatory frameworks;
- 5) the development of AI literacy initiatives addressing metacognitive risks and technological dependence.

These directions point toward a model of AI development in which considerations of disability are integrated into the foundational stages of system design rather than treated as subsequent corrections. The dialogue between AI ethics and disability studies remains essential for shaping digital environments that support diversity and equitable participation.



References

- Ahmed, S. (2012). *On being included: Racism and diversity in institutional life*. Duke University Press. <https://doi.org/10.1215/9780822395324>
- AI Now Institute (2019). *Disability, bias, and AI: A report of the AI Now Disability & AI Working Group*. AI Now Institute. <https://ainowinstitute.org/publication/disabilitybiasai-2019>
- 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>
- Barnes, C. (1998). The social model of disability: A sociological phenomenon ignored by sociologists. In T. Shakespeare (Ed.), *The disability reader: Social science perspectives* (pp. 65–78). Cassell.
- Bender, E. M. & Friedman, B. (2018). Data statements for natural language processing: Toward mitigating system bias and enabling better science. *Transactions of the Association for Computational Linguistics*, 6, 587–604. https://doi.org/10.1162/tacl_a_00041
- Bender, E. M., Gebru, T., McMillan-Major, A. & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In M. C. Elish, W. Isaac & R. S. Zemel (Eds.), *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FACt '21)* (pp. 610–623). Association for Computing Machinery. <https://doi.org/10.1145/3442188.3445922>
- Blasi, D. E., Anastasopoulos, A. & Neubig, G. (2022). Systematic inequalities in language technology performance across the world's languages. In S. Muresan, P. Nakov & A. Villavicencio (Eds.), *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 5486–5505). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2022.acl-long.376>
- Bokor, T. (2023). A mesterséges intelligencia alkalmazása az oktatásban. *Vezetéstudomány*, 54(3), 114–129. <https://unipub.lib.uni-corvinus.hu/8710/>
- Bommasani, R., Hudson, D. A., Adeli, E., Altman, R., Arora, S., von Arx, S., Bernstein, M. S., Bohg, J., Bosselut, A., Brunskill, E., Brynjolfsson, E., Buch, S., Card, D., Castellon, R., Chatterji, N., Chen, A., Creel, K., Davis, J. Q., Demszky, D., ... Liang, P. (2022). On the opportunities and risks of foundation models. *Journal of Machine Learning Research*, 23(1), 1–214. <https://doi.org/10.48550/arXiv.2108.07258>
- Borsodi, Zs. & Virányi, A. (2024). Tanulás és technológia: a mesterséges intelligencia szerepe az oktatásban, különös tekintettel a gyógypedagógiai alkalmazásra. *Új Pedagógiai Szemle*, 74(11–12), 34–59. https://epa.oszk.hu/00000/00035/00255/pdf/EPA00035_upsz_2024_11-12_034-059.pdf
- Buçinca, Z., Malaya, M. B. & Gajos, K. Z. (2021). To trust or to think: Cognitive forcing functions can reduce overreliance on AI in AI-assisted decision-making. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW2), 1–21. <https://doi.org/10.1145/3449287>
- Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 1–12. <https://doi.org/10.1177/2053951715622512>
- Campbell, F. K. (2009). *Contours of ableism: The production of disability and abledness*. Palgrave Macmillan. <https://doi.org/10.1057/9780230245181>
- Carabantes, M. (2020). Black-box artificial intelligence: An epistemological and critical analysis. *AI & Society*, 35(2), 309–317. <https://doi.org/10.1007/s00146-019-00888-w>
- Charlton, J. I. (1998). *Nothing About Us Without Us: Disability Oppression and Empowerment*. University of California Press. <https://doi.org/10.1525/9780520925441>
- Dang, Q. & Li, G. (2026). Unveiling trust in AI: the interplay of antecedents, consequences, and cultural dynamics. *AI & Society*, 41, 669–692. <https://doi.org/10.1007/s00146-025-02477-6>
- Demeter, Zs. & Mező, K. (2023). A mesterséges intelligencia pedagógiai használatára vonatkozó hajlandóság vizsgálat a gyógypedagógus hallgatók körében. *Különleges Bánásmód*, 9(2), 31–45. <https://doi.org/10.18458/KB.2023.2.31>
- Diakopoulos, N. (2016). Accountability in algorithmic decision making. *Communications of the ACM*, 59(2), 56–62. <https://doi.org/10.1145/2844110>
- Doshi-Velez, F. & Kim, B. (2017). *Towards a rigorous science of interpretable machine learning*. arXiv. <https://doi.org/10.48550/arXiv.1702.08608>
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). "So what if ChatGPT wrote it?" Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>

- El Morr, C., Kundi, B., Mobeen, F., Taleghani, S., El-Lahib, Y. & Gorman, R. (2024). AI and disability: A systematic scoping review. *Health Informatics Journal*, 30(3). <https://doi.org/10.1177/14604582241285743>
- Erevelles, N. (2011). *Disability and difference in global contexts: Enabling a transformative body politic*. Palgrave Macmillan. <https://doi.org/10.1057/9781137001184>
- European Parliament and Council of the European Union (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act)*. Official Journal of the European Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>
- Fernandes, D., Villa, S., Nicholls, S., Haavisto, O., Buschek, D., Schmidt, A., Kosch, T., Shen, C. & Welsch, R. (2025). AI makes you smarter but none the wiser: The disconnect between performance and metacognition. *Computers in Human Behavior*, 158, 108779. <https://doi.org/10.1016/j.chb.2025.108779>
- Floridi, L. (2013). *The Ethics of Information*. Oxford University Press. <https://doi.org/10.1093/acprof:oso/9780199641321.001.0001>
- Foley, A. & Melese, F. (2025). Disabling AI: Power, exclusion, and disability. *British Journal of Sociology of Education*, 1–22. <https://doi.org/10.1080/01425692.2025.2519482>
- Gallegos, I. O., Rossi, R. A., Barrow, J., Tanjim, M., Kim, S., Deroncourt, F. & Ahmed, N. K. (2023). *Bias and fairness in large language models: A survey*. *Computational Linguistics*, 50(3), 1097–1179. https://doi.org/10.1162/coli_a_00524
- Garland-Thomson, R. (2002). *The politics of staring: Visual rhetorics of disability in popular photography*. In S. L. Snyder, B. J. Brueggemann & R. Garland-Thomson (Eds.), *Disability Studies: Enabling the Humanities* (pp. 56–75). Modern Language Association of America.
- Gebru, T., Morgenstern, J., Vecchione, B., Vaughan, J. W., Wallach, H., Daumé III, H. & Crawford, K. (2021). Datasheets for Datasets. *Communications of the ACM*, 64(12), 86–92. <https://doi.org/10.1145/3458723>
- Ghassemi, M., Oakden-Rayner, L. & Beam, A. L. (2021). The false hope of current approaches to explainable artificial intelligence in health care. *The Lancet Digital Health*, 3(11), e745–e750. [https://doi.org/10.1016/s2559-7500\(21\)00208-9](https://doi.org/10.1016/s2559-7500(21)00208-9)
- Glazko, K., Mohammed, Y., Kosa, B., Potluri, V. & Mankoff, J. (2024). *Identifying and improving disability bias in GPT-based resume screening*. In *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency (FAccT '24)* (pp. 687–700). Association for Computing Machinery. <https://doi.org/10.1145/3630106.3658933>
- Gomes, S., Lopes, J. M. & Nogueira, E. (2025). *Anthropomorphism in artificial intelligence: A game-changer for brand marketing*. *Future Business Journal*, 11(2), Article 2. <https://doi.org/10.1186/s43093-025-00423-y>
- Google Developers (2024). *Fairness: Mitigating bias*. Google Machine Learning Crash Course. <https://developers.google.com/machine-learning/crash-course/fairness/mitigating-bias> (13. 10. 2025.).
- Hamraie, A. & Fritsch, K. (2019). Crip Technoscience Manifesto. *Catalyst: Feminism, Theory, Technoscience*, 5(1), 1–34. <https://doi.org/10.28968/cftt.v5i1.29607>
- Horváth, L. (2023). Feltáró szakirodalmi áttekintés a mesterséges intelligencia oktatási használatáról. *Pannon Digitális Pedagógia*, 3(1), 5–17. <https://doi.org/10.56665/PADIPE.2023.1.1>
- Hu, T., Kyrchenko, Y., Rathje, S., Collier, N., van der Linden, S., & Roozenbeek, J. (2025). Generative language models exhibit social identity biases. *Nature Computational Science*, 5, 65–75. <https://doi.org/10.1038/s43588-024-00741-1>
- Könczei, Gy. & Hernádi, I. (2011). A fogyatékoságtudomány főfogalma és annak változásai. In Z. É. Nagy (Ed.), *Az akadályozott és az egészségtudományok közötti emberek élethelyzete Magyarországon: Kutatási eredmények a TÁMOP 5.4.1 projekt kutatási pillérében* (pp. 7–28). Nemzeti Család- és Szociálpolitikai Intézet.
- Jain, A. (2025). *Designing for ethical and inclusive AI through a human-centered design lens*. *Global Business & Economics Journal*. <https://doi.org/10.70924/f83n6wqz/kdxoiwt>
- Jobin, A., Lenca, M. & Vayena, E. (2019). The global landscape of AI ethics guidelines. *Nature Machine Intelligence*, 1(9), 389–399. <https://doi.org/10.1038/s42256-019-0088-2>
- Joshi, P., Santy, S., Budhiraja, A., Bali, K. & Choudhury, M. (2020). The state and fate of linguistic diversity and inclusion in the NLP world. In D. Jurafsky, J. Chai, N. Schluter, & J. Tetreault (Eds.), *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics* (pp. 6282–6293). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.acl-main.560>
- Jurafsky, D. & Martin, J. H. (2025). *Speech and language processing: An introduction to natural language processing, computational linguistics, and speech recognition with language models* (3rd ed.) [Online manuscript]. <https://web.stanford.edu/~jurafsky/slp3/> (9. 10. 2025.).

- Kamiran, F. & Calders, T. (2012). Data preprocessing techniques for classification without discrimination. *Knowledge and Information Systems*, 33(1), 1–33. <https://doi.org/10.1007/s10115-011-0463-8>
- Kamruzzaman, M. & Kim, G. L. (2025). The impact of disability disclosure on fairness and bias in LLM-driven candidate selection. *The International FLAIRS Conference Proceedings*, 38(1). <https://doi.org/10.32473/flairs.38.1.138911>
- Kleinberg, J. & Raghavan, M. (2021). Algorithmic monoculture and social welfare. *Proceedings of the National Academy of Sciences*, 118(22), e2018340118. <https://doi.org/10.1073/pnas.2018340118>
- Kooli, C., & Chakraoui, R. (2025). AI-driven assistive technologies in inclusive education: benefits, challenges, and policy recommendations. *Sustainable Future*, 2025(10), 101042. <https://doi.org/10.1016/j.sftr.2025.101042>
- Kruger, J. & Dunning, D. (1999). Unskilled and unaware of it: How difficulties in recognizing one's own incompetence lead to inflated self-assessments. *Journal of Personality and Social Psychology*, 77(6), 1121–1134. <https://doi.org/10.1037/0022-3514.77.6.1121>
- Liao, Q. V. & Vaughan, J. W. (2023). *AI transparency in the age of LLMs: A human-centered research roadmap*. Microsoft Research. <https://doi.org/10.48550/arXiv.2306.01941> and <https://doi.org/10.1162/99608f92.8036d03b>
- Lipton, Z. C. (2018). The mythos of model interpretability. *Communications of the ACM*, 61(10), 36–43. <https://doi.org/10.1145/3233231>
- Li, Y., Du, M., Song, R., Wang, X. & Wang, Y. (2023). *A survey on fairness in large language models*. [Preprint]. [arXiv. https://arxiv.org/abs/2308.10149](https://arxiv.org/abs/2308.10149) (10. 11. 2026.).
- Mehrabi, N., Morstatter, F., Saxena, N., Lerman, K. & Galstyan, A. (2021). A survey on bias and fairness in machine learning. *ACM Computing Surveys*, 54(6), Article 115. <https://doi.org/10.1145/3457607>
- Mező, K. (2025). A mesterséges intelligencia gyógypedagógiai célú felhasználási lehetőségei. *Különleges Bánásmód – Interdiszciplináris folyóirat*, 11(2), 143–153. <https://doi.org/10.18458/KB.2025.2.143>
- McGuire, C. (2020). *Measuring difference, numbering normal: Setting the standards for disability in the interwar period*. Manchester University Press. <https://doi.org/10.7765/9781526143167>
- Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence*, 267, 1–38. <https://doi.org/10.1016/j.artint.2018.07.007>
- Mitchell, M., Wu, S., Zaldivar, A., Barnes, P., Vasserman, L., Hutchinson, B., Spitzer, E., Raji, I. D., & Gebru, T. (2019). Model cards for model reporting. In d. boyd & J. H. Morgenstern (Eds.), *Proceedings of the Conference on Fairness, Accountability, and Transparency (FAT*)* (pp. 220–229). Association for Computing Machinery. <https://doi.org/10.1145/3287560.3287596>
- Mittelstadt, B. D., Allo, P., Taddeo, M., Wachter, S. & Floridi, L. (2016). The ethics of algorithms: Mapping the debate. *Big Data & Society*, 3(2), 1–21. <https://doi.org/10.1177/2053951716679679>
- Naseer, A., Ahmad, N. R. & Chishti, M. A. (2025). Psychological impacts of AI dependence: Assessing the cognitive and emotional costs of intelligent systems in daily life. *Review of Applied Management and Social Sciences*, 8(1), 291–307. <https://doi.org/10.47067/ramss.v8i1.458>
- Navigli, R., Conia, S. & Ross, B. (2023). Biases in large language models: Origins, inventory, and implications. *ACM Journal of Data and Information Quality*, 15(2), Article 10. 1–21. <https://doi.org/10.1145/3597307>
- Nazer, L. H., Zatarah, R., Waldrip, S., Ke, J. X. C., Moukheiber, M., Khanna, A. K., Hicklen, R. S., Moukheiber, L., Moukheiber, D., Ma, H. & Mathur, P. (2023). Bias in artificial intelligence algorithms and recommendations for mitigation. *PLOS Digital Health*, 2(6), e0000278. <https://doi.org/10.1371/journal.pdig.0000278>
- New York City Bar Association (2025). *The impact of the use of artificial intelligence on people with disabilities*. Presidential Task Force on Artificial Intelligence and Digital Technologies. [https://www.nycbar.org/reports/the-impact-of-the-use-of-ai-on-people-with-disabilities/\(15.11.2025\).](https://www.nycbar.org/reports/the-impact-of-the-use-of-ai-on-people-with-disabilities/(15.11.2025).)
- Obermeyer, Z., Powers, B., Vogeli, C. & Mullainathan, S. (2019). Dissecting racial bias in an algorithm used to manage the health of populations. *Science*, 366, 447–453. <https://doi.org/10.1126/science.aax2342>
- OECD (2025). *Governing with Artificial Intelligence: How Artificial Intelligence Is Accelerating the Digital Government Journey*. Organisation for Economic Co-operation and Development. https://www.oecd.org/en/publications/2025/06/governing-with-artificial-intelligence_398fa287/full-report/how-artificial-intelligence-is-accelerating-the-digital-government-journey_d9552dc7.html

- Oliver, M. (1990). *The politics of disablement*. Macmillan Education. <https://doi.org/10.1007/978-1-349-20895-1>
- Pennycook, G., Ross, R. M., Koehler, D. J. & Fugelsang, J. A. (2017). Dunning–Kruger effects in reasoning: Theoretical implications of the failure to recognize incompetence. *Psychonomic Bulletin & Review*, 24(6), 1774–1784. <https://doi.org/10.3758/s13423-017-1242-7>
- Phutane, M., Jung, H., Kim, M., Mitra, T. & Vashistha, A. (2025). *ABLEIST: Intersectional Disability Bias in LLM-Generated Hiring Scenarios*. arXiv. <https://doi.org/10.48550/arXiv.2510.10998>
- Prates, M. O. R., Avelar, P. H. & Lamb, L. C. (2020). Assessing gender bias in machine translation: A case study with Google Translate. *Neural Computing and Applications*, 32(10), 6363–6381. <https://doi.org/10.1007/s00521-019-04144-6>
- Shakespeare, T. (2014). *Disability rights and wrongs revisited* (2nd ed.). Routledge. <https://doi.org/10.4324/9781315887456>
- Shew, A. (2020). Ableism, technoableism, and future AI. *IEEE Technology and Society Magazine*, 39(1), 40–85. <https://doi.org/10.1109/MTS.2020.2967492>
- Soldatic, K., Lee, M., Tunggal, E., Liao, A. & Magee, L. (2025). Rethinking digital and AI inclusion: Participatory and intersectionality-informed methods for disability and migrant justice. *Frontiers in Sociology*, 10, 1593330. <https://doi.org/10.3389/fsoc.2025.1593330>
- Stahl, B. C. (2021). *Artificial Intelligence for a Better Future*. Springer. <https://doi.org/10.1007/978-3-030-69978-9>
- Tilmes, N. (2022). Disability, fairness, and algorithmic bias in AI recruitment. *Ethics and Information Technology*, 24, Article 21. <https://doi.org/10.1007/s10676-022-09633-2>
- Treviranus, J. (2023). Inclusive Design: Valuing Difference, Recognizing Complexity. In M. H. Rioux, A. Buettgen, E. Zubrow, & J. Viera (Eds.), *Handbook of Disability: Critical Thought and Social Change in a Globalizing World* (pp. 1–24). Springer. https://doi.org/10.1007/978-981-16-1278-7_98-1
- Trewin, S., Basson, S., Muller, M., Branham, S., Treviranus, J., Gruen, D., Hebert, D., Lyckowski, N. & Manser, E. (2019). Considerations for AI fairness for people with disabilities. *AI Matters*, 5(3), 40–63. <https://doi.org/10.1145/3362077.3362086>
- United Nations (2006). *Convention on the Rights of Persons with Disabilities (UNCRPD)*. United Nations General Assembly, A/RES/61/106. <https://www.un.org/disabilities/documents/convention/convoptprot-e.pdf>
- WebAIM (2023). *The WebAIM Million – The 2023 report on the accessibility of the top 1000000 home pages*. <https://webaim.org/projects/million/2023> (22. 12. 2025.).
- World Health Organization (2001). *International classification of functioning, disability and health: ICF*. World Health Organization. <https://iris.who.int/handle/10665/42407>
- World Health Organization & World Bank (2011). *World report on disability*. WHO. <https://www.who.int/publications/b/online-reader/30358>
- World Wide Web Consortium (2023). *Web Content Accessibility Guidelines (WCAG) 2.2*. <https://www.w3.org/TR/WCAG22/> (15. 12. 2025.).
- Young, S. (2014). *I'm not your inspiration, thank you very much* [TEDx talk]. TEDxSydney. https://www.ted.com/talks/stella_young_i_m_not_your_inspiration_thank_you_very_much
2025. évi LXXV. törvény az Európai Unió mesterséges intelligenciáról szóló rendeletének magyarországi végrehajtásáról (Act LXXV of 2025 on the Hungarian implementation of the European Union's regulation on artificial intelligence). *Magyar Közlöny*, 2025/127. <https://njt.hu/jogszabaly/2025-75-00-00>

Köszönetnyilvánítás

A tanulmány elkészítését az MTA Közoktatás-fejlesztési Kutatási Programja támogatta.