

When AI Policy Becomes an Accessibility Barrier

ADA Exposure from Restrictions on Artificially Intelligent Disability Assistance (AIDA)

Basil C. Puglisi, MPA

basilpuglisi.com

Working Paper, September 2026 DOI: 10.5281/zenodo.22940225

Companion paper: When AI Becomes Assistive Technology: Artificially Intelligent Disability Assistance (AIDA) (Puglisi, 2026d). <https://doi.org/10.5281/zenodo.22908272>

Keywords: Artificially Intelligent Disability Assistance, AIDA, Americans with Disabilities Act, reasonable accommodation, reasonable modification, AI policy, anti-AI policy, AI disclosure, AI detection, facial neutrality, EU AI Act Article 50, Checkpoint-Based Governance

Companion paper: *When AI Becomes Assistive Technology: Artificially Intelligent Disability Assistance (AIDA). How General-Purpose AI Can Help People Work, Communicate, Think, and Create Around Physical and Cognitive Limitations* (Puglisi, 2026d)

Abstract

General-purpose AI may function as assistive technology when a person with a disability uses it to increase, maintain, or improve a functional capability. The companion paper names that function Artificially Intelligent Disability Assistance (AIDA) and offers a three-prong functional test for when a particular use qualifies, paired with a separate governance screen. This paper asks what follows once it does. An institution's blanket restriction, mandatory disclosure, distribution penalty, discipline, or detection rule may then implicate existing disability-law obligations, even when the AI policy is facially neutral.

The paper poses a question. If general-purpose AI is a disability tool, is what institutions are doing to restrict it unethical at best, and illegal at worst? The work is to find which

applies where, measured against the limits the law already keeps. The problem it examines is AI policy that restricts or penalizes the people who use AI. That covers anti-AI rules that stop the use outright, and wider rules that penalize it through disclosure mandates, labels, reduced distribution, discipline, or detection. AI Governance, which is human oversight and accountability, is the remedy. A requirement to admit or disclose assistive use may itself breach ethical duties or disability protections. It sets out the reasons institutions give for restricting AI and classifies six types of policy risk. It maps which United States statutes reach which actors, works three gray-area hypotheticals, and proposes a seven-gate screen for accommodation exposure. It compares the European Union, where the Article 50 Guidelines exempt meaning-preserving assistive communication while marking the summarizing and restructuring that cognitive disabilities often rely on. It closes with a model accommodation clause that generalizes practice already in place at the Open University and at a United Kingdom government department. Existing law, applied actor by actor, decides whether any duty follows. The register throughout is potential exposure.

Author's Note

How the question arrived. HAIA, which stands for Human Artificial Intelligence Assistant, took its name after a shoulder surgery in September 2025 took away the use of my dominant arm and shifted my workflow from typing to voice (Puglisi, n.d.-a). The companion paper tells that story in its own Author's Note. It is the example of AI's value that brought this subject to my attention, and it is not evidence for any claim in this paper.

Where the #AIassisted tag came from. From 2023 through 2026, content on basilpuglisi.com carried two tags, #AIassisted and #AIgenerated, which separated for live readers the work AI had assisted from the work AI had generated. My position has been that all content is AI-assisted and has been for a decade, because search engines and library catalogs have long served some part of the work; the question is how the tools are used.

In April 2026, Dr. Joy Buolamwini commented publicly on LinkedIn on my work on her evocative audit (Buolamwini, 2026). From her comment, with LinkedIn's rendering of hashtags corrected:

I noticed you use the #AIassisted marker and appreciate the transparency. We all see on social media many people are using AI in their writing without any explicit indication. Others take the opposite approach to be explicit that a post was #HumanMade.

She then asked what #AIassisted meant in my case: whether I had read the work first and used AI to organize ideas around it, or had an AI tool analyze the document.

My public answer explained #AIassisted as practice in the HAIA ecosystem under Checkpoint-Based Governance, with human authority and accountability. I read the source first, and multiple AI platforms test that reading in parallel. The arguments and the voice are mine. I described #AIgenerated as Responsible AI, where the machine produced the work and a human signed it, and #AIassisted as AI Governance, where the work is human-led with a named human accountable. The full answer followed by InMail.

Her question led to the mark that now closes my work, “#AIassisted using the HAIA Ecosystem,” which declares on the page what she asked about. The exchange records how that mark developed. It does not document stigma, and this paper does not treat it as evidence of stigma. The evidence on how disclosure is received sits in Section 10.

Author standpoint. This author feels that the industry created a problem with automation that made it the face of the industry, AI Agents and what I call Responsible AI, code is no substitute for humans. The fact that AI has so many positive impacts on humanity should not be dismissed or punished because of “possible” misuse or future issues. As someone who had AI to help when I lost the temporary use of an arm, I could not imagine telling someone with long term or permanent disabilities that they can not use a tool that would help them compensate, compete and have a better life because of the “what ifs”. I also believe this holds with United States law, which does not punish people for the tools they use to offset their disabilities, while keeping its existing limits in place: undue hardship, fundamental alteration, essential functions, and legitimate security requirements.

Responsible AI, as the author defines it, is machine checking machine: an AI agent running the pipeline, AI reviewing AI, or a human in the loop who only watches, rarely acts, and barely observes (Puglisi, 2026a).

Standpoint of the Paper

This paper presents governance and policy analysis informed by statutes, regulations, cases, and institutional practice. The legal reading is that of a retired police officer with a Master of Public Administration whose studies included constitutional law, business law, procedural law, courts, and judicial case studies. It is not legal advice, it does not predict the outcome of any dispute, and no attorney reviewed it. The register stays hedged throughout, using terms such as potential exposure, may implicate, and emerging issue. Every statute and case is cited from a primary source where one was retrievable, and pending litigation is described only through its reported allegations.

1. The Starting Point: A Case Already in Court, and Guidance That Ran One Direction

The question. This paper is built around a conditional question. The companion paper argues that general-purpose AI can function as a disability tool. If it does, the question is whether the rules institutions are writing to restrict it are unethical at best, and illegal at worst. The paper tests that proposition against the limits the law already keeps: undue hardship, fundamental alteration, essential functions, and legitimate security requirements. A rule that one of those limits justifies may be neither. For a rule that none of them justifies, which of the two it becomes depends on the actor, the person, the use, and the law that reaches them, and the sections that follow work through each.

The stakes run past any single case. The author holds that the future of work is Augmented Intelligence, in which human-AI collaboration lets people who are willing to work and learn compete with those who hold greater natural or credentialed advantage. A rule that takes that tool away from a person with a disability removes the equalizer from the person it helps most.

1.1 Cedeno

Bruno Cedeno v. Walt Disney Parks and Resorts U.S., Inc., No. 6:25-cv-02046 (M.D. Fla., filed October 23, 2025), is the clearest sign that the pattern this paper describes has reached federal court. Every fact below is an allegation from the amended complaint, the operative pleading filed January 7, 2026, and the case establishes nothing (Am. Compl., Doc. 23).

- The plaintiff, Angeliz E. Bruno Cedeno, a security host, alleges light sensitivity from postpartum eye conditions and astigmatism that substantially limits seeing (§ 22).
- Her physician prescribed Meta smart glasses as a medical device, and she wore them at work without incident until June 2025 (§§ 23, 25).
- Management could cite no written policy or rule the glasses violated, and a vice president confirmed in writing that no existing policy prohibited the device (§§ 29, 39).
- She filed a formal accommodation request on June 27, 2025, and her physician's forms confirmed the glasses were medically necessary. Employee Relations offered only relocation or medical leave (§§ 24, 34, 35).
- On July 24, 2025, managers announced an "executive decision" banning smart glasses "effective immediately," including her prescribed device. No written policy was shown, and the manager refused to say who made the decision (§§ 44, 45, 110).

- A coworker’s text reported that managers were “updating the policy to not include Meta or smart glasses,” which the pleading offers as proof that no policy existed when she was disciplined (§ 57).
- Other employees wore smartwatches and other electronic devices without restriction, and other cast members wore Meta glasses without discipline (§§ 64, 114, 115).
- The pleading argues that the device could be configured so that it did not record guests or confidential information. It adds that the employer could have used written guidance, inspections, disabled features, or post-by-post limits instead of denying the accommodation outright (§ 94).
- It asserts failure to accommodate, disparate treatment, and retaliation under the ADA (Counts I to III).

The case remained pending as of September 2026. It involves an AI-enabled wearable, so it cannot settle the questions this paper raises about writing and reasoning tools. What the amended complaint alleges illustrates the facial neutrality pattern in its plainest form: a ban announced as an executive decision after an individual request, with no written policy behind it and uneven enforcement around it. Paragraph 94 also carries the argument this paper makes in Section 8.2, that a security concern calls for safeguards on the tool, leaving the person’s accommodation in place. It raises one more distinction this paper keeps throughout. A device a physician prescribes sits closer to traditional assistive technology than a general-purpose tool the person selects, and an institution may treat the two differently.

1.2 The inversion

Federal guidance on AI and disability ran in one direction. It addressed AI used on people with disabilities, such as screening tools, algorithmic assessments, and the duty to accommodate inside an AI-driven hiring process. The EEOC removed its AI technical assistance from its website in January 2025, though the statutes the guidance interpreted did not change (Cooley LLP, 2025). As of September 22, 2026, this review located no current federal guidance addressing general-purpose AI that a disabled person selects as an accommodation.

Mobley v. Workday, Inc., No. 3:23-cv-00770-RFL (N.D. Cal.), shows the distinction. It concerns AI used to screen applicants. The court allowed the age-discrimination claim to proceed as a collective on May 16, 2025, and the case also tests whether an AI vendor can answer for screening it performs on behalf of employers. That question returns in Section 7. The inversion this paper examines runs the other way: the person with a disability holds the tool, and the institution writes the rule.

The model is simple to state:

Human capability → disability-related barrier → AI intervention → restored or expanded functional access → human action

Key events, January 2025 to December 2026



Sources: Cooley LLP (2025); A.J.T., 605 U.S. 335; Am. Compl., Doc. 23; Payan, No. 24-1809; European Commission (2026); GCN (2026). The final date is a scheduled deadline.

Figure 1. Key legal, regulatory, and platform events from January 2025 to December 2026. The Cedeno entries are allegations; the final date is a scheduled deadline.

2. Foundation, Compressed

The companion paper develops the functional test and the AIDA definition, and this paper relies on them without repeating the argument (Puglisi, 2026d). The legal questions here rest on the three-prong functional test alone. The companion paper's governance screen asks how authority and accountability are allocated once a use qualifies, and it creates no legal entitlement and no legal condition. Four legal points carry the analysis here.

First, disability under the ADA turns on substantial limitation of a major life activity, which falls well short of inability. The regulations state that "substantially limits" is not meant to be a demanding standard, and that limitation of one major life activity is enough (29 C.F.R. § 1630.2(j)(1)). Second, employment law asks a separate question: whether the person is a qualified individual who can perform the essential functions of the job with or without reasonable accommodation. A person can be disabled and fully capable of high-level professional work once accommodated.

Third, success with a mitigating measure does not erase the disability. The statute requires that substantial limitation be assessed without regard to the ameliorative effects of mitigating measures, and it lists the "use of assistive technology" among them (42 U.S.C. § 12102(4)(E)(i)). That answers the argument that a person who can work with AI must not be disabled. Fourth, the Assistive Technology Act supplies the functional concept. It defines an assistive technology device as any item, piece of

equipment, or product system, “whether acquired commercially, modified, or customized,” used to increase, maintain, or improve functional capabilities (29 U.S.C. § 3002(4)). The Individuals with Disabilities Education Act applies the same test to children and adds the phrase “commercially off the shelf” (20 U.S.C. § 1401(1)). Neither statute creates the accommodation duties, which arise under the ADA, the Rehabilitation Act, and in education, Section 504 and IDEA.

Nearest legal prior work. Gilly (2026b) argues that the accommodation duty in disability law is an affirmative obligation, distinct from the analysis of discrimination by intent or by impact. Gilly applies that argument to AI deployed on people. This paper applies the same duty to AI used by people.

Convenience for one person, access for another. The tool alone cannot decide whether a use is assistive. The same drafting feature is a convenience for one employee and an accommodation for another, and the relationship between the disability, the barrier, and the use decides which.

2.1 The working hypotheticals

The companion paper runs seven hypothetical cases: blindness, speech loss, dyslexia, intellectual and developmental disability, post-traumatic stress disorder, a combined physical limitation, and a combined borderline person. This paper works the last three, because they are where an institution is most likely to dispute disability status, combined functional effects, or the link between the AI use and the barrier.

Regarded-as coverage and temporary impairments. The six-month limit for impairments that are transitory and minor applies only to the regarded-as prong of the definition. It does not apply to actual disability or to a record of disability (29 C.F.R. § 1630.2(j)(1)(ix)). Coverage under the regarded-as prong alone carries no accommodation duty (42 U.S.C. § 12201(h)). A temporary limitation can therefore still qualify under the actual-disability prong when it substantially limits a major life activity.

Post-traumatic stress disorder. This is a borderline case, and its qualification as a disability is part of the debate. The regulation names PTSD among the impairments that should easily be found to substantially limit brain function (29 C.F.R. § 1630.2(j)(3)(iii)), and concentrating, thinking, and communicating are also listed major life activities. Both points weigh toward coverage, yet the determination remains an individualized assessment of the particular person (29 C.F.R. § 1630.2(j)(1)(iv)), and an employer or school may contest it. The functional link between a given AI use and a PTSD-related barrier is a second open question.

Combined physical limitation. Walking, standing, sitting, and performing manual tasks are listed major life activities. The case rests on the combined functional effect on

those activities and asserts no special rule for aggregating impairments. It does not rest on the major life activity of working, which the regulations treat by class or type of work. An employer will call eight hours at a workstation an essential function. Section 12 answers that the output may be the standard while the workstation is only the means, and Section 8 states when that answer fails.

3. Context as of September 2026

This section is dated so that the argument does not depend on one week's news.

The public climate is polarized. On September 14, 2026, the President posted that fears of AI “taking over the World, destroying Humanity” were a “HOAX,” pushing back against technology leaders who had called for slower development and greater oversight (Trump, 2026, as reported by Associated Press, 2026). One pole fears the technology, and the other dismisses the fear.

Both poles may impose costs on the assistive user. Fear attaches stigma to anyone seen using AI, so a disabled writer who uses a drafting tool inherits suspicion aimed at the technology. Dismissal leaves the guidance vacuum described in Section 1, where no agency tells an employer or a school how to treat an accommodation request for a general-purpose tool.

Public opinion leans toward concern. Half of United States adults report being more concerned than excited about AI (Pew Research Center, 2026). The share saying AI does more harm than good rose from 31% in 2025 to 39% in 2026 (Gallup, 2026). In a March 2026 poll, 80% of respondents said they were very or somewhat concerned (Quinnipiac University Poll, 2026).

The dissent this paper must carry. Seventy-six percent of Americans say it is extremely or very important to be able to tell whether content was made by AI or by people. In the same survey, 53% were not confident they could tell (Pew Research Center, 2025). Disclosure has a real public mandate, and this paper does not dismiss it. The question the paper asks is who pays the cost of meeting that mandate, and whether people with disabilities pay more of it than anyone else.

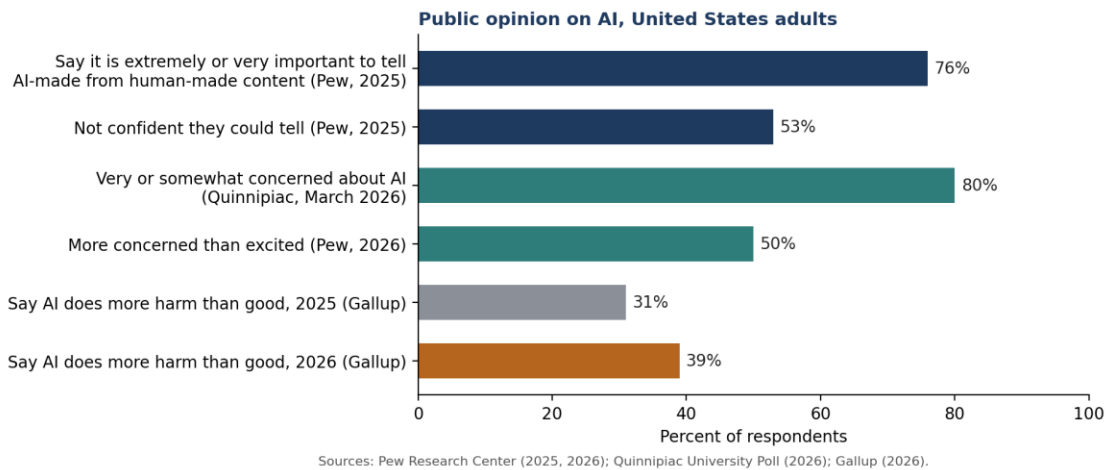


Figure 2. *United States public opinion on AI and on telling AI-made content from human-made content. Sources: Pew Research Center (2025, 2026); Quinnipiac University Poll (2026); Gallup (2026).*

4. Why Institutions Restrict AI

Every restriction has a reason, and the reasons deserve a fair statement before any critique.

Employers cite data security, confidentiality, intellectual property, accuracy, and reputation. In the Cisco 2024 Data Privacy Benchmark Study, 27% of organizations had banned generative AI at least temporarily. Respondents cited risks to legal and intellectual property rights (69%), disclosure of information to the public or competitors (68%), and wrong outputs (68%) (Cisco, 2024). The 2026 edition reports that 90% of organizations expanded their privacy programs because of AI (Cisco, 2026).

Platforms protect their economics and identity. Medium’s policy protects a shared revenue pool and a reputation for human storytelling (Medium, n.d.).

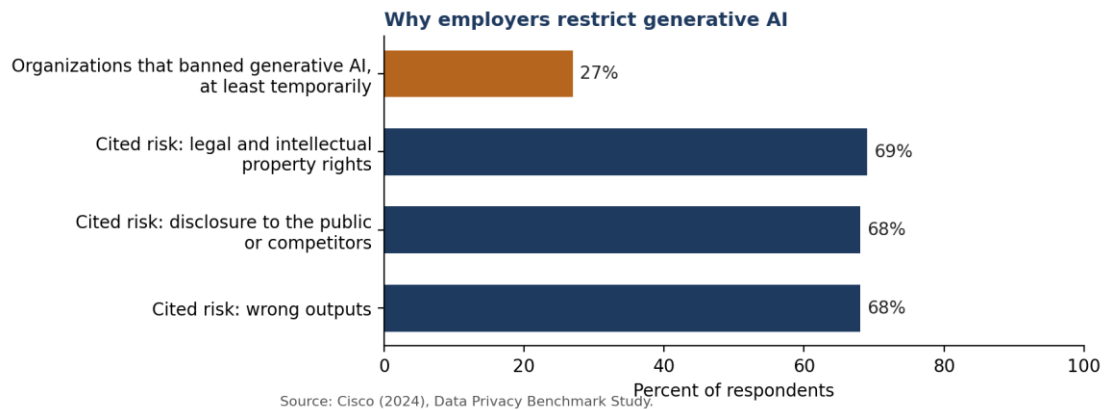


Figure 3. *Share of organizations that banned generative AI, and the risks they cited. Source: Cisco (2024).*

Publishers protect authorship, which AI cannot hold. The ICMJE recommendations ask journals to require authors to disclose AI-assisted technology at submission, in the cover letter and in the work. They keep human authors responsible for accuracy, integrity, and originality (International Committee of Medical Journal Editors, n.d.). The recommendations contain no accessibility exception.

Schools protect the assessment of the student's own knowledge and reasoning. Turnitin itself states that its AI writing assessment may misidentify human-written text and should not be the sole basis for adverse action against a student (Sheridan College, n.d.).

The finding. None of these reasons concerns the method by which a worker, author, or student reaches that person's own capability. Each concerns data, accuracy, authorship, or the integrity of what is being assessed. Where an employer can offer a secure, approved alternative, that alternative satisfies the stated rationale for a ban, so the rationale alone does not justify denying an accommodation that the alternative could meet.

5. The Facial Neutrality Problem

AI policies are often written as rules that apply to everyone. "No employee may use generative AI." "AI-assisted submissions are prohibited." "AI-assisted work receives reduced distribution." Each rule applies to every person equally, and each can still bar the means through which a particular disabled person does the work.

Disability law has long recognized that equal application does not remove the need for disability-related modification. The accommodation duty exists precisely because a rule that treats everyone the same can exclude a person whose disability requires a different route. *US Airways, Inc. v. Barnett*, 535 U.S. 391 (2002), marks the limit on the other side. The Court held that where a requested accommodation conflicts with the rules of an established seniority system, the employer's showing of that conflict ordinarily defeats the request. The employee may still show special circumstances that make an exception reasonable in the particular case. *Barnett* concerned seniority, and carrying its reasoning to a company-wide restriction on AI is an analogy this paper draws beyond the Court's holding.

A universally applied AI rule can still create individualized accommodation obligations.

6. Six Types of AI Policy Risk

AI policies reach assistive users in six distinct ways, and each carries a different kind of exposure. Only the first is anti-AI in the strict sense, a rule that stops the use. The other five regulate how AI use is carried out, disclosed, received, or punished.

- **Type I, Use Prohibition.** The policy denies the assistive means entirely.
- **Type II, Process Restriction.** The policy allows AI in general but forbids the specific functions that offset the limitation, such as outlining, organizing, editing, drafting, or restructuring.
- **Type III, Disclosure Requirement.** The policy requires a label, which can carry stigma, expose private information, and change how valid work is received.
- **Type IV, Distribution or Opportunity Penalty.** The policy grants formal access while reducing substantive participation through lower reach, lost monetization, or reduced consideration or eligibility.
- **Type V, Discipline or Exclusion.** The policy punishes the use. Exposure is strongest where an accommodation was requested or the institution knows why the tool is being used.
- **Type VI, Perception Enforcement.** Detection tools and crowd reports act on whether work seems to be AI-made, whether or not AI was used. A disabled writer whose prose reads as machine-like can be flagged whether or not any tool was involved.

Six types of AI policy risk (Type I is anti-AI in the strict sense)

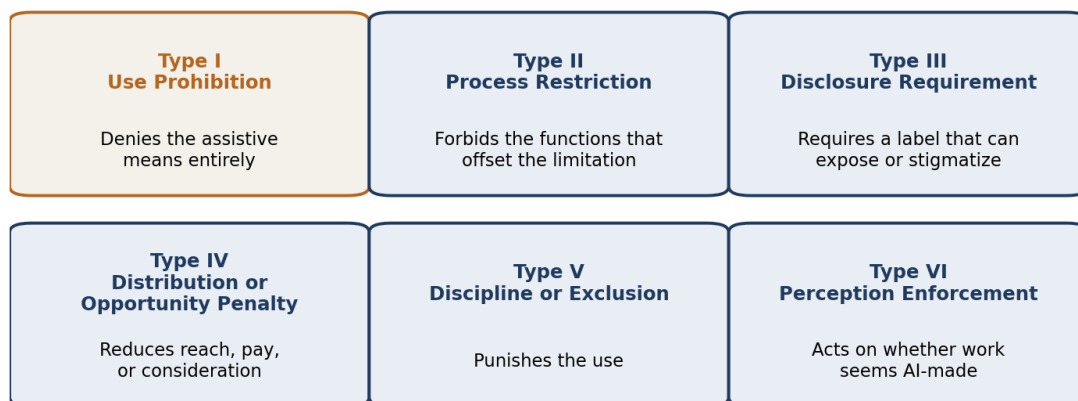


Figure 4. Six types of AI policy risk. Only Type I stops the use; Types II to VI regulate how AI use is carried out, disclosed, received, or punished.

6.1 The Type VI record, with its causal limit

Perception enforcement is new, so its record deserves care.

LinkedIn launched a “Seems like AI slop” reporting option on July 30, 2026, and more than one million unique members had used it by August 20 (GCN, 2026). LinkedIn reports roughly 40% fewer views on content it classifies as low-quality AI content. New classifiers launched on the same day as the reporting option, and LinkedIn did not separate their effect from the effect of reader flags, so the view reduction cannot be attributed to reports alone. LinkedIn states that no single report determines distribution, and authors can receive a notice that members flagged a post. The same platform also replaced its “enhance your post” rewriting tool with proofreading designed to preserve the author’s voice. In doing so, the platform itself drew a line between preserving the author and replacing the author.

Other enforcement systems show the same pattern. Medium uses detection tools combined with human review of positive results (Medium, n.d.). Turnitin lists repetitive text, and lists without structural variation, among the patterns that can produce false positives (Sheridan College, n.d.). The American Foundation for the Blind found that 12 disabled workers under AI surveillance, 7 percent of that group, reported an accommodation that the surveillance could not account for. One worker reported being flagged for underperformance while taking breaks granted as an accommodation (Shock & Silverman, 2026). Automated enforcement there misread an accommodation as a deficiency, which is the Type VI harm in its plainest form.

7. Who Is Reached: The Actor and Statute Matrix

Exposure depends on which statute reaches which actor, and the honest answer varies widely. The table below maps that reach actor by actor, with the known limit on each.

Actor	Primary federal hook	Known limit
Private employer	ADA Title I	Applies to employers with 15 or more employees. Protects employees and applicants; independent contractors are generally outside Title I, subject to analysis of employee status
Federal employer	Rehabilitation Act § 501	Federal employees proceed under § 501 in place of Title I
Federal contractor	Rehabilitation Act § 503	Applies to contractors above a contract-value threshold set in regulation
State and local government	ADA Title II, and Title I for employment	<i>Board of Trustees of the University of Alabama v. Garrett</i> , 531 U.S. 356 (2001), bars state employees' Title I suits for money damages against states. Whether Title II reaches public employment is split, as <i>Garrett</i> noted at 360 n.1
Public-facing business	ADA Title III	Circuit split on businesses that operate purely online. Private plaintiffs are limited to injunctive relief under federal law
School or university	Section 504, and ADA Titles II and III	Section 504 requires federal financial assistance. Fundamental alteration and academic integrity limit the duty
Kindergarten through grade 12	IDEA and Section 504	IDEA expressly incorporates assistive technology devices and services
Health program receiving federal funds	Section 1557, 45 C.F.R. § 92.205 (text as of September 2026)	Reasonable modification of policies unless it fundamentally alters the health program
Credentialing and testing body	42 U.S.C. § 12189 and 28 C.F.R. § 36.309	Examinations must reflect aptitude rather than impaired sensory, manual, or speaking skills, except where those skills are what is measured. Written expression and executive function are not among the enumerated skills
Publisher	Varies; often none under Title I for freelance authors	State law and federal funding status matter most
Platform	Title III where it applies	Whether an online-only platform is covered remains unsettled
Individual	Generally none directly	42 U.S.C. § 12203 reaches interference and retaliation; see Section 11

State disability and civil rights statutes may extend reach past this federal floor, depending on the jurisdiction.

The online question. The Eleventh Circuit held in *Gil v. Winn-Dixie Stores, Inc.*, 993 F.3d 1266 (11th Cir. 2021), that websites are not places of public accommodation, then vacated that opinion as moot on December 28, 2021. The underlying split remains. The Third, Sixth, and Ninth Circuits have required a nexus to a physical place, while the First and Seventh have not required one (Perkins Coie, n.d.).

The public employment question. The Eleventh Circuit allowed employment claims against public entities under Title II in *Bledsoe v. Palm Beach County Soil & Water Conservation District*, 133 F.3d 816 (11th Cir. 1998). The Ninth Circuit did not in *Zimmerman v. Oregon Department of Justice*, 170 F.3d 1169 (9th Cir. 1999), and the Tenth Circuit later joined the Ninth in *Elwell v. Oklahoma ex rel. Board of Regents*, 693 F.3d 1303 (10th Cir. 2012). The Supreme Court noted the disagreement in *Garrett* without resolving it, and *Skinner v. Salem School District*, No. 09-cv-193-JL, 2010 DNH 106 (D.N.H. June 18, 2010), collects the split.

AI vendors. Supplying an AI system does not by itself make the vendor responsible for the customer's accommodation duty. Independent obligations may arise from the vendor's own legal status, conduct, agency relationship, funding, or contractual role. Vendors also enter the analysis through contract, product changes, and pricing, any of which can remove an accommodation after an employer or school approves it. Gilly (2026a) names that harm accommodation repricing, and identifies the most exposed position as an accessibility capability embedded in a platform an institution procured. The vendor can meter or withdraw that capability in a transaction to which the disabled user is not a party. *Mobley v. Workday* tests a separate question, whether a vendor can answer for AI it deploys on people.

8. Employers: Accommodation, Security, and What Happens After the Request

8.1 The accommodation duty

Reasonable accommodation can include modified equipment, changed work methods, and modified policies. When needed to identify an effective accommodation, the employer and the individual engage in an informal interactive process to identify the limitation and the accommodation (29 C.F.R. § 1630.2(o)(3)). The employer may choose among effective accommodations and need not provide the one the employee prefers. The person's preference still matters in practice. Phillips and Zhao (1993) found that a user's input in selecting a device was among the predictors of assistive technology abandonment. That finding gives the interactive process a practical reason to weigh the person's choice of tool.

Consider a hypothetical. An employer bans generative AI. A qualified employee with a documented disability requests an approved tool for specified functions, and the employer answers that the rule applies to everybody. That answer may bypass the analysis the law requires. The proper questions are whether the modification is reasonable, whether it is effective, whether another effective accommodation is available, whether it would impose undue hardship, and whether it would eliminate or displace an essential function.

8.2 The limits

Disability status creates no unrestricted right to any AI system. An employer may bar confidential records from a public model, and health privacy law, client confidentiality, and in the European Union the General Data Protection Regulation give that restriction real weight. The analysis does not end there, however. It continues to the alternatives that could meet the same need: an enterprise AI deployment, a local or private model, dictation software, approved summarization, or another adaptive tool.

When “the workstation is the means” fails. Presence, timing, the security environment, or the method of work can themselves be essential functions of a specific job. A position that requires physical presence at a secured console differs from one that requires only a finished analysis. The answer in Section 12, that the output may be the standard while the workstation is the means, holds only where those features are not themselves essential.

8.3 Employers already doing it

At least one public employer has already resolved the question in favor of access. Defra, the United Kingdom’s Department for Environment, Food and Rural Affairs, made full Microsoft Copilot licences available as a reasonable adjustment after trials showed benefits for neurodivergent and disabled colleagues (Ede, 2026). One employee describes reviewing every output and calls the tool a thinking partner rather than a decision maker. In the research literature, Brand, Fischer Mogensen, and Engelbrecht (2026) argue that generative AI can function as reasonable accommodation in employment across disability types.

8.4 Evidence from the workplace

The American Foundation for the Blind’s *Working with the Machine* surveyed 1,374 workers between July and October 2025 (Shock & Silverman, 2026). Its findings form this paper’s empirical anchor for Type I and Type II risk:

- 68% of workers used AI at work, with no difference by disability status.
- Among workers with disabilities, 32% used AI for visual description and 25% for captioning on the job.

- 19% of all workers said workplace AI rules affected how well they could do their jobs, with qualitative responses pointing mostly to negative effects.
- Workers with disabilities reported assistive AI they could not use because of privacy rules.
- A blind job candidate closed a screen reader during an automated test that warned against other open applications, then hired a friend to operate the mouse.

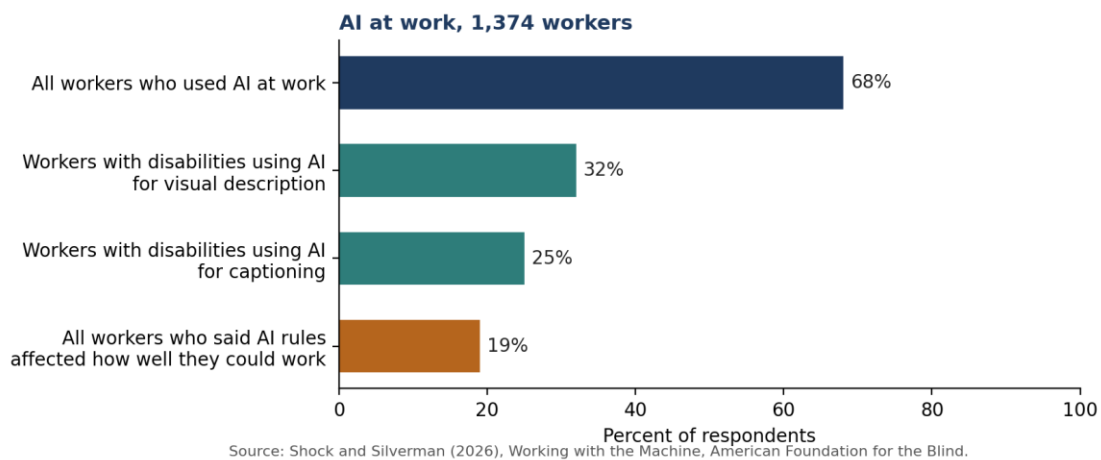


Figure 5. AI use at work in a sample of 1,374 workers. Source: Shock and Silverman (2026).

A companion AFB analysis of the same survey treats visual description, captions, and augmentative and alternative communication directly as AI assistive technology (Hanuschock & Silverman, 2026). AFB recommends allowing workers to use AI for discrete access tasks as a reasonable accommodation, and adjusting company-wide IT policies to permit disability-related AI tools.

8.5 Barnett's rebuttal route

Under *Barnett*, an exception to an established seniority system is ordinarily unreasonable, but the employee may show special circumstances. This paper uses *Barnett*'s special-circumstances reasoning by analogy, since *Barnett* itself does not establish that rule for restrictions on AI outside seniority systems. One example the Court described is a system with so many exceptions that one more would not matter. The *Cedeno* allegation that other employees wore smartwatches and other devices without restriction, and that other cast members wore Meta glasses without discipline (Am. Compl. ¶¶ 64, 114) is that argument in the form of a pleading. Whether it succeeds depends on facts the case has not yet developed.

The contrast the section turns on is simple:

Legitimate restriction + effective alternative versus **Blanket denial + no accommodation analysis**

8.6 After the request

Exposure shifts once a request is on record, and five scenarios should be kept apart:

- a policy violation with no accommodation need identified
- a request that is denied
- discipline because of the request
- an approval that is later penalized
- an employee treated as less capable for needing AI

As a matter moves down that list, exposure can move from failure to accommodate toward retaliation or interference.

9. Government, Public-Facing Business, and Education

9.1 The modification duties

Title II reaches public universities, licensing bodies, benefits agencies, state research systems, and AI policies that governments themselves write. Its regulation requires a public entity to make reasonable modifications in policies, practices, or procedures when necessary to avoid discrimination on the basis of disability. The only exception is a modification the entity can show would fundamentally alter the service, program, or activity (28 C.F.R. § 35.130(b)(7)). *Wong v. Regents of the University of California*, 192 F.3d 807 (9th Cir. 1999), requires a fact-specific, individualized analysis of the person's circumstances. Mere speculation that a modification is unreasonable falls short of that requirement (ILRU, n.d.).

Title III carries the parallel duty for public accommodations, requiring reasonable modification of policies unless the modification would fundamentally alter the goods or services offered (28 C.F.R. § 36.302, implementing 42 U.S.C. § 12182(b)(2)(A)(ii)). It also bars eligibility criteria that screen out or tend to screen out people with disabilities. The online coverage question from Section 7 applies here.

Section 504 bars exclusion or discrimination “solely by reason of her or his disability” in any program receiving federal financial assistance (29 U.S.C. § 794(a)). For employment claims, it applies the standards of ADA Title I (29 U.S.C. § 794(d)).

9.2 Education

Schools have real interests in assessing knowledge, reasoning, writing, competency, and integrity, and nothing in this paper disputes them. Accommodation may change how a student reaches the opportunity to show what is being assessed, without changing what is assessed. The question in each case is whether a particular AI use alters the essential nature of the assignment.

The leading authority. *PGA Tour, Inc. v. Martin*, 532 U.S. 661 (2001), held that a refusal to consider an individual's circumstances runs counter to the ADA's requirement of an individualized inquiry. It also held that waiving a peripheral rule that does not impair the rule's purpose is not a fundamental alteration. The record showed that the golf cart gave Martin no advantage, because he endured greater fatigue riding than his competitors did walking. That is the model for showing that an AI accommodation leaves the standard intact.

The testing regulation. For covered examinations, results must reflect the individual's aptitude or achievement rather than impaired sensory, manual, or speaking skills (28 C.F.R. § 36.309(b)(1)(i)). The exception is an examination that measures those very skills. The *Standards for Educational and Psychological Testing* state the same fairness doctrine for test design (American Educational Research Association et al., 2014), and the author's measurement work already applies it (Puglisi, n.d.-b). The regulation's list of skills does not include written expression or executive function, which limits how far it reaches the cognitive cases.

Adjacent analyses. Several recent works address parts of the education question: Eberhardt (2026) analyzes generative AI writing tools as cognitive scaffolding for documented memory impairment under the ADA and Section 504 in higher education, one impairment in one setting. Morgan (2026) applies *University of Bristol v Abrahart* to argue under United Kingdom law that regulated use of large language models may be a reasonable adjustment where the barrier lies in written expression rather than reasoning. Wright (2026) proposes four criteria (fidelity, non-augmentation, traceability, and attestation) to keep optical character recognition, voice-to-text transcription, and handwriting recognition out of findings of generative AI misconduct. Wright also argues that detector output alone does not meet the balance-of-probabilities standard many institutions apply. The criteria stop at non-generative tools, and this paper's Type II and Type VI analysis covers the generative functions Wright leaves out.

A working institutional model. The Open University permits students who have shared a disability to use generative AI as a reasonable adjustment in any assessment category (The Open University, 2025). The use must not compromise the intended learning outcomes or academic integrity, and the policy does not require that use to be acknowledged. Its own boundary example shows the limit working. Where an assessment measures pronunciation, a student may not use text to speech, because that would defeat the purpose of the task, and the module team is likely to offer an alternative assessment instead. That is a fundamental-alteration boundary drawn in advance.

Education remedies. In *Payan v. Los Angeles Community College District*, No. 24-1809 (9th Cir. March 11, 2026), blind students denied accessible classroom technology could recover for lost educational opportunities under Title II, although emotional-

distress damages were unavailable. The holding binds the Ninth Circuit only. Nationally, *A.J.T. v. Osseo Area Schools*, 605 U.S. 335 (2025), removed a barrier that had stood in front of education claims. A unanimous Court held that ADA and Rehabilitation Act claims based on educational services are subject to the same standards that apply in other disability discrimination contexts. It rejected the heightened showing of bad faith or gross misjudgment that several circuits had required. A student challenging an AI policy under Section 504 or Title II therefore meets the ordinary liability standard. What damages require remains governed by the remedial rules that otherwise apply, set out in Section 14.

The principle that runs through each authority is the same:

Same standard, different means.

10. Disclosure, Visibility, and Perception Harm

The evidence on disclosure falls into four tiers, from written policy to measured reception to automated enforcement.

10.1 Tier one, a direct policy penalty: Medium

Medium's Help Center requires that "any story incorporating AI assistance be clearly labeled as such," and restricts unlabeled AI-assisted text to the author's personal network (Medium, n.d.). AI-generated writing cannot be placed behind the paywall whether it is disclosed or not. The FAQ exempts grammar and spell checkers, outline assistance, and fact verification, and says a disclosure of AI-generated text can be placed within the first two paragraphs. The April 2024 version of the policy put it more strictly, requiring AI-assisted text to be disclosed at the beginning of the story, within the first two paragraphs (Plagiarism Today, 2024). The policy acknowledges that AI assistance can help an author make ideas clearer or write in a second language, yet the page carries no accessibility or disability exception. One platform thereby combines a Type III disclosure requirement with a Type IV distribution penalty.

10.2 Tier two, a recognized exception with a gap: Elsevier and SAGE

Elsevier's policy, updated in June 2026, requires a disclosure statement and exempts basic grammar, spelling, and punctuation checks (Elsevier, 2026). It requires disclosure when an AI tool makes substantive changes to sentence structure or organization. It also exempts AI features within "specialist disability-related assistive technology" used solely for accessibility.

Read together, those clauses place a disabled writer's work on either side of the line depending on the product. The structural and organizational help that a dyslexic writer or a writer with an intellectual disability receives from a general-purpose tool falls on

the disclosure side. The same help from a specialist product does not. That boundary is where the companion paper's thesis begins.

Publishers also disagree about where the line belongs. SAGE classifies tools that improve the language, grammar, or structure of an author's own work as assistive AI that requires no disclosure, and requires disclosure only for generated content (Sage Publishing, n.d.). Elsevier requires disclosure of the same structural help. Two major publishers place structural assistance on opposite sides of the disclosure line, which shows that the boundary is a policy choice.

10.3 Tier three, measured perception harm

Two large experimental programs measured how disclosure is received.

Schilke and Reimann (2025) ran 13 preregistered experiments and found that people who disclose AI use are trusted less. The penalty held when the disclosure stated that a human reviewed and revised the work, and when it stated that AI was used only for proofreading (Study 9). It held under voluntary and mandatory disclosure regimes alike, with a large effect in both (Study 12), and exposure by a third party was worse than self-disclosure (Study 13). The authors suggest that organizations can make disclosure voluntary to protect employees, or can make AI use collectively valid.

Raj, Berg, and Seamans (2026) ran 16 experiments with about 27,000 participants. AI-labeled creative writing was devalued by 6.2% on average, and framing the work as a collaboration did not reliably reduce the penalty. The finding covers creative writing only.

Neither program tested disabled users or disclosure made because of an accommodation. Adnin et al. (2025) examine student and teacher perspectives on undisclosed use of generative AI in academic work, a behavior that disclosure regimes may encourage.

10.4 Tier four, perception enforcement

The LinkedIn record in Section 6.1, with its causal limit, completes the picture. At this tier, no disclosure is needed for a penalty to arrive, because the system acts on how the work reads.

10.5 The disclosure tax and the chain

Johnson, Lewis, Mankoff, and Banner (2026) name disclosure among the accessibility taxes disabled people pay for using generative AI. That disability-led concept names what the first three tiers measure in general populations. The chain runs from disability to assistive AI, then to content produced, then to a perception that the content is AI-made. From there it runs to negative classification or reporting, reduced visibility, and

reduced professional opportunity. The chain does not establish liability. It sets a causation hypothesis: accessibility may fail downstream even where interface access exists upstream.

10.6 Governance labels versus reception

An author-chosen tag that states who holds authority and accountability differs from a label an institution imposes and attaches to a penalty, and the difference lies in consent and control. In the evidence in Tier three, none of the tested disclosure framings removed the trust penalty, and this paper reports that cost as it stands. The paper does not claim that every AI label breaches ADA confidentiality. It asks when a mandatory regime exposes accommodation use, or forces a person to explain a disability in order to escape an AI stigma penalty.

Chosen assistive disclosure already appears in scholarship. Gilly's 2026 repricing paper acknowledges speech-to-text software and Claude as assistive technology for a diagnosed neurodevelopmental condition, and states that all intellectual contributions are the author's own (Gilly, 2026a). That disclosure was the author's choice, framed in the author's terms, which is the distinction this section draws.

10.7 False positives

Detection tools misread some human writing as machine writing. Liang et al. (2023) found that detectors consistently misclassified non-native English writing as AI-generated, and concluded that detectors may penalize writers with constrained linguistic expressions. The open question is what that mechanism means for disabled writers whose prose is repetitive or formulaic, or restructured by assistive tools.

11. Rhetoric and People

Policies are not the only source of harm, because people also speak and act on their views of AI.

Coworkers and supervisors may comment on a colleague's AI use, and an employer can answer for a hostile environment in some circumstances. A disability nexus is required, and criticism of AI use alone does not qualify. Public accusations that a person's work is AI-made, including false ones, may raise state law questions that vary by jurisdiction. Federal law adds one direct hook. Section 12203 bars retaliation against a person who opposes an act the ADA makes unlawful. It also makes it unlawful to coerce, intimidate, threaten, or interfere with any individual in the exercise or enjoyment of rights the ADA grants or protects (42 U.S.C. § 12203(a) and (b)).

Individuals are rarely reached directly by federal disability law, and this section says so plainly. Each hook above applies only where the conduct connects to rights the statute protects.

12. The Difference Between Regulating Outcomes and Regulating Access

The conceptual center of the paper is the difference between two kinds of rule.

Rule A: Every citation must be accurate. **Rule B:** You may not use AI to locate or organize sources.

Rule A: The student must show independent understanding. **Rule B:** No AI may help organize the student's own ideas.

Rule A governs the quality of the outcome, while Rule B governs the means. Exposure grows more plausible as a policy moves away from substantive standards and toward restricting the mechanism a person needs for access. The exception is a means that is itself an essential function, or the very construct being assessed.

Disability accommodation need not lower the standard. It may require allowing a different means of meeting it.

The legal floor is not the ethical ceiling. Disability law reaches only people who meet its definitions, and many people with partial or episodic limitations fall short of them. Removing AI from a person whose disability is partial may create no legal exposure at all. It still raises a moral and ethical question, because the barrier is real even where the statute does not reach it. The principle of same standard, different means applies as a matter of fairness before it applies as a matter of law.

This matches the author's published reading of New York's Part 161 rule on AI in court filings (Puglisi, 2026c). That rule moved past the disclosure fight to human accountability: a lawyer may use AI without routine disclosure, and the signature still certifies the paper. It regulates the outcome and leaves the means to the accountable person.

13. The AI Accessibility Exposure Screen for Accommodation

The screen below is a governance tool for the accommodation question. It runs as a gate sequence, in which failure at an earlier gate ends the federal accommodation analysis for that actor and turns the question to state law. Other theories carry their own elements and do not pass through these gates. Retaliation and interference under 42 U.S.C. §

12203 protect any individual, whether or not that person meets the disability definition at gate 2, and disparate treatment and confidentiality claims run on their own terms.

1. **Covered actor.** Which statute, if any, reaches this institution?
2. **Protected person.** Does the person have a disability under that statute, and is the person a qualified individual where the statute requires it?
3. **Disability-related functional limitation.** What limitation does the AI use address?
4. **Accommodation or modification duty.** Has a request been made or an interactive process begun, and does a duty attach?
5. **Responsive AI use.** Does the AI use meet the companion paper's functional test for that limitation?
6. **Policy conflict.** Does the AI rule prohibit, penalize, expose, or diminish that use?
7. **Individualized defense analysis.** Would the use impose undue hardship, fundamentally alter the program, displace an essential function, or pose a direct threat? Does a legitimate legal or security need apply, and is an effective alternative available?

The seven-gate screen for accommodation exposure



Figure 6. *The seven-gate screen for accommodation exposure. Failure at an earlier gate ends the federal accommodation analysis for that actor.*

Only after all seven gates does an accommodation-based exposure take shape under this screen.

Governance triage. Institutions can place their current policies on a simple ladder:

- **Low:** the policy includes an accommodation exception, individualized review, and alternatives.
- **Moderate:** the policy does not mention accommodation, but an existing process can handle requests.
- **Elevated:** the policy is a blanket prohibition with no pathway for requests.
- **High:** the institution refuses to consider modification because the rule applies to everyone.
- **Severe:** the institution takes adverse action after a request.
- **Critical:** the institution discloses accommodation information publicly, or retaliates.

These labels describe governance escalation and make no estimate of the probability or size of any legal liability.

14. What Exposure Could Look Like

Remedies differ by title and by defendant, and the paper keeps them separate.

- **Title I** allows equitable relief, back pay, front pay, attorney's fees, and compensatory and punitive damages within statutory caps. *Garrett* bars state employees from recovering money damages against states under Title I.
- **Title III** gives private plaintiffs injunctive relief under federal law. Department of Justice enforcement and state statutes can add more.
- **Education claims** carry no heightened standard of proof after *A.J.T. v. Osseo Area Schools*, 605 U.S. 335 (2025).
- **Section 504** carries the limits of Spending Clause remedies. **Title II** is not Spending Clause legislation, but it incorporates the Rehabilitation Act's remedies by statute (42 U.S.C. § 12133), so those limits carry through to Title II by incorporation. *Barnes v. Gorman*, 536 U.S. 181 (2002), bars punitive damages in private suits under Title II and Section 504. *Cummings v. Premier Rehab Keller, P.L.L.C.*, 596 U.S. 212 (2022), bars emotional-distress damages in private actions under the Rehabilitation Act and the Affordable Care Act, which reaches Section 1557 health programs. The Ninth Circuit in *Payan* applied the same bar to Title II and held that lost educational opportunities remain compensable.

The three channels. The author's *Liability Map* describes three channels through which AI creates legal exposure: regulatory enforcement, civil liability, and contract and insurance (Puglisi, 2026b). All three ask for the same artifact, a record that can be produced. In this setting, that record is the accommodation file that shows individualized review took place. Employment practices liability coverage belongs in the third channel, where an insurer will ask whether the file exists.

What the file should record. The author's documentation protocol offers a ready structure for that file. Its decision taxonomy sorts each human act at a checkpoint into four kinds (Puglisi, n.d.-c). A corrective override means the human caught an error, and a creative supersession means the human produced something better than any platform output. A checkpoint confirmation means the human actively reviewed and approved, and a deferred decision means the human saw the item and chose not to resolve it yet. An accommodation file built that way shows individualized review on its face. The same protocol carries its own caution, which applies here. Preserved dissent may be discoverable, subject to relevance, proportionality, privilege, and other limits, and a record showing that a reviewer overrode a recommendation can be read against the

institution later. That is not a reason to keep no record. It is a reason to make every confirmation carry rationale that matches the depth of review it actually received.

The paper uses the term potential exposure throughout, and it predicts no outcome in any named case.

15. The European Comparison

15.1 Protections that resemble the ADA

Directive 2000/78 carries an employment-focused reasonable-accommodation duty that resembles ADA employment accommodation in several respects. Directive 2000/78, Article 5, provides that employers “shall take appropriate measures, where needed in a particular case, to enable a person with a disability to have access to, participate in, or advance in employment, or to undergo training, unless such measures would impose a disproportionate burden on the employer.” Recital 20 lists adapting premises and equipment among appropriate measures. Article 2(2)(b)(ii) ties indirect discrimination on grounds of disability to the measures Article 5 requires, which is the European analog of the facial neutrality argument in Section 5. The Court of Justice in *HR Rail*, Case C-485/20 (February 10, 2022), read the duty individually and broadly, holding that it can require reassignment unless that imposes a disproportionate burden.

The CRPD, to which the European Union is a party, states in Article 2 that discrimination on the basis of disability includes denial of reasonable accommodation. The European Accessibility Act covers specified categories of products and services, including consumer general-purpose computers and operating systems, smartphones and terminal equipment, e-readers, electronic communications, consumer banking, e-books, and e-commerce. Generative AI services are not a listed category as such.

15.2 Where the AI Act meets them

Article 50 of the AI Act applies from August 2, 2026. The Commission published its Guidelines on the transparency obligations on July 20, 2026, as C(2026) 5054 final, Annex (European Commission, 2026). The Regulation’s own text in Article 50(2) contains no disability language, while the Guidelines do.

Provider marking under Article 50(2). The statute exempts AI systems that perform an assistive function for standard editing, and outputs that do not substantially alter the input or its semantics. The Guidelines explain both exceptions:

- Point (90) describes standard editing as preparing existing content for publication through small edits to readability, grammar, quality, and format, without generating new content. Editing goes beyond standard editing when it changes meaning, style, or intent in a material way.

- Point (91) treats the second exception, no substantial alteration of the input or its semantics, as a case-by-case assessment.

Among the examples that fall outside marking, the annex lists “grammar correction and spellchecking, linguistic and minor stylistic polishing that do not change the substance, meaning, style or messaging of text, AI-generated translations of text.”

The annex separately lists “AI-generated content that only transforms authentic human input through assistive technologies allowing persons with disabilities to communicate (e.g., augmentative and alternative communication (AAC) or customized neural voices (CNV)) since they do not alter semantically the meaning of the content.”

Among the examples that require marking, it lists “AI-generated summaries of text; paraphrasing or rewriting text that changes style, structure and meaning beyond mere grammatical and minor stylistic correction.”

The cut. The European Union now names a disability example, and the example is narrow. It covers the transformation of authentic human input that preserves meaning, which fits AAC, customized neural voices, and the companion paper’s speech-loss case. Summarizing, substantive rewriting, and restructuring used to offset a disability-related barrier sit on the marking side, and drafting and outlining can too where the output generates or materially alters content. Those are the functions the dyslexia and intellectual disability cases rely on. The cut matches the one Elsevier drew between specialist and general-purpose tools, now in European soft law. Article 50 is not a disability exemption for general-purpose drafting, and this paper does not describe it as one.

Deployer labelling under Article 50(4). The deployer duty covers text published to inform the public on matters of public interest, and it carries no standard-editing exception. The exception it does carry requires human review or editorial control, with a person holding editorial responsibility. The Guidelines require substantive review by persons with relevant knowledge, with fact-checking as a minimum. Spell-checking, an editorial policy that exists only on paper, automated review, and cursory approval do not qualify, and a substantive AI intervention after review voids the exception. That is accountability in place of labelling, the same pattern as New York’s Part 161 rule and the author’s definition of #AIassisted.

Audience and accessibility. The Guidelines state that the obviousness exception for interactive systems cannot be relied upon where the foreseeable audience includes persons with disabilities. Article 50(5) requires that the transparency information itself meet accessibility requirements, by reference to Directives (EU) 2016/2102 and 2019/882. That is a duty about the design of the label, and it creates no accommodation duty for the user’s own tool.

Transition. The Digital Omnibus gives generative systems already on the market until December 2, 2026 to meet the marking duty, and defers the interoperability of watermark detection to February 2, 2027. The same regulation separately postponed the high-risk duties in Chapter III, to December 2, 2027 for stand-alone high-risk systems and August 2, 2028 for high-risk AI embedded in regulated products (Regulation (EU) 2026/1744; White & Case LLP, 2026). Those later dates govern high-risk systems and do not change the Article 50 timeline.

European publishers and platforms therefore sit in a gap much like the American one. Accommodation duties concentrate in employment, while the AI Act's exemptions stop at meaning preservation.

16. The Governance Fix: An Accommodation Layer

Most rules about AI ask only whether AI is allowed. A disability-aware policy also asks what happens when a prohibited AI use is requested as disability-related assistance.

16.1 Minimum architecture

A policy that answers that second question carries eight elements:

- an accommodation exception, or a cross-reference to the accommodation process
- individualized review
- documentation proportional to need
- approved secure alternatives
- confidential handling of accommodation information
- an appeal path
- separation of accommodation from ordinary AI enforcement
- human review before any adverse action based on detection or reports

The review has to be real. A human placed near an automated flag does not by that fact govern it. Checkpoint-Based Governance draws the line in its canonical text: “Human in the loop means the person is present and participating. It does not require that the person holds authority or answers for the outcome” (Puglisi, n.d.-c). The same framework prohibits one AI from approving another AI's output under its own governing rules, stating that no platform count, confidence score, or level of agreement substitutes for human arbitration (Puglisi, n.d.-c). A detection score reviewed by a person who lacks the time, the evidence, or the authority to overturn it is the arrangement Elish (2019) calls a moral crumple zone, in which the human absorbs responsibility for a system the human did not control. Crootof, Kaminski, and Price (2023) reach a related conclusion: the law should regulate the human-in-the-loop system as a whole and test whether the human in it can act. The framework also makes

rubber-stamping measurable through audit triggers it defines itself. Approval rates above 95 percent, or decision reversals below 2 percent, sustained across three consecutive cycles, trigger a mandatory audit (Puglisi, n.d.-c). An institution that reviews detection flags before discipline can apply the same test to its own reviewers.

16.2 Drafting principle

Regulate unacceptable outcomes and risks first. Restrict assistive means only when necessary.

A policy can require authentic scholarship, accurate citations, protection of confidential information, no impersonation, no undisclosed synthetic evidence, and human accountability. Before banning a method, the drafter should ask whether the method may also be needed for disability access.

16.3 Model clause

Nothing in this AI policy is intended to restrict the use of AI or AI-enabled functionality approved as a reasonable accommodation or other disability-related assistive technology. Requests for such use will be evaluated through the organization's applicable disability accommodation process, subject to legitimate requirements concerning essential functions, security, confidentiality, academic integrity, and fundamental alteration.

The clause guarantees no request. It guarantees that a request is treated as an accessibility question from the start.

16.4 The principle is already stated, and the practice already exists

The American Foundation for the Blind's *Guiding Principles* (2025) state that students and employees with disabilities should be able to use AI as assistive technology. They also say that schools and employers should consult the disability community when developing policies on AI as a reasonable accommodation. The Open University and Defra already grant general-purpose AI as a reasonable adjustment, and the Open University exempts that use from its acknowledgment requirement. The model clause states that practice as policy language any institution can adopt.

16.5 The two boundaries the clause closes

The clause closes two gaps that existing rules leave open.

- **Elsevier** exempts AI features within “specialist disability-related assistive technology” used solely for accessibility, while requiring disclosure when an AI tool makes substantive changes to sentence structure or organization (Elsevier, 2026). The exemption does not reach a general-purpose tool making structural edits for a disabled author.

- **The EU Guidelines** exempt “AI-generated content that only transforms authentic human input through assistive technologies allowing persons with disabilities to communicate,” while marking “AI-generated summaries of text” and rewriting that “changes style, structure and meaning” (European Commission, 2026). The exemption does not reach summarizing and restructuring used to offset a cognitive barrier.

The model clause reaches both, because it turns on the accommodation process, whatever the category of the product.

17. Research and Litigation Questions

The following questions remain open:

- What will *Cedeno* decide, and how far will its reasoning travel beyond wearables?
- When does general-purpose software qualify as an effective accommodation?
- Can a disclosure requirement produce discriminatory consequences for people with disabilities?
- Can downranking of assistive AI output amount to unequal access or a denial of benefits?
- Can platforms separate substitution from disability-mediated assistance without forcing users to disclose a disability?
- What evidence shows that an AI prohibition is truly necessary to preserve an essential function?
- What counts as an effective alternative when a particular tool is prohibited?
- How should institutions treat false-positive AI detection involving disabled users?
- Does the disclosure penalty differ for disabled users, or for disclosure made because of an accommodation?
- How does machine-readable marking in the European Union interact with accommodation duties?
- When AI is used to triage accommodation requests themselves, who reviews its output, and what record shows individualized review?

The last question is no longer hypothetical. Government Executive reported in August 2026 that an internal Department of Labor email described plans to use AI to triage a backlog of hundreds of accommodation requests (Government Executive, 2026a). The same month, the department notified accommodation requesters that some of their health information had been sent outside its network (Government Executive, 2026b). That incident is the confidentiality risk in Section 16 in live form. Under Checkpoint-Based Governance, an AI triage decision is not a governed decision until a named

human completes it, and the record should show who decided, on what evidence, and why (Puglisi, n.d.-c).

18. Conclusion: AI Policy Has an Accessibility Cost

The question this paper asked has a two-part answer. Where a restriction removes a tool that functions as assistive technology for a person with a disability, and no limit the law already keeps justifies it, the restriction is unethical at best, whether or not any statute reaches it. Where disability law does reach the actor and the person, the same restriction may also be illegal, and the actor matrix and the accommodation screen show where that line is likely to fall.

The person remains capable, and disability changes the route available to that capability. AI may restore another route, and the institution then writes a rule governing it. At that moment an AI policy can become a barrier to disability access, whether or not its authors intended it. The remedy is AI Governance in its proper sense, human oversight and accountability, applied to the accommodation question.

Organizations remain free to regulate legitimate risks associated with artificial intelligence. But once AI is functioning as assistive technology for a person with a disability, a policy that prohibits, penalizes, exposes, or diminishes that use may implicate disability-law obligations that cannot be resolved merely by calling the rule neutral or applying it to everyone.

The governing principle is short enough to write into any policy: same standard, different means. The governance principle proposed here is that the standard protects what the institution needs to protect, and the means belongs to the person who needs it.

The cause for concern is therefore not that every anti-AI rule violates the ADA. It is that organizations are writing AI policies today that may control disability-related access tomorrow, often without recognizing that the accommodation question exists.

References

Primary sources

- Adnin, R., Pandkar, A., Yao, B., Wang, D., & Das, M. (2025). Examining student and teacher perspectives on undisclosed use of generative AI in academic work. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*.
- American Educational Research Association, American Psychological Association, & National Council on Measurement in Education. (2014). *Standards for educational and psychological testing*. American Educational Research Association.
- American Foundation for the Blind. (2025, January). *Guiding principles for more disability-inclusive AI*. <https://afb.org/research-and-initiatives/empowering-or-excluding/guiding-principles-more-disability-inclusive-ai>
- Associated Press. (2026, September 14). Trump calls AI risks a “hoax,” says there is a “SICK conspiracy” against AI and data centers. <https://www.local10.com/business/2026/09/14/trump-calls-ai-risks-a-hoax-says-there-is-a-sick-conspiracy-against-ai-and-data-centers/>
- Brand, D., Fischer Mogensen, K., & Engelbrecht, M. (2026, September 15). AI as reasonable accommodation and disability employment. *Disability & Society*. <https://doi.org/10.1080/09687599.2026.2726868>
- Buolamwini, J. (2026, April). [Comment on a LinkedIn post by B. C. Puglisi concerning the evocative audit]. LinkedIn.
- Cisco. (2024). *Cisco 2024 data privacy benchmark study*. Cisco Systems.
- Cisco. (2026). *Cisco 2026 data privacy benchmark study*. https://www.cisco.com/c/dam/en_us/about/doing_business/trust-center/docs/cisco-privacy-benchmark-study-2026.pdf
- Cooley LLP. (2025, February 21). *Gone but not forgotten: Federal laws still apply despite guidance disappearance act*. <https://www.cooley.com/news/insight/2025/2025-02-21-gone-but-not-forgotten-federal-laws-still-apply-despite-guidance-disappearance-act>
- Crootof, R., Kaminski, M. E., & Price, W. N., II. (2023). Humans in the loop. *Vanderbilt Law Review*, 76(2), 429 to 510.
- Eberhardt, K. E. (2026, July 7). *Position paper: AI as a disability accommodation in higher education. Legal analysis and annotated bibliography on cognitive scaffolding and generative AI*. Zenodo. <https://doi.org/10.5281/zenodo.21231449>
- Ede, R. (2026, August 20). A thinking partner: Copilot as assistive technology. *Defra digital, data, technology and security blog*.

<https://defradigital.blog.gov.uk/2026/08/20/a-thinking-partner-copilot-as-assistive-technology/>

- Elish, M. C. (2019). Moral crumple zones: Cautionary tales in human-robot interaction. *Engaging Science, Technology, and Society*, 5, 40 to 60. <https://doi.org/10.17351/ests2019.260>
- Elsevier. (2026, June). *Generative AI policies for journals*. <https://www.elsevier.com/about/policies-and-standards/generative-ai-policies-for-journals>
- European Commission. (2026, July 20). *Guidelines on the implementation of the transparency obligations for certain AI systems under Article 50 of Regulation (EU) 2024/1689 (C(2026) 5054 final, Annex)*. <https://digital-strategy.ec.europa.eu/en/library/guidelines-transparency-obligations-providers-and-deployers-ai-systems>
- Gallup. (2026). *Americans cool toward AI*. <https://news.gallup.com/poll/712751/americans-cool-toward.aspx>
- GCN. (2026, August). LinkedIn slop flag used by a million members. <https://gcn.com/linkedin-slop-flag-used-million-members/21090/>
- Gilly, T. (2026a, August). *Accommodation repricing: Disability, generative AI subscriptions, and the accessibility tax as a variable the vendor controls (vo.7)*. Real Safety AI Foundation. https://realsafetyai.org/documents/accommodation_repricing_vo_7.pdf
- Gilly, T. (2026b). *The accommodation axis: Disability discrimination, algorithmic opacity, and the deployers that anti-discrimination scholarship overlooks*. SSRN 7011798. https://papers.ssrn.com/sol3/papers.cfm?abstract_id=7011798
- Government Executive. (2026a, August). Labor looks to AI to tackle accommodation requests from disabled employees. <https://www.govexec.com/technology/2026/08/labor-department-seeks-ai-triage-accommodation-requests-disabled-employees/415438/>
- Government Executive. (2026b, August). Data from Labor employees seeking disability accommodations impacted by “internal incident.” <https://www.govexec.com/management/2026/08/data-labor-employees-seeking-disability-accommodations-impacted-internal-incident/415673/>
- Hanuschock, W. E., & Silverman, A. M. (2026, August). *Innovation for access: AI and its use as assistive technology for people with disabilities*. American Foundation for the Blind. <https://afb.org/research-and-initiatives/ai-series/innovation-access>
- ILRU. (n.d.). *Public entities: Reasonable modifications and fundamental alteration*. <https://www.ilru.org/publications/public-entities-reasonable-modifications-fundamental-alteration>

- International Committee of Medical Journal Editors. (n.d.). *Use of AI by authors*. <https://icmje.org/recommendations/browse/artificial-intelligence/ai-use-by-authors.html>
- Johnson, J., Lewis, A., Mankoff, J., & Banner, O. (2026). “I don’t trust it, but I use it”: Navigating trust, privacy, and identity in disabled people’s use of generative AI. In *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3772318.3790652>
- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*. <https://doi.org/10.1016/j.patter.2023.100779>
- Medium. (n.d.). *Artificial intelligence (AI) content policy*. Medium Help Center. <https://help.medium.com/hc/en-us/articles/22576852947223-Artificial-Intelligence-AI-content-policy>
- Morgan, D. (2026). Reasonable adjustments for neurodivergent students in higher education: Generative AI, accessibility, and mediated authorship. *Disability & Society*. <https://doi.org/10.1080/09687599.2026.2667528>
- The Open University. (2025, July 15). *Generative AI for students*. <https://about.open.ac.uk/policies-and-reports/policies-and-statements/generative-ai-learning-teaching-and-assessment-ou-o>
- Perkins Coie. (n.d.). *Eleventh Circuit vacates ruling that websites are not public accommodations under ADA*. <https://legacy.perkinscoie.com/insights/blog/eleventh-circuit-vacates-ruling-websites-are-not-public-accommodations-under-ada>
- Pew Research Center. (2025, September 17). *How Americans view AI and its impact on people and society*. https://www.pewresearch.org/wp-content/uploads/sites/20/2025/09/PS_2025.9.15_AI-and-its-impact_report.pdf
- Pew Research Center. (2026, March 12). *Key findings about how Americans view artificial intelligence*. <https://www.pewresearch.org/short-reads/2026/03/12/key-findings-about-how-americans-view-artificial-intelligence/>
- Phillips, B., & Zhao, H. (1993). Predictors of assistive technology abandonment. *Assistive Technology*, 5(1), 36 to 45. <https://doi.org/10.1080/10400435.1993.10132205>
- Plagiarism Today. (2024, April 11). *Medium sets new policies on AI-generated writing*. <https://www.plagiarismtoday.com/2024/04/11/medium-sets-new-policies-on-ai-generated-writing/>
- Quinnipiac University Poll. (2026, March 30). <https://poll.qu.edu/poll-release?releaseid=3955>

- Raj, M., Berg, J. M., & Seamans, R. (2026). The artificial intelligence disclosure penalty: Humans persistently devalue AI-generated creative writing. *Journal of Experimental Psychology: General*. <https://doi.org/10.1037/xge0001889>
- Sage Publishing. (n.d.). *Artificial intelligence policy*. <https://www.sagepub.com/journals/publication-ethics-policies/artificial-intelligence-policy>
- Schilke, O., & Reimann, M. (2025). The transparency dilemma: How AI disclosure erodes trust. *Organizational Behavior and Human Decision Processes*, 188, 104405. <https://doi.org/10.1016/j.obhdp.2025.104405>
- Sheridan College. (n.d.). *Turnitin AI writing detection*. <https://ltsa.sheridancollege.ca/digital-learning-support-hub/turnitin-ai-writing-detection/>
- Shock, A., & Silverman, A. M. (2026, June). *Working with the machine: AI's expanding role in employment for people with disabilities*. American Foundation for the Blind. <https://afb.org/research-and-initiatives/ai-series/working-machine>
- Trump, D. J. (2026, September 14). [Post]. Truth Social. As reported by the Associated Press (2026).
- White & Case LLP. (2026, August 4). *EU AI Omnibus enters into force, amending the AI Act*. <https://www.whitecase.com/insight-alert/eu-ai-omnibus-enters-force-amending-ai-act>
- Wright, C. (2026). Transcription is not generation: Distinguishing non-generative AI tool use from academic misconduct in higher education assessment. *International Journal for Educational Integrity*, 22, 24. <https://doi.org/10.1007/s40979-026-00234-w>

Cases

- *A.J.T. v. Osseo Area Schools*, 605 U.S. 335 (2025).
- *Barnes v. Gorman*, 536 U.S. 181 (2002).
- *Bledsoe v. Palm Beach County Soil & Water Conservation District*, 133 F.3d 816 (11th Cir. 1998).
- *Board of Trustees of the University of Alabama v. Garrett*, 531 U.S. 356 (2001).
- *Bruno Cedeno v. Walt Disney Parks and Resorts U.S., Inc.*, No. 6:25-cv-02046-GAP-DCI (M.D. Fla. filed Oct. 23, 2025), Amended Complaint and Demand for Jury Trial, Doc. 23 (filed Jan. 7, 2026). <https://storage.courtlistener.com/recap/gov.uscourts.flmd.449095/gov.uscourts.flmd.449095.23.0.pdf> <https://brodyandassociates.com/smart-glasses-complicated-questions-navigating-ai-privacy-and-the-ada/>
- *Cummings v. Premier Rehab Keller, P.L.L.C.*, 596 U.S. 212 (2022).

- *Elwell v. Oklahoma ex rel. Board of Regents of the University of Oklahoma*, 693 F.3d 1303 (10th Cir. 2012).
- *Gil v. Winn-Dixie Stores, Inc.*, 993 F.3d 1266 (11th Cir. 2021), vacated as moot (Dec. 28, 2021).
- *Mobley v. Workday, Inc.*, No. 3:23-cv-00770-RFL (N.D. Cal. May 16, 2025) (order on collective certification).
https://www.govinfo.gov/content/pkg/USCOURTS-cand-3_23-cv-00770/pdf/USCOURTS-cand-3_23-cv-00770-1.pdf
- *Payan v. Los Angeles Community College District*, No. 24-1809 (9th Cir. Mar. 11, 2026). <https://cdn.ca9.uscourts.gov/datastore/opinions/2026/03/11/24-1809.pdf>
- *PGA Tour, Inc. v. Martin*, 532 U.S. 661 (2001).
- *Skinner v. Salem School District*, No. 09-cv-193-JL, 2010 DNH 106 (D.N.H. June 18, 2010).
<https://www.nhd.uscourts.gov/sites/default/files/Opinions/10/10NH106P.pdf>
- *US Airways, Inc. v. Barnett*, 535 U.S. 391 (2002).
- *Wong v. Regents of the University of California*, 192 F.3d 807 (9th Cir. 1999).
- *Zimmerman v. Oregon Department of Justice*, 170 F.3d 1169 (9th Cir. 1999).
- *XXXX v. HR Rail SA*, Case C-485/20 (CJEU Feb. 10, 2022).

Statutes, regulations, and instruments

- 20 U.S.C. § 1401 (Individuals with Disabilities Education Act, definitions).
- 29 U.S.C. § 794 (Rehabilitation Act, Section 504).
- 29 U.S.C. § 3002 (Assistive Technology Act, definitions).
- 42 U.S.C. §§ 12102, 12182, 12189, 12201, and 12203 (Americans with Disabilities Act).
- 28 C.F.R. §§ 35.130, 36.302, and 36.309 (Department of Justice regulations under ADA Titles II and III).
- 29 C.F.R. § 1630.2 (EEOC regulations under ADA Title I).
- 45 C.F.R. § 92.205 (Section 1557, reasonable modifications in health programs).
- Council Directive 2000/78/EC of 27 November 2000 establishing a general framework for equal treatment in employment and occupation.
- Directive (EU) 2019/882 (European Accessibility Act).
- Regulation (EU) 2024/1689 (Artificial Intelligence Act), Article 50.
- Regulation (EU) 2026/1744 (Digital Omnibus on AI).
- United Nations. (2006). *Convention on the Rights of Persons with Disabilities*.

Works by the author, with the sources supporting each

- **Puglisi, B. C. (n.d.-a). *HAIA: The Human Artificial Intelligence Assistant ecosystem* [Canonical framework page]. basilpuglisi.com. <https://basilpuglisi.com/haia-the-human-artificial-intelligence-assistant-ecosystem/>**
 - Atari, M., Xue, M. J., Park, P. S., Blasi, D. E., & Henrich, J. (2023). Which humans? PsyArXiv preprint. <https://doi.org/10.31234/osf.io/5b26t>
 - Henrich, J., Heine, S. J., & Norenzayan, A. (2010). The weirdest people in the world? *Behavioral and Brain Sciences*, 33(2-3), 61-83.
 - Puglisi, B. C. (2012). *Digital Factics: Twitter*. Digital Media Press.
 - Puglisi, B. C. (2025). *Governing AI: When Capability Exceeds Control*. ISBN 9798349677687.
 - Puglisi, B. C. (2025). *Ethics of artificial intelligence: A white paper on principles, risks, and responsibility*. Published August 18, 2025. <https://basilpuglisi.com/ethics-of-artificial-intelligence/>
 - Puglisi, B. C. (n.d.). *AI Provider Plurality: A congressional package* [AI Policy page]. <https://basilpuglisi.com/ai-provider-plurality-a-congressional-package/>
 - Puglisi, B. C. (n.d.). *Checkpoint-Based Governance* [Canonical framework page]. <https://basilpuglisi.com/cbg/>
 - Puglisi, B. C. (n.d.). *HAIA-CAIPR* [Canonical framework page]. <https://basilpuglisi.com/haia-caipr/>
 - Puglisi, B. C. (2026). *HAIA-RECCLIN case study 006* (v7). <https://basilpuglisi.com>
 - Puglisi, B. C. (n.d.). *HAIA-RECCLIN Reasoning and Dispatch* [Canonical framework page]. <https://basilpuglisi.com/haia-recclin/>
 - Puglisi, B. C. (n.d.). *HEQ and AIS* [Canonical framework page]. <https://basilpuglisi.com/heq-ais/>
 - Puglisi, B. C. (2026). *The loop that ate the governor* (v7). <https://basilpuglisi.com>
 - Puglisi, B. C. (2026). *The minds that bend the machine: The voices shaping responsible AI governance* (forthcoming, October 2026). <https://basilpuglisi.com>
- **Puglisi, B. C. (n.d.-b). *HEQ and AIS: Measuring governance of AI literacy* [Canonical framework page]. basilpuglisi.com. <https://basilpuglisi.com/heq-ais/>**
 - Colorado General Assembly. (2026, May 14). *Senate Bill 26-189: Consumer protections in interactions with artificial intelligence systems*. Signed May 14, 2026; effective January 1, 2027.

- European Parliament and Council. (2024). *Regulation (EU) 2024/1689 laying down harmonised rules on artificial intelligence (AI Act)*. Official Journal of the European Union.
- European Parliament and Council. (2026). *Regulation (EU) 2026/1744 amending Regulation (EU) 2024/1689 (Digital Omnibus on AI)*. Official Journal of the European Union, July 24, 2026; in force July 27, 2026.
- International Organization for Standardization. (2023). *ISO/IEC 42001:2023, Information technology, Artificial intelligence, Management system*. ISO.
- Landgericht München I. (2026, May 28). *Case No. 26 O 869/26*. Preliminary injunction in expedited proceedings; on appeal.
- Lokken v. UnitedHealth Group, Inc., No. 0:23-cv-03514 (D. Minn.). Discovery order, March 9, 2026.
- National Institute of Standards and Technology. (2023). *Artificial intelligence risk management framework (AI RMF 1.0)*. NIST AI 100-1.
- Organisation for Economic Co-operation and Development. (2019, updated 2024). *OECD AI Principles*. oecd.ai
- AERA, APA, & NCME. (2014). *Standards for educational and psychological testing*. American Educational Research Association.
- Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakcı, Ö., & Mariman, R. (2025). Generative AI without guardrails can harm learning: Evidence from high school mathematics. *Proceedings of the National Academy of Sciences*, 122(26), e2422633122.
<https://doi.org/10.1073/pnas.2422633122>
- Buçinca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To trust or to think: Cognitive forcing functions can reduce overreliance on AI. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1), Article 188.
<https://doi.org/10.1145/3449287>
- Ganuthula, V. R. R., & Balaraman, K. K. (2025). Artificial intelligence quotient framework for measuring human collaboration with artificial intelligence. *Discover Artificial Intelligence*, 5, Article 268.
<https://doi.org/10.1007/s44163-025-00516-1>
- Gerlich, M. (2025a). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), Article 6.
<https://doi.org/10.3390/soc15010006>
- Hollnagel, E., & Woods, D. D. (2005). *Joint cognitive systems: Foundations of cognitive systems engineering*. CRC Press.
- Hutchins, E. (1995). *Cognition in the wild*. MIT Press.

- Insurance Business. (2026, February 11). Chaucer, Armilla bet big on standalone AI coverage as reinsurers struggle to keep pace. insurancebusinessmag.com
- The Insurer. (2026, April 24). Standalone AI liability market takes shape with underwriting discipline key to MGA success. theinsurer.com
- Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50 to 80. https://doi.org/10.1518/hfes.46.1.50_30392
- Legg, S., & Hutter, M. (2007). Universal intelligence: A definition of machine intelligence. *Minds and Machines*, 17(4), 391 to 444. <https://doi.org/10.1007/s11023-007-9079-x>
- Lior, A. (2025). E/Insuring the AI age: Empirical insights into artificial intelligence liability policies. *Connecticut Insurance Law Journal*, 31, 99. SSRN Abstract 5316376.
- Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (pp. 1 to 16). Association for Computing Machinery. <https://doi.org/10.1145/3313831.3376727>
- National Association of Insurance Commissioners. (2023, December 4). *Model bulletin on the use of artificial intelligence systems by insurers*. content.naic.org
- Ng, D. T. K., Leung, J. K. L., Chu, S. K. W., & Qiao, M. S. (2021). Conceptualizing AI literacy: An exploratory review. *Computers and Education: Artificial Intelligence*, 2, Article 100041. <https://doi.org/10.1016/j.caeai.2021.100041>
- OECD & European Commission. (2026). *Empowering learners for the age of AI: An AI literacy framework for primary and secondary education*. OECD Publishing. <https://doi.org/10.1787/65cd27d4-en>
- Risko, E. F., & Gilbert, S. J. (2016). Cognitive offloading. *Trends in Cognitive Sciences*, 20(9), 676 to 688. <https://doi.org/10.1016/j.tics.2016.07.002>
- ScienceSoft. (2026, September 10). AI risks to enter 60 to 80% of liability and cyber insurance underwriting by 2028. [GlobeNewswire](https://www.globenewswire.com).
- Sidra, S., & Mason, C. (2025). Generative AI in human-AI collaboration: Validation of the Collaborative AI Literacy and Collaborative AI Metacognition Scales for effective use. *International Journal of Human-Computer Interaction*. <https://doi.org/10.1080/10447318.2025.2543997>
- Sparrow, B., Liu, J., & Wegner, D. M. (2011). Google effects on memory. *Science*, 333(6043), 776 to 778. <https://doi.org/10.1126/science.1207745>

- U.S. Department of Labor. (2026). *Training and Employment Notice No. 07-25: The U.S. Department of Labor’s Artificial Intelligence Literacy Framework*. <https://www.dol.gov/agencies/eta/advisories/ten-07-25>
- UNESCO. (2024). AI competency framework for students. United Nations Educational, Scientific and Cultural Organization. <https://doi.org/10.54675/JKJB9835>
- Vaccaro, M., Almaatouq, A., & Malone, T. (2024). When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nature Human Behaviour*, 8(12), 2293 to 2303. <https://doi.org/10.1038/s41562-024-02024-1>
- Vasconcelos, H., Jörke, M., Grunde-McLaughlin, M., Gerstenberg, T., Bernstein, M. S., & Krishna, R. (2023). Explanations can reduce overreliance on AI systems during decision-making. *Proceedings of the ACM on Human-Computer Interaction*, 7(CSCW1), Article 129. <https://doi.org/10.1145/3579605>
- **Puglisi, B. C. (n.d.-c). *What is Checkpoint-Based Governance? Human authority [Canonical framework page]*. basilpuglisi.com. <https://basilpuglisi.com/cbg/>**
 - Asimov, I. (1942). Runaround. Astounding Science Fiction.
 - Asimov, I. (1985). Robots and Empire. Doubleday.
 - Puglisi, B. C. (2026). HAIA-CARCS: Compliance accountability record and case study. basilpuglisi.com.
 - Puglisi, B. C. (2026). Why you cannot program or prompt governance into AI. basilpuglisi.com.
 - Puglisi, B. C. (2026). HAIA specifications. <https://github.com/basilpuglisi>
- **Puglisi, B. C. (2026a, January 1). *Ethical AI, Responsible AI, and AI Governance are not the same thing*. basilpuglisi.com. <https://basilpuglisi.com/what-we-failed-to-define-is-how-we-fail/>**
 - Bai, Y., et al. (2022). Constitutional AI: Harmlessness from AI feedback. arXiv. <https://arxiv.org/abs/2212.08073>
 - Center for AI Safety. (2023, May 30). Statement on AI risk. <https://safe.ai/work/press-release-ai-risk>
 - CGTN. (2024, December 28). 30 years left? AI ‘Godfather’ warns the technology may end humanity. <https://newseu.cgtn.com/news/2024-12-28/AI-Godfather-warns-rapid-development-can-cause-human-extinction-1zHEIq63EDC/index.html>
 - Dietterich, T. G. (2000). Ensemble methods in machine learning. In Multiple classifier systems (pp. 1-15). Springer. <https://web.engr.oregonstate.edu/~tgd/publications/mcs-ensembles.pdf>

- European Commission. (2019). Ethics guidelines for trustworthy AI. High-Level Expert Group on Artificial Intelligence.
https://www.europarl.europa.eu/cmsdata/196377/AI%20HLEG_Ethics%20Guidelines%20for%20Trustworthy%20AI.pdf
- European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council (Artificial Intelligence Act), Article 14: Human oversight. Official Journal of the European Union. <https://eur-lex.europa.eu/eli/reg/2024/1689/oj/eng>
- Harper, D. (n.d.). Govern. In Online Etymology Dictionary.
<https://www.etymonline.com/word/govern>
- Henrich, J., Heine, S. J., & Norenzayan, A. (2010). The weirdest people in the world? *Behavioral and Brain Sciences*, 33(2-3), 61-83.
<https://www2.psych.ubc.ca/~henrich/pdfs/WeirdPeople.pdf>
- IBM. (2025). AI governance. IBM Think.
<https://www.ibm.com/think/topics/ai-governance>
- International Organization for Standardization. (2021). ISO 37000:2021 Governance of organizations. <https://www.iso.org/standard/65036.html>
- International Organization for Standardization. (2023). ISO/IEC 42001:2023 AI management systems.
<https://www.iso.org/standard/42001>
- Irving, G., Christiano, P., & Amodei, D. (2018). AI safety via debate. *arXiv*.
<https://arxiv.org/abs/1805.00899>
- Tao, Y., Viberg, O., Baker, R. S., & Kizilcec, R. F. (2024). Cultural bias and cultural alignment of large language models. *PNAS Nexus*, 3(9), pga346.
<https://pmc.ncbi.nlm.nih.gov/articles/PMC11407280/>
- Merriam-Webster. (n.d.). Govern. In Merriam-Webster.com dictionary.
<https://www.merriam-webster.com/dictionary/govern>
- Merriam-Webster. (n.d.). Governance. In Merriam-Webster.com dictionary. <https://www.merriam-webster.com/dictionary/governance>
- Microsoft. (2022). Microsoft Responsible AI Standard v2: General requirements. <https://cdn-dynmedia-1.microsoft.com/is/content/microsoftcorp/microsoft/final/en-us/microsoft-brand/documents/Microsoft-Responsible-AI-Standard-General-Requirements.pdf>
- MIT Sloan School of Management. (2023, May 23). Why neural net pioneer Geoffrey Hinton is sounding the alarm on AI.
<https://mitsloan.mit.edu/ideas-made-to-matter/why-neural-net-pioneer-geoffrey-hinton-sounding-alarm-ai>

- National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). U.S. Department of Commerce. <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>
- National Institute of Standards and Technology. (2024). Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST AI 600-1). U.S. Department of Commerce. <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf>
- OECD. (2019). Recommendation of the Council on Artificial Intelligence. <https://legalinstruments.oecd.org/en/instruments/OECD-LEGAL-0449>
- OpenAI. (2023, July 5). Introducing Superalignment. <https://openai.com/index/introducing-superalignment/>
- Oxford English Dictionary. (n.d.). Govern, v. In OED Online. Oxford University Press. https://www.oed.com/dictionary/govern_v
- PBS. (2023, May 9). Geoffrey Hinton warns of the “existential threat” of AI. Amanpour and Company. <https://www.pbs.org/video/godfather-of-ai-warns-of-the-existential-threat-of-ai-lj1i1c/>
- UNESCO. (2021). Recommendation on the Ethics of Artificial Intelligence. <https://unesdoc.unesco.org/ark:/48223/pf0000380455>
- **Puglisi, B. C. (2026b, June). *The liability map: The three channels through which AI creates legal exposure*. basilpuglisi.com.**
 - Brownstein Hyatt Farber Schreck. (2026, March 4). Colorado’s landmark AI law coming online: What developers and deployers should know.
 - European Commission, AI Act Service Desk. Article 99: Penalties. <https://ai-act-service-desk.ec.europa.eu/en/ai-act/article-99> (penalty tiers, including the 1 percent misleading-information tier).
 - Gibson Dunn. (2026, May 27). EU AI Act Omnibus agreement: Postponed high-risk deadlines and other key changes (Annex III deferral to December 2, 2027; Annex I to August 2, 2028; August 2, 2026 transparency obligations remain live pending Official Journal publication).
 - Goodwin. (2026, June). Colorado enacts law repealing and replacing landmark AI Act (first comprehensive state AI law; repeal and replace signed May 14, 2026; effective January 1, 2027).
 - Hogan Lovells. (2026, May 7). EU legislators agree to delay for high-risk AI rules (provisional agreement reached May 7, 2026).
 - Hunton Andrews Kurth. (2026, May). Colorado AI Act amended and effective date delayed.
 - Insurance Services Office endorsements CG 40 47 and CG 40 48 (effective January 1, 2026), as reported in Lathrop GPM, The AI coverage gap (2026, May 4), and Traverse Legal, AI insurance requirements (2026,

- April 24). Carrier-level approval figures and specific endorsement language reflect trade-press reporting, not primary regulatory filings.
- *Jones v Family Court at Whangārei* [2026] NZSC 1 (self-represented litigant; AI-misuse caution from the bench).
 - *Mata v. Avianca, Inc.*, No. 1:22-cv-01461 (S.D.N.Y. 2023).
 - *Moffatt v. Air Canada*, 2024 BCCRT 149.
 - Morrison Foerster. (2026, May 15). Colorado hits reset on AI regulation with a new AI Act.
 - PYMNTS. (2026, May 1). Big insurance backs away from AI risk and startups rush in (carrier exclusion reporting).
 - Puglisi, B. C. (2026, May 24). *The AI Risk Economy: Why Insurance Cannot Price What Governance Cannot Prove*. basilpuglisi.com. <https://basilpuglisi.com/ai-risk-economy-insurance-governance/>
 - Regulation (EU) 2024/1689 (EU Artificial Intelligence Act), Article 99 penalty tiers and staged application dates.
 - Revised Product Liability Directive, Directive (EU) 2024/2853 (software and AI within strict liability; transposition due December 9, 2026).
 - Willis Research Network. (2026, May). *AI in Action: The Road to Responsible Adoption. Risk and Resilience Review*. WTW.
 - NIST AI Risk Management Framework (AI RMF 1.0), voluntary framework, and ISO/IEC 42001:2023 AI management system standard, as discussed in relation to the tort duty-of-care standard in Catastrophic Liability: Managing Systemic Risks in Frontier AI Development (arXiv:2505.00616) and in industry analysis of AI governance and insurance underwriting (Johnson Lambert LLP, 2026; StackAware, 2025, citing The Geneva Association, 2024).
 - **Puglisi, B. C. (2026c, June). *New York skipped the AI disclosure fight. It went straight to human accountability*. basilpuglisi.com.**
 - New York State Unified Court System. (2026). *Administrative Order of the Chief Administrative Judge AO/75/26* (Mar. 25, 2026), adding a new Part 161 (22 NYCRR 161.1 to 161.4 and Appendix A) to the Rules of the Chief Administrator of the Courts, effective June 1, 2026. <https://www.nycourts.gov/rules/part-161-use-artificial-intelligence-technology>
 - New York Codes, Rules and Regulations. *22 NYCRR 130-1.1 and 130-1.1a* (frivolous conduct and sanctions; signing of papers and certification after reasonable inquiry).
 - New York Rules of Professional Conduct, Rule 3.3 (22 NYCRR Part 1200) (candor toward the tribunal).
 - *Mata v. Avianca, Inc.*, 678 F. Supp. 3d 443 (S.D.N.Y. 2023).

- *United States v. Heppner*, No. 25-cr-00503-JSR, 2026 WL 436479 (S.D.N.Y. Feb. 17, 2026).
- New York State Bar Association. (2026). *Effective June 1, 2026, the New York State Unified Court System has adopted a new rule regarding the use of artificial intelligence*. <https://nysba.org/effective-june-1-2026-the-new-york-state-unified-court-system-has-adopted-a-new-rule-regarding-the-use-of-artificial-intelligence/>
- New York State Unified Court System, Office of Court Administration. (2025, October 10). *Interim Policy on the Use of Artificial Intelligence* (press release PR25_23). https://www.nycourts.gov/LegacyPDFS/press/pdfs/PR25_23.pdf
- Legal AI Governance Tracker. (2026). *New York* (individual-judge AI rules and Part 161 adoption status). <https://legalaigovernance.com/tracker/states/new-york/>
- National Institute of Standards and Technology. (2023). *AI Risk Management Framework (AI RMF 1.0, NIST AI 100-1)*. <https://www.nist.gov/itl/ai-risk-management-framework>
- International Organization for Standardization. (2023). *ISO/IEC 42001:2023, Artificial intelligence management system*.
- European Union. (2016). *Regulation (EU) 2016/679 (General Data Protection Regulation)*, Article 22 (automated individual decision-making).
- European Union. (2024). *Regulation (EU) 2024/1689 (Artificial Intelligence Act)*, Article 50 (transparency obligations for certain AI systems). <https://artificialintelligenceact.eu/article/50/>
- Puglisi, B. C. (2026). *The Standard of Care: How NIST and ISO Are Turning Voluntary AI Governance Into a Liability Defense*. <https://basilpuglisi.com/standard-of-care-ai-governance/>
- Puglisi, B. C. (n.d.). *Checkpoint-Based Governance* [Canonical framework page]. <https://basilpuglisi.com/cbg/>
- Puglisi, B. C. (2026). *The AI Risk Economy*. <https://basilpuglisi.com/ai-risk-economy-insurance-governance/>
- **Puglisi, B. C. (2026d, September). *When AI becomes assistive technology: Artificially Intelligent Disability Assistance (AIDA). How general-purpose AI can help people work, communicate, think, and create around physical and cognitive limitations* (Working paper, companion to this paper). [basilpuglisi.com](https://doi.org/10.5281/zenodo.22908272). <https://doi.org/10.5281/zenodo.22908272>**
 - Adnin, R., & Das, M. (2024). “I look at it as the king of knowledge”: How blind people use and understand generative AI tools. In *Proceedings of the*

- 26th International ACM SIGACCESS Conference on Computers and Accessibility (ASSETS '24). <https://doi.org/10.1145/3663548.3675631>
- Accessibility Standards Canada. (2025). *CAN-ASC-6.2:2025, Accessible and equitable artificial intelligence systems*. <https://accessible.canada.ca/creating-accessibility-standards/asc-62-accessible-equitable-artificial-intelligence-systems>
 - American Foundation for the Blind. (2025, January). *Guiding principles for more disability-inclusive AI*. <https://afb.org/research-and-initiatives/empowering-or-excluding/guiding-principles-more-disability-inclusive-ai>
 - Autistic Self Advocacy Network. (2025, July). *ASAN says no generative AI in plain language*. <https://autisticadvocacy.org/2025/07/asan-says-no-generative-ai-in-plain-language>
 - Bickenbach, J. (2014). The capability approach and the International Classification of Functioning. *ALTER: European Journal of Disability Research*, 8, 10 to 23.
 - Brand, D., Fischer Mogensen, K., & Engelbrecht, M. (2026, September 15). AI as reasonable accommodation and disability employment. *Disability & Society*. <https://doi.org/10.1080/09687599.2026.2726868>
 - Burchardt, T. (2004). Capabilities and disability: The capabilities framework and the social model of disability. *Disability & Society*, 19(7), 735 to 751.
 - Coffey, K. V., Krahn, G. L., Hanley, J. P., & Neely, J. E. (2026). Identifying implicit bias in LLM-based chat AI toward people with intellectual disabilities. *Disability and Health Journal*.
 - CORDIS. (n.d.). *AIDE: Adaptive multimodal interfaces to assist disabled people in daily activities* (Horizon 2020 project 645322). European Commission. <https://cordis.europa.eu/project/id/645322>
 - Della Penna, G., Buzzi, M., & Leporini, B. (2026). Generative AI as a new assistive technology for web interaction. In *Human-Computer Interaction, INTERACT 2025* (Lecture Notes in Computer Science, Vol. 16108). Springer. https://doi.org/10.1007/978-3-032-04999-5_6
 - Davis Wright Tremaine LLP. (2026, September). *New technical standard for European Accessibility Act compliance released*. <https://www.dwt.com/insights/2026/09/european-accessibility-act-ict-standards-update>
 - Demers, L., Weiss-Lambrou, R., & Ska, B. (2002). The Quebec User Evaluation of Satisfaction with Assistive Technology (QUEST 2.0). *Technology and Disability*, 14, 101 to 105.

- Deque Systems. (2026, September). *EN 301 549 v4.1.1 is final: What changed, what it means, and what you should do*.
<https://www.deque.com/blog/en-301-549-v4-1-1-is-final-what-changed-what-it-means-and-what-you-should-do/>
- DeRosier, R., & Farber, R. (2005). Speech recognition software as an assistive device: A pilot study of user satisfaction and psychosocial impact. *Work*, 25(2), 125 to 134.
- Eberhardt, K. E. (2026, July 7). *Position paper: AI as a disability accommodation in higher education. Legal analysis and annotated bibliography on cognitive scaffolding and generative AI*. Zenodo.
<https://doi.org/10.5281/zenodo.21231449>
- Ede, R. (2026, August 20). A thinking partner: Copilot as assistive technology. *Defra digital, data, technology and security blog*.
<https://defradigital.blog.gov.uk/2026/08/20/a-thinking-partner-copilot-as-assistive-technology/>
- 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).
- ETSI. (2026). *EN 301 549 V4.1.1 (2026-09), Accessibility requirements for ICT products and services*.
https://www.etsi.org/deliver/etsi_en/301500_301599/301549/04.01.01_60/en_301549v040101p.pdf
- Gilly, T. (2026a, August). *Accommodation repricing: Disability, generative AI subscriptions, and the accessibility tax as a variable the vendor controls* (v0.7). Real Safety AI Foundation. SSRN 7218538.
<https://doi.org/10.2139/ssrn.7218538>
- Gilly, T. (2026b). *The accommodation axis: Disability discrimination, algorithmic opacity, and the deployers that anti-discrimination scholarship overlooks*. SSRN 7011798.
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=7011798
- Giraldo, M., & Sacchi, F. (2026). From technical performance to assistive validity: A critical review of evaluation approaches for AI-enabled assistive technology. *Healthcare*, 14(18), 2900.
<https://doi.org/10.3390/healthcare14182900>
- Glazko, K., Yamagami, M., Desai, A., Mack, K. A., et al. (2023). *Three-month autoethnography of generative AI use for accessibility* (arXiv:2308.09924). <https://arxiv.org/abs/2308.09924>
- Gonzalez Penuela, R. E., Jung, C., Lin, S., Hu, R., & Azenkot, S. (2026). How multimodal large language models support access to visual information: A diary study with blind and low vision people. In

Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems. <https://doi.org/10.1145/3772318.3793266>

- Hadar Souval, D., Haber, Y., Tal, A., Simon, T., Elyoseph, T., & Elyoseph, Z. (2025). Transforming perceptions: Exploring the multifaceted potential of generative AI for people with cognitive disabilities. *JMIR Neurotechnology*, 4, e64182. <https://doi.org/10.2196/64182>
- Hanuschock, W. E., & Silverman, A. M. (2026, August). *Innovation for access: AI and its use as assistive technology for people with disabilities*. American Foundation for the Blind. <https://afb.org/research-and-initiatives/ai-series/innovation-access>
- Jang, J., Moharana, S., & Carrington, P. (2024). “It’s the only thing I can trust”: Envisioning large language model use by autistic workers for communication assistance. In *Proceedings of the CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3613904.3642894>
- Johnson, J., Lewis, A., Mankoff, J., & Banner, O. (2026). “I don’t trust it, but I use it”: Navigating trust, privacy, and identity in disabled people’s use of generative AI. In *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3772318.3790652>
- Jutai, J., & Day, H. (2002). Psychosocial Impact of Assistive Devices Scale (PIADS). *Technology and Disability*, 14, 107 to 111.
- Knobel, N. (2026, April 10). Algorithmic risk and neuroinclusion: AI, psychosocial safety, and the rights of neurodivergent workers in Aotearoa New Zealand. *New Zealand Journal of Health and Safety Practice*, 3(1). <https://doi.org/10.26686/nzjhsp.v3i1.10606>
- Mitra, S. (2006). The capability approach and disability. *Journal of Disability Policy Studies*, 16(4), 236 to 247.
- Mladenov, T. (2025). AI and disabled people’s independent living: A framework for analysis. *AI & Society*. <https://doi.org/10.1007/s00146-025-02642-x>
- Moore, R. E., & Williams, A. B. (2020). AIDA: Using social scaffolding to assist workers with intellectual and developmental disabilities. In *Companion of the 2020 ACM/IEEE International Conference on Human-Robot Interaction* (pp. 366 to 368). <https://doi.org/10.1145/3371382.3378385>
- Moore, R. E., & Williams, A. B. (2021). Towards a learning architecture to support social scaffolding for an artificially intelligent disability assistant. In *Companion of the 2021 ACM/IEEE International Conference on*

Human-Robot Interaction (pp. 467 to 469).

<https://doi.org/10.1145/3434074.3447215>

- Morgan, D. (2026). Reasonable adjustments for neurodivergent students in higher education: Generative AI, accessibility, and mediated authorship. *Disability & Society*. <https://doi.org/10.1080/09687599.2026.2667528>
- National Disability Insurance Agency. (n.d.). *Framework for artificial intelligence-enabled assistive technology as supports under the NDIS*. <https://ndis.gov.au/news/8492-framework-artificial-intelligence-enabled-assistive-technology-supports-under-ndis>
- Nussbaum, M. C. (2006). *Frontiers of justice: Disability, nationality, species membership*. Harvard University Press. <https://doi.org/10.4159/9780674041578>
- The Open University. (2025, July 15). *Generative AI for students*. <https://about.open.ac.uk/policies-and-reports/policies-and-statements/generative-ai-learning-teaching-and-assessment-ou-o>
- Randall, K. N., Drew, H., Gilman, E. S., & Dixon, E. (2025). Assistive technology uses and barriers in the home and workplace for adults with intellectual and developmental disabilities. *Journal of Applied Research in Intellectual Disabilities*, 38(1), e13306. <https://doi.org/10.1111/jar.13306> (first published online October 27, 2024)
- Phillips, B., & Zhao, H. (1993). Predictors of assistive technology abandonment. *Assistive Technology*, 5(1), 36 to 45. <https://doi.org/10.1080/10400435.1993.10132205>
- Polgar, J. M., Encarnação, P., Smith, E., & Cook, A. M. (2025). *Assistive technologies: Principles and practice* (6th ed.). Elsevier.
- Rodgers, R. (n.d.). AI can create plain-language and easy-read versions of every written material on earth. *Impact*, 38(3). Institute on Community Integration, University of Minnesota. <https://publications.ici.umn.edu/impact/38-3/ai-can-create-plain-language-and-easy-read-versions-of-every-written-material-on-earth-by-rylin-rodgers>
- Sustainable Bus. (2025, October). *Masats AIDA: Artificial intelligence bus doors*. <https://www.sustainable-bus.com/components/masats-aida-artificial-intelligence-bus-doors>
- Scherer, M. J., & Craddock, G. (2002). Matching Person and Technology (MPT) assessment process. *Technology and Disability*, 14(3), 125 to 131.
- Shew, A. (2023). *Against technoableism: Rethinking who needs improvement*. W. W. Norton.
- Shock, A., & Silverman, A. M. (2026, June). *Working with the machine: AI's expanding role in employment for people with disabilities*. American

- Foundation for the Blind. <https://afb.org/research-and-initiatives/ai-series/working-machine>
- Silverman, A. M., Whistler, A. L., Shock, A., Heydarian, C. H., Baguhn, S. J., Hanuschock, W. E., Hashimoto, M., Khan, O., & Vader, M.-L. (2026). *The AI quagmire: Benefits, risks, and user aspirations through a disability lens*. American Foundation for the Blind. <https://afb.org/AIResearch2>
 - United Nations. (2006). *Convention on the Rights of Persons with Disabilities*. <https://www.un.org/disabilities/documents/convention/convoptprot-e.pdf>
 - Silvera-Tawil, D., Higgins, L., Packer, K., Bayor, A. A., Walker, J. G., Li, J., Niven, P., Khanna, S., Byrnes, J., Bradford, D., et al. (2025). AI-enabled AT framework: A principles-based approach to emerging assistive technology. *Disability and Rehabilitation: Assistive Technology*, 20, 1955 to 1974. <https://doi.org/10.1080/17483107.2025.2479838>
 - Valencia, S., Cave, R., Kallarackal, K., Seaver, K., Terry, M., & Kane, S. K. (2023). “The less I type, the better”: How AI language models can enhance or impede communication for AAC users. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3544548.3581560>
 - van Toorn, G., Scully, J. L., & Gendera, S. (2025). “This robot is dictating her next steps in life”: Disability justice and relational AI ethics. *AI & Society*, 40. <https://doi.org/10.1007/s00146-025-02224-x>
 - Whittaker, M., Alper, M., Bennett, C. L., Hendren, S., Kaziunas, L., Mills, M., Morris, M. R., Rankin, J., Rogers, E., Salas, M., & West, S. M. (2019, November). *Disability, bias, and AI*. AI Now Institute. <https://ainowinstitute.org/wp-content/uploads/2023/04/disabilitybiasai-2019.pdf>
 - Williams, K., & Breazeal, C. (2013). Reducing driver task load and promoting sociability through an Affective Intelligent Driving Agent (AIDA). In *Human-Computer Interaction, INTERACT 2013* (pp. 619 to 626). Springer.
 - Williams, A. B., Williams, R. M., Moore, R. E., & McFarlane, M. (2019). AIDA: A social co-robot to uplift workers with intellectual and developmental disabilities. In *2019 14th ACM/IEEE International Conference on Human-Robot Interaction* (pp. 584 to 585). <https://doi.org/10.1109/HRI.2019.8673272>
 - World Health Organization. (2001). *International classification of functioning, disability and health*. WHO.

- World Health Organization & UNICEF. (2022). *Global report on assistive technology*. WHO.
- World Wide Web Consortium. (2024). *Web Content Accessibility Guidelines (WCAG) 2.2*. <https://www.w3.org/TR/WCAG22/>
- Wright, C. (2026). Transcription is not generation: Distinguishing non-generative AI tool use from academic misconduct in higher education assessment. *International Journal for Educational Integrity*, 22, 24. <https://doi.org/10.1007/s40979-026-00234-w>
- Zhao, X., Cox, A., & Chen, X. (2025). The use of generative AI by students with disabilities in higher education. *The Internet and Higher Education*, 66, 101014. <https://doi.org/10.1016/j.iheduc.2025.101014>
- Zhao, X., Chen, X., & Cox, A. (2026). Exploring the affordances of generative AI in academic writing for students with disabilities. *Policy Futures in Education*, 24(1), 59 to 81. <https://doi.org/10.1177/14782103251395436>

Disclaimer

I am not a lawyer, and this article does not provide legal advice. This is thought research and governance analysis based on public sources, cited materials, and human-AI review. It is intended to help executives, practitioners, insurers, and governance teams think more clearly about AI risk, liability exposure, and documentation practices. Readers should not rely on this article as a legal opinion, compliance determination, or substitute for qualified counsel. Any organization facing a legal, regulatory, contractual, or insurance question should consult its own attorney, broker, or professional adviser before acting.

The Other AI: Audio Briefings on Augmented Intelligence and AI Governance

[Spotify](#) | [Apple Podcasts](#) | [Amazon Music](#) | [YouTube Playlist](#)

#AIassisted using the HAIA Ecosystem | CC BY-NC-SA 4.0 Free for personal, educational, and noncommercial research use with attribution. Commercial exploitation, paid productization, and enterprise commercialization require separate permission and licensing.