

Against Prohibition: Generative AI, Disabled Students, and the Case for Critically Governed Inclusion in Higher Education

Nicki James Shepherd

Department of Creative Industries and Computing, Burnley College, Burnley, Lancashire, United Kingdom

Corresponding Author: Nicki James Shepherd E-mail: nick.shepherd@uk-research.org

ARTICLE INFO

Received: May 13th, 2026

Accepted: June 27th 2026

Published: July, 29th 2026

Volume: 4

Issue: 3

DOI: <https://doi.org/10.61424/issej.v4i3.963>

KEYWORDS

Generative artificial intelligence, disability, higher education, assistive technology, academic integrity, Capability Approach

ABSTRACT

Since the public release of large language models in late 2022, universities have faced pressure to prohibit or severely restrict student use of generative artificial intelligence. This paper argues that outright bans are the wrong response, and that they are wrong in a particular way: they redistribute harm towards disabled students while doing little to address the problems they claim to solve. Using conceptual analysis grounded in the social model of disability, the Capability Approach, and the author's Critical Disability Framework for AI in Education (CDF-AIED), the paper develops three claims. First, prohibition is a category error that treats a general-purpose cognitive technology as if it were a discrete cheating device, and its enforcement mechanisms, particularly AI-detection software, produce discriminatory false positives. Second, the documented biases of AI systems, which this author has examined elsewhere, justify governance rather than exclusion, since a banned technology cannot be audited, contested, or redesigned from within the institution. Third, generative AI offers a distinctive set of affordances for disabled students, including executive function scaffolding, format transformation, communication support, and reduced dependence on formal disclosure, that no previous assistive technology has combined in one tool. The discussion sketches a future in which AI functions as a governed support infrastructure and specifies the institutional conditions, from co-design to disaggregated evaluation, under which that future becomes defensible rather than naive. The paper concludes that the biases of AI are an argument for institutional presence and power over these systems, not for surrendering both through prohibition.

1. Introduction

When ChatGPT reached general release in November 2022, the first institutional reflex in many universities was refusal. Departments rewrote assessment regulations overnight, procurement teams blocked access on campus networks, and academic integrity offices reclassified the use of a language model as a form of contract cheating (Cotton et al., 2024; Sullivan et al., 2023). The reflex was understandable. Institutions that had spent two decades building quality assurance systems around the authored essay watched a freely available tool produce plausible essays on demand. Something had to be done, and prohibition is the fastest thing an institution knows how to do.

This paper argues that prohibition is nevertheless the wrong response, and that the case against it is strongest exactly where the debate has looked least: at disabled students. The academic integrity literature has treated generative AI mainly as a threat to assessment validity (Rudolph et al., 2023; Dwivedi et al., 2023). The critical literature, to which I have contributed, has treated it mainly as a carrier of ableist and racialised assumptions (Shepherd, 2025; Benjamin, 2019; Whittaker et al., 2019). Both literatures are right about the risks. Neither has fully confronted an awkward fact: for a significant group of disabled students, generative AI is already functioning as assistive technology, and a ban on the tool is, in practice, a ban on their adjustments.

I want to be precise about the argument, because it is easy to caricature. I am not claiming that generative AI is benign. In earlier work I set out a framework, the Critical Disability Framework for AI in Education (CDF-AIED), for examining how educational AI systems encode a normed, nondisabled learner in their training data, distribute their benefits unevenly, entangle personalisation with surveillance, and exclude disabled people from design power (Shepherd, 2025). Nothing in this paper retracts that analysis. The claim here is narrower and, I think, harder to dismiss: the biases of AI systems are an argument for governing them, not for banning them, because a banned system is a system the institution has surrendered any power to shape. Prohibition does not remove the technology from students' lives. It removes the institution from the technology's governance.

The research question guiding the paper is therefore: on what grounds, and under what institutional conditions, should universities permit rather than prohibit student use of generative AI, given both its documented biases and its emerging role as support technology for disabled students? Section 2 reviews the three literatures the question sits between. Section 3 describes the conceptual methodology. Section 4 presents the analysis in three parts: why bans fail on their own terms, why bias supports governance rather than exclusion, and what generative AI offers disabled students. Section 5 discusses what a governed future could look like and the conditions that would make it defensible. Section 6 concludes with limitations and directions for research.

A note on positionality, consistent with my earlier work: I write as a disabled practitioner-researcher. I use some of these tools daily, not as a novelty but as a working adjustment, and that experience shapes the analysis without, I hope, exempting it from scrutiny (Hamraie & Fritsch, 2019).

2. Literature Review

2.1 Generative AI and the academic integrity debate

The first wave of scholarship on generative AI in higher education was dominated by the question of cheating. Cotton et al. (2024) framed the arrival of ChatGPT as a direct challenge to assessment integrity and surveyed institutional countermeasures from detection software to assessment redesign. Rudolph et al. (2023) asked whether the technology spelled the end of traditional assessment altogether, and Dwivedi et al. (2023), in a large multidisciplinary commentary, catalogued both opportunity and threat across research, practice, and policy. Survey evidence gathered during the first year of availability showed a consistent pattern: majority student use alongside majority institutional restriction (Sullivan et al., 2023; Chan, 2023). Policy responses have since softened from prohibition towards principled permission, a direction visible in the Russell Group (2023) principles and in UNESCO's (2023) global guidance, both of which favour transparency and assessment redesign over bans. What this literature has barely addressed is disability. Its implicit student is a nondisabled student choosing between honest and dishonest effort, and its equity concerns have centred on access to paid tools rather than on who bears the burden of enforcement.

2.2 Critical disability studies and algorithmic harm

A second literature examines what algorithmic systems do to disabled people. Its foundations lie in the social model of disability, which locates disability in disabling structures rather than in impaired individuals (Oliver, 1990; Shakespeare, 2014; Thomas, 2007), and in the Capability Approach, which evaluates social arrangements by the real freedoms they afford (Sen, 1999; Nussbaum, 2011; Terzi, 2005). Applied to computation, this tradition has documented discriminatory design across welfare automation (Eubanks, 2018), search (Noble, 2018), and the broad pattern Benjamin (2019) calls the New Jim Code, with Whittaker et al. (2019) providing the most direct treatment of disability bias in AI systems. Crip technoscience insists that disabled people are experts and designers of the

sociotechnical systems they navigate, not passive recipients of them (Hamraie & Fritsch, 2019; Kafer, 2013), and Fricker (2007) supplies the concept of epistemic injustice for cases where that expertise is discounted. In education specifically, Williamson (2017) and Selwyn (2022) have traced how data-driven systems encode a politics of normalisation, and my own earlier work synthesised these strands into the CDF-AIED, a four-dimensional framework for evaluating educational AI from a disability justice perspective (Shepherd, 2025). This literature is strong on harm and weak on the question this paper asks: what disabled students lose when institutions respond to harm with exclusion.

2.3 Assistive technology and disabled students in higher education

A third literature examines disabled students' actual conditions in higher education. It documents an environment of persistent structural barriers: fixed-format teaching materials (Seale, 2014), disclosure processes that are stigmatising, slow, and deterrent (Goode, 2007; Madriaga, 2007; Eccles et al., 2018), and a psycho-emotional labour of self-management that Thomas (2007) terms barriers to being. Against this background, the assistive technology evidence is encouraging but qualified: McNicholl et al. (2021), in a systematic review, found technology use associated with gains in participation, autonomy, and wellbeing, alongside recurring problems of cost, training, and inadequate institutional support. Emerging work on large language models in education notes their potential for individualised support while cataloguing risks of error, bias, and dependence (Kasneci et al., 2023; Bearman & Ajjawi, 2023). What is missing across these literatures is integration: the integrity debate ignores disability, the critical literature stops at critique, and the assistive technology literature has not yet caught up with a tool that arrived faster than any evaluation cycle. This paper addresses that gap.

3. Methodology

This is a conceptual and normative paper, a recognised mode of theoretical study in which existing bodies of thought are synthesised to structure a problem that empirical work has not yet framed (Jabareen, 2009). The approach follows the method of my earlier framework paper (Shepherd, 2025): arguments and evidence are drawn from distinct literatures, brought into contact around a defined question, and assessed against an explicit evaluative standard. Three methodological commitments shape the analysis. First, the evaluative standard is the Capability Approach as applied to disability and education (Sen, 1999; Terzi, 2005): policies are judged by whether they expand or contract the real educational freedoms of the students with the least, relative to actually existing alternatives rather than imagined neutral baselines. Second, the CDF-AIED framework (Shepherd, 2025) is used as an analytical lens rather than re-derived; its four dimensions (the construction of the normal learner, the distribution of adaptive benefits, the politics of disclosure and surveillance, and design power) organise the assessment of both prohibition and permission. Third, because the paper is conceptual, its results are analytical claims rather than empirical findings; they are presented in Section 4 with the evidence that supports them and are offered as hypotheses for the empirical programme described in Section 6.

4. Results

4.1 Why prohibition fails on its own terms

Institutional bans on generative AI have rested on three rationales. The first is assessment integrity: if a model can produce a passing essay, unsupervised written assessment no longer certifies individual competence (Cotton et al., 2024). The second is equity: students with paid access to more capable models gain an advantage over those without, so restricting everyone appears fairer than permitting a stratified free-for-all. The third is data protection: student inputs to commercial models may be retained, reused for training, and processed outside institutional and jurisdictional control (UNESCO, 2023). Each rationale identifies a real problem. The difficulty is that in each case prohibition is a poor instrument for the problem named, and in each case the enforcement burden falls unevenly.

The enforcement problem comes first. A ban on generative AI is unenforceable in any straightforward sense, because the technology leaves no reliable forensic trace. The detection tools marketed to universities do not detect AI use; they estimate the statistical predictability of prose, and this estimate is systematically confounded with markers of linguistic difference. Liang et al. (2023) found that widely used GPT detectors misclassified more than half of a sample of essays by non-native English writers as machine-generated, while performing near-perfectly on essays by native speakers. There is reason to expect these tools to behave no better with the atypical syntax, formulaic scaffolding, or assistive-software artefacts present in the writing of many disabled students. A dyslexic student who has always relied on predictable sentence frames, or an autistic student whose prose style is unusually regular, writes exactly the kind of text a perplexity-based detector flags. The result is a familiar pattern from critical algorithm

studies: an enforcement technology adopted to protect fairness that redistributes suspicion towards those already furthest from the institutional norm (Eubanks, 2018; Benjamin, 2019). Under a ban regime, the students most likely to face a misconduct allegation are not the students most likely to have cheated. They are the students whose unassisted writing least resembles the statistical centre of the training distribution. Prohibition does not merely fail to protect disabled students; it manufactures a new mechanism for accusing them.

The displacement problem follows. Bans do not stop use; they stratify it. Students with private devices, personal subscriptions, and confident digital literacies continue to use the tools quietly and skilfully. Students who depend on institutional infrastructure, or who are more rule-anxious, comply. A policy that is widely violated but selectively enforced is not a neutral policy. It converts a technology gap into a compliance gap, and compliance gaps are always patterned by disability, class, and prior familiarity with institutional risk (Madriaga, 2007). There is a further displacement that matters specifically for disabled students. Where AI use is prohibited generally but permitted as an individual adjustment, the disabled student who uses it must disclose, evidence, and register that use through the same bureaucratic apparatus that already deters disclosure (Goode, 2007; Eccles et al., 2018). The ban thereby recreates, for a new technology, the oldest problem in disability support: help that is available only at the price of administrative self-identification, with all the stigma and delay that entails.

Underlying these failures is a category error. Prohibition treats generative AI as a discrete device with a single function, like a phone in an exam hall, when it is better understood as a general-purpose cognitive technology that is already being embedded in word processors, search engines, screen readers, and operating systems. Within a few years there will be no clean behavioural line between “using AI” and “using a computer”. A rule that cannot specify what it prohibits cannot be fairly enforced, and a rule that in practice prohibits the default tools of digital writing is a rule against participation itself. Bearman and Ajjawi (2023) make a related pedagogical point: students will spend their working lives alongside opaque computational systems, and an education that pretends otherwise fails at its own vocational promises. The regulatory question worth asking is not whether students may touch the technology but which uses, in which contexts, with what transparency, serve the purposes of assessment and learning (Russell Group, 2023; UNESCO, 2023).

4.2 Bias justifies governance, not exclusion

The strongest argument for restriction does not come from integrity officers. It comes from the critical tradition in which my own work sits. Large language models are trained on corpora that encode normative assumptions about what competent communication looks like, and their fine-tuning processes reward outputs that predominantly neurotypical evaluators judged helpful (Shepherd, 2025). Adaptive systems built on such models inherit a model of the normal learner and can read disabled students’ patterns as deficits (Oliver, 1990; Hamraie & Fritsch, 2019). Personalisation architectures gather behavioural data that can function as undisclosed disability profiling (Whittaker et al., 2019). If all of that is true, and I have argued at length that it is, why is exclusion not the safer course?

The answer begins with an observation about where critique can operate. An institution that licenses, procures, or formally permits a technology acquires levers over it: contractual data-protection terms, accessibility requirements, audit rights, the ability to demand disaggregated evaluation, and a legitimate forum in which students can report harm. An institution that bans a technology has none of these levers, while its students continue to use the technology on consumer terms, under consumer surveillance, with no institutional record that use is occurring at all. The ban does not protect disabled students from biased AI. It ensures that their encounters with biased AI happen in private, unsupported, and invisible to the equality duties that would otherwise apply.

This is why the CDF-AIED framework, though built as an instrument of critique, points away from prohibition. Its four dimensions are all questions that can only be asked of systems an institution actually engages with (Shepherd, 2025). One cannot audit the training assumptions of a tool one refuses to procure. One cannot demand co-design of a tool one pretends students do not use. Critical disability studies has long insisted that disabled people be present within sociotechnical systems as experts and designers rather than positioned outside them as objects of protection (Hamraie & Fritsch, 2019; Kafer, 2013). A ban is protective positioning of precisely the kind that tradition rejects.

The Capability Approach supplies the evaluative standard that prohibition lacks. On this test, the question is never “does this technology contain bias?”, since every educational technology, including the lecture, the essay, and the timed examination, contains bias. The question is comparative and distributional: relative to the actually existing alternative, does regulated access to this technology expand the capabilities of the students with the least, and can its harms be constrained by governance (Sen, 1999; Terzi, 2005)? What the capability test rules out is the move prohibition depends on: judging the technology against an imagined neutral baseline, when the actual baseline, higher education as currently arranged, systematically underserves disabled students already (Madriaga, 2007; Seale, 2014).

4.3 What generative AI offers disabled students

The affordances described here extend a well-documented finding in the assistive technology literature: that technology use is associated with meaningful gains in participation, autonomy, and psychological wellbeing for disabled students, alongside persistent problems of cost, training, and institutional support (McNicholl et al., 2021; Seale, 2014). What is new about generative AI is not any single function but the combination of functions in one conversational tool that requires no specialist procurement, no diagnostic gatekeeping, and no per-function training course. Four affordances stand out.

The first is executive function scaffolding. A large share of the disabled student population in UK higher education is neurodivergent, and executive function demands, planning an essay, decomposing a project into steps, initiating a task, re-entering work after interruption, are among the most consistently reported barriers (Goodley, 2014). Generative AI performs exactly this scaffolding on demand: it can turn an assignment brief into a sequenced plan, generate the first rough sentence that breaks task-initiation paralysis, and summarise where a piece of work stopped so that a student returning after a fatigue episode does not pay the full cognitive cost of re-entry. These are not luxuries or shortcuts. They are functions for which universities currently employ specialist study skills tutors, funded through disability support schemes, with waiting lists measured in weeks. The model does a worse job than a good specialist tutor. It does an incomparably better job than the waiting list.

The second is format transformation. For students with dyslexia, visual impairment, or processing differences, the persistent barrier is not the content of higher education but its fixation in single formats: the dense thirty-page PDF, the uncaptioned recording, the diagram with no description (Seale, 2014). Universal Design for Learning has argued for decades that material should be available in multiple means of representation, and institutions have implemented this slowly, unevenly, and mostly retrospectively (Holmes et al., 2019). A language model collapses the transformation cost to near zero at the point of use: the same source can be summarised, simplified, expanded, restructured, read aloud, or interrogated through questions, at whatever grain the reader needs. This shifts control of format from the producer to the reader, which is where the social model of disability has always said it belongs (Oliver, 1990). The barrier was never the student’s reading; it was the assumption that one fixed document suits every reader.

The third is communication load. Autistic students and others with communication-related disabilities pay a continuous, largely invisible tax in higher education: composing the email to the tutor in the expected register, decoding what feedback actually requires, drafting the seminar contribution, judging the tone of an appeal. Thomas (2007) calls the cumulative psycho-emotional weight of such demands barriers to being, and they shape attainment and attrition as surely as any physical barrier. Generative AI functions here as a translation layer in both directions: it can render institutional language into plain meaning and render the student’s meaning into institutional language. One can object that this trains conformity to neurotypical norms, and the objection has force. But the immediate alternative is not a university that accepts autistic communication as it stands. The immediate alternative is the student paying the full tax alone, or not sending the email at all.

The fourth affordance is structural rather than functional, and it is the one the literature has most seriously underpriced: support without disclosure. Every formal adjustment in UK higher education is gated by disclosure: the student must identify as disabled, produce medical evidence, and negotiate support through processes documented as stigmatising, slow, and deterrent, with substantial numbers of eligible students never disclosing at all (Goode, 2007; Eccles et al., 2018; Madriaga, 2007). Generative AI is the first broadly capable support technology whose use requires no disclosure whatsoever. The student who cannot face the evidence process, or is still years from a diagnosis, or comes from a community in which disability identification carries particular cost, obtains a meaningful fraction of the support function anyway. In earlier work I emphasised the dark side of this same property: systems that infer difficulty from behaviour can build disability profiles without consent or legal protection (Shepherd, 2025). Both

things are true, and the difference lies in where the inference happens and who holds the record. A student privately using a general-purpose tool discloses nothing to their institution. An institutional analytics platform inferring “attentional difficulties” from that student’s clickstream is a different technology, a different data controller, and a different ethics. Policy that cannot distinguish these two cases, and blanket prohibition cannot, is not yet policy about the thing itself.

None of these affordances comes free, and a paper arguing against prohibition has a particular duty not to write advertising copy. Three costs deserve explicit entry in the ledger. First, normalisation: a tool that translates atypical communication into standard register makes individual lives easier while leaving the standard itself unchallenged, a dynamic Kafer (2013) would recognise as curative logic in software form. Second, dependence and skill formation: scaffolding that substitutes for a developing capability, rather than supporting it, can hollow out the very capacities education exists to build, and we do not yet have good longitudinal evidence about where that line falls. Third, quality and error: models produce confident falsehoods, and students with the least academic capital, including many disabled students, are least positioned to catch them (Kasneci et al., 2023). These are serious costs. They are also, without exception, arguments for governed and pedagogically embedded use rather than for prohibition, since every one of them worsens when use is unsupported, untaught, and hidden.

5. Discussion

5.1 *From tool to infrastructure*

The near future of generative AI in education will not look like students visiting a chatbot website. It will look like ambient capability: models embedded in the word processor, the reading list, the virtual learning environment, the library catalogue, and the operating system’s accessibility layer. The policy window in which universities can shape this embedding is open now and will not stay open long. My earlier warning applies with full force: assumptions poured into infrastructure set hard, and are far more expensive to contest after the fact than before (Shepherd, 2025; Williamson, 2017). The choice facing institutions is not whether AI becomes part of educational infrastructure. It is whether disabled students’ interests are engineered into that infrastructure or bolted on afterwards, as they were for the web, for e-learning platforms, and for lecture capture in turn (Seale, 2014).

5.2 *A defensible institutional settlement*

A defensible settlement would, at minimum, have five features. I present them as commitments rather than as a checklist, because checklists are how accessibility became compliance theatre in the first place (Shepherd, 2025).

The first is a rebuttable presumption of permitted use. Assessment design, not blanket regulation, marks the specific contexts in which unassisted performance is genuinely what is being certified, and those contexts are secured by design (supervised, oral, practical) rather than by detection software. This aligns with the direction of the Russell Group (2023) principles and UNESCO (2023) guidance, but states plainly what they imply: the default is access.

The second is recognition of AI use as a reasonable adjustment where it functions as one. Under the Equality Act 2010, institutions already owe anticipatory duties to disabled students. An institution that prohibits a technology a disabled student reasonably relies on for access, without providing an equivalent, is on increasingly weak legal as well as ethical ground. Adjustment status also protects disabled students from the misconduct exposure described in Section 4.1: use pursuant to an adjustment cannot coherently be an offence.

The third is procurement with teeth. Institutional licensing of AI tools should be conditional on contractual terms covering training-data transparency sufficient for audit, retention limits on student inputs, prohibition of secondary profiling, and accessibility conformance, with the CDF-AIED dimensions supplying the evaluative frame (Shepherd, 2025). This is the lever prohibition throws away.

The fourth is disaggregated evaluation as a standing condition. Every institutional deployment should report learning and experience outcomes broken down by disability status and type, because the affordances set out in Section 4.3

are hypotheses to be tested, not entitlements to be assumed, and because the central finding of the critical literature is that aggregate benefit routinely conceals distributional harm (Sen, 1999; Terzi, 2005; Noble, 2018).

The fifth is design power. Disabled students belong on the governance bodies, procurement panels, and evaluation teams that shape institutional AI, not as consultees after decisions but as decision-makers within them, for the epistemic reasons Fricker (2007) and Hamraie and Fritsch (2019) supply: they know things about navigating hostile systems that no one else in the room knows.

5.3 Bridging the gap, honestly described

What might the resulting future actually deliver? The realistic promise is not that AI closes the disability attainment gap; gaps of that kind are produced by structures, and tools do not dismantle structures (Oliver, 1990). The realistic promise is that AI takes over a substantial share of the compensatory labour that disabled students currently perform unpaid and alone: the reformatting, the planning, the translating, the re-entering, the composing of institutional selves. That labour is real, it is unevenly levied, and it is currently invisible in every attainment statistic. A support infrastructure that absorbs even part of it would widen what Terzi (2005) calls the real freedom to pursue one's own learning project, which is the outcome that matters, whether or not the headline metrics move.

There is a longer-range possibility worth naming, though it belongs to speculation rather than analysis. Interfaces that adapt to the reader dissolve the economic argument for the fixed, single-format curriculum, and once every text is fluid, the category of the "non-standard" reader begins to lose its meaning. That would be the social model's ambition realised in infrastructure: not disabled students accommodated within a normed system, but a system with no single norm to deviate from. Nothing guarantees that trajectory, and the surveillance-heavy alternative is at least as likely. Which future arrives is a governance question, which is to say a political one, which is to say one that universities forfeit the right to influence on the day they opt for prohibition.

6. Conclusion

The question universities have been asking, whether to permit generative AI, was never quite the right question. The technology's presence in students' lives is not within institutional control; the terms of its presence are. This paper has argued that outright bans fail on their own rationales, expose disabled students to discriminatory enforcement, tax them with disclosure burdens, and surrender the governance levers through which biased systems could be audited and redesigned. It has argued that the same critical framework that documents AI's ableist assumptions supports engagement over exclusion, because critique needs presence. And it has set out, honestly costed, what generative AI already does for disabled students, along with the institutional settlement under which those benefits could be secured against the harms.

Three limitations should be stated. First, this is a conceptual and normative argument, not an empirical evaluation; the affordance claims in Section 4.3 rest on adjacent assistive technology evidence (McNicholl et al., 2021) and on practitioner experience, and they require direct testing through participatory research with disabled students using these specific tools. Second, disability is not unitary, and the argument's weight varies across it: the case is strongest for neurodivergent students and students with print and communication disabilities, and weaker where AI offers less distinctive support. Third, the analysis is anchored in UK higher education and its particular legal duties; the balance of argument may differ in jurisdictions with weaker data protection or no anticipatory adjustment duty, where the surveillance costs of engagement rise.

Three research directions follow. First, empirical testing of the affordance claims through mixed-methods and participatory research with disabled students using generative AI in real study contexts. Second, evaluation of institutional policy variation: as universities settle on divergent AI regulations, comparative study of their disability-disaggregated outcomes becomes possible and urgent. Third, development of the procurement and audit instruments sketched in Section 5.2 into operational tools, work for which the CDF-AIED provides a starting frame.

I remain, as in my earlier work, unwilling to be either an evangelist or an abolitionist about these systems. They are biased, extractive, and frequently wrong, and they are also the most capable support technology that has ever been available to a disabled student at zero marginal cost on a Tuesday night before a deadline. Both facts are true at once. Policy made from only one of them will be bad policy. The task, and it is urgent, is to build the governed middle: an infrastructure in which the tools are present, their terms are contested, their benefits are measured where they land,

and the people with the most at stake hold real power over their design. Prohibition builds nothing. It only decides who faces the technology alone.

Funding: This research received no external funding.

Conflicts of Interest: The authors declare no conflict of interest.

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers

Artificial Intelligence (AI) Use Disclosure: The authors declare that no artificial intelligence tools were used in the preparation of this manuscript.

References

- Bearman, M., & Ajjawi, R. (2023). Learning to work with the black box: Pedagogy for a world with artificial intelligence. *British Journal of Educational Technology*, 54(5), 1160–1173.
- Benjamin, R. (2019). *Race after technology: Abolitionist tools for the New Jim Code*. Polity Press.
- Chan, C. K. Y. (2023). A comprehensive AI policy education framework for university teaching and learning. *International Journal of Educational Technology in Higher Education*, 20, Article 38.
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of ChatGPT. *Innovations in Education and Teaching International*, 61(2), 228–239.
- Dwivedi, Y. K., Kshetri, N., Hughes, L., Slade, E. L., Jeyaraj, A., Kar, A. K., Baabdullah, A. M., Koohang, A., Raghavan, V., Ahuja, M., Albanna, H., Albashrawi, M. A., Al-Busaidi, A. S., Balakrishnan, J., Barlette, Y., Basu, S., Bose, I., Brooks, L., Buhalis, D., ... Wright, R. (2023). “So what if ChatGPT wrote it?” Multidisciplinary perspectives on opportunities, challenges and implications of generative conversational AI for research, practice and policy. *International Journal of Information Management*, 71, 102642. <https://doi.org/10.1016/j.ijinfomgt.2023.102642>
- Eccles, S., Hutchings, M., Hunt, C., & Heaslip, V. (2018). Disclosure of ‘disability’ in higher education: A qualitative study of student perceptions. *Widening Participation and Lifelong Learning*, 20(4), 191–208. <https://doi.org/10.5456/WPLL.20.4.191>
- Equality Act 2010, c. 15 (UK). <https://www.legislation.gov.uk/ukpga/2010/15>
- Eubanks, V. (2018). *Automating inequality: How high-tech tools profile, police and punish the poor*. St. Martin's Press.
- Fricker, M. (2007). *Epistemic injustice: Power and the ethics of knowing*. Oxford University Press.
- Goode, J. (2007). ‘Managing’ disability: Early experiences of university students with disabilities. *Disability & Society*, 22(1), 35–48. <https://doi.org/10.1080/09687590601056204>
- Goodley, D. (2014). *Dis/ability studies: Theorising disablism and ableism*. Routledge.
- Hamraie, A., & Fritsch, K. (2019). Crip technoscience manifesto. *Catalyst: Feminism, Theory, Technoscience*, 5(1), 1–34. <https://doi.org/10.28968/cftt.v5i1.29607>
- Holmes, W., Bialik, M., & Fadel, C. (2019). *Artificial intelligence in education: Promises and implications for teaching and learning*. Center for Curriculum Redesign.
- Jabareen, Y. (2009). Building a conceptual framework: Philosophy, definitions, and procedure. *International Journal of Qualitative Methods*, 8(4), 49–62. <https://doi.org/10.1177/160940690900800406>
- Kafer, A. (2013). *Feminist, queer, crip*. Indiana University Press.
- Kasneci, E., Sessler, K., Küchemann, S., Bannert, M., Dementieva, D., Fischer, F., Gasser, U., Groh, G., Günemann, S., Hüllermeier, E., Krusche, S., Kutyniok, G., Michaeli, T., Nerdel, C., Pfeiffer, J., Poquet, O., Sailer, M., Schmidt, A., Seidel, T., ... Kasneci, G. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>
- Liang, W., Yuksekgonul, M., Mao, Y., Wu, E., & Zou, J. (2023). GPT detectors are biased against non-native English writers. *Patterns*, 4(7), 100779. <https://doi.org/10.1016/j.patter.2023.100779>
- Madriaga, M. (2007). Enduring disablism: Students with dyslexia and their pathways into UK higher education and beyond. *Disability & Society*, 22(4), 399–412. <https://doi.org/10.1080/09687590701337942>
- McNicholl, A., Casey, H., Desmond, D., & Gallagher, P. (2021). The impact of assistive technology use for students with disabilities in higher education: A systematic review. *Disability and Rehabilitation: Assistive Technology*, 16(2), 130–143.
- Noble, S. U. (2018). *Algorithms of oppression: How search engines reinforce racism*. New York University Press.
- Nussbaum, M. C. (2011). *Creating capabilities: The human development approach*. Harvard University Press.
- Oliver, M. (1990). *The politics of disablement*. Macmillan.

- Russell Group. (2023). *Russell Group principles on the use of generative AI tools in education*. Russell Group.
- Rudolph, J., Tan, S., & Tan, S. (2023). ChatGPT: Bullshit spewer or the end of traditional assessments in higher education? *Journal of Applied Learning and Teaching*, 6(1), 342–363.
- Seale, J. (2014). *E-learning and disability in higher education: Accessibility research and practice* (2nd ed.). Routledge.
- Selwyn, N. (2022). The future of AI and education: Some cautionary notes. *European Journal of Education*, 57, 620–631. <https://doi.org/10.1111/ejed.12532>
- Sen, A. (1999). *Development as freedom*. Oxford University Press.
- Shakespeare, T. (2014). *Disability rights and wrongs revisited* (2nd ed.). Routledge.
- Shepherd, N. J. (2025). Disability, data, and design: Toward a critical framework for evaluating AI in inclusive education. *IRE Journals*, 8(12), 2218–2230. <https://doi.org/10.64388/IREV8I12-1719101>
- Sullivan, M., Kelly, A., & McLaughlan, P. (2023). ChatGPT in higher education: Considerations for academic integrity and student learning. *Journal of Applied Learning and Teaching*, 6(1), 31–40.
- Terzi, L. (2005). Beyond the dilemma of difference: The capability approach to disability and special educational needs. *Journal of Philosophy of Education*, 39(3), 443–459. <https://doi.org/10.1111/j.1467-9752.2005.00447.x>
- Thomas, C. (2007). *Sociologies of disability and illness: Contested ideas in disability studies and medical sociology*. Palgrave Macmillan.
- UNESCO. (2023). *Guidance for generative AI in education and research*. UNESCO.
- 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). *Disability, bias, and AI*. AI Now Institute.
- Williamson, B. (2017). *Big data in education: The digital future of learning, policy and practice*. SAGE.