

AGAINST THE UNCRITICAL ADOPTION OF ‘AI’ TECHNOLOGIES IN ACADEMIA

Olivia Guest^{1,2}, Marcela Suarez¹, Barbara C. N. Müller³, Edwin van Meerkerk⁴, Arnoud Oude Groote Beverborg⁵, Ronald de Haan⁶, Andrea Reyes Elizondo⁷, Mark Blokpoel^{1,2}, Natalia Scharfenberg^{1,2}, Annelies Kleinherenbrink^{1,8}, Ileana Camerino¹, Marieke Woensdregt^{1,2}, Dagmar Monett⁹, Jed Brown¹⁰, Lucy Avraamidou¹¹, Juliette Alenda-Demoutiez¹², Felienne Hermans¹³, and Iris van Rooij^{1,2,14}

¹Department of Cognitive Science and Artificial Intelligence, Radboud University, The Netherlands

²Donders Institute for Brain, Cognition, and Behaviour, Radboud University, The Netherlands

³Communication & Media, Behavioural Science Institute, Radboud University, The Netherlands

⁴Radboud Institute for Culture and History, Radboud University, The Netherlands

⁵Pedagogical and Educational Sciences, Radboud University, The Netherlands

⁶Institute for Logic, Language and Computation, University of Amsterdam, The Netherlands

⁷CWTS & LUCAS, Faculties of Social Sciences & Humanities, Leiden University, The Netherlands

⁸Gender and Diversity, Radboud University, The Netherlands

⁹Computer Science Dept., Berlin School of Economics and Law, Germany

¹⁰Department of Computer Science, University of Colorado, USA

¹¹Institute for Science Education and Communication, University of Groningen, The Netherlands

¹²Department of Economics and Business Economics, Radboud University, The Netherlands

¹³Vrije Universiteit Amsterdam, The Netherlands

¹⁴Department of Linguistics, Cognitive Science, and Semiotics, Aarhus University, Denmark

Abstract: *Under the banner of progress, products have been uncritically adopted or even imposed on users — in past centuries with tobacco and combustion engines, and in the 21st with social media. For these collective blunders, we now regret our involvement or apathy as scientists, and society struggles to put the genie back in the bottle. Currently, we are similarly entangled with artificial intelligence (AI) technology. For example, software updates are rolled out seamlessly and non-consensually, Microsoft Office is bundled with chatbots, and we, our students, and our employers have had no say, as it is not considered a valid position to reject AI technologies in our teaching and research. This is why in June 2025, we co-authored an Open Letter calling on our employers to reverse and rethink their stance on uncritically adopting AI technologies. In this position piece, we expound on why universities must take their role seriously to a) counter the technology industry’s marketing, hype, and harm; and to b) safeguard higher education, critical thinking, expertise, academic freedom, and scientific integrity. We include pointers to relevant work to further inform our colleagues.*

Keywords: *higher education; artificial intelligence; digital technology; critical analysis; open letter; policy*

1 Overview

The culture of AI is imperialist and seeks to expand the kingdom of the machine. The AI community is well organized and well funded, and its culture fits its dreams: it has high priests, its greedy businessmen, its canny politicians. The U.S. Department of Defense is behind it all the way. And like the communists of old, AI scientists believe in their revolution; the old myths of tragic hubris don't trouble them at all.

Tony Solomonides and Les Levidow (1985, pp. 13–14)

This paper sets out our expert position on artificial intelligence (AI) technologies permeating the higher education sector, demonstrating how this directly erodes our ability to function (see also our *Open Letter*, Guest, van Rooij, et al. 2025). The harms to our fields and students that directly result from the technology sector corrupting our practices unchecked and unimpeded are manifold: from conflicts of interest that go unreported or are worn as badges of honour by colleagues (Mohamed Abdalla and Moustafa Abdalla 2021), to the mushrooming of chatbots in software we are coerced to use, such as in Microsoft Office. This is highly problematic, as academia is meant to be a refuge for knowledge production independent from ulterior motives, weaving together teaching and research. Research funding and academic freedom are presently under a worldwide attack (ALLEA 2025; Kinzelbach et al. 2025; KNAW 2021, 2025). The technology industry is taking advantage of us, sometimes even speaking through us, to convince our students that these AI technologies are useful (or necessary) and not harmful. Therefore, we argue that university leaders and administrators must act to help us collectively turn back the tide of garbage software, which fuels harmful tropes (e.g. so-called lazy students) and false frames (e.g. so-called efficiency or inevitability) to obtain market penetration and increase technological dependency.

When it comes to the AI technology industry, we refuse their frames, reject their addictive and brittle technology, and demand that the sanctity of the university both as an institution and a set of values be restored. If we cannot even in principle be free from external manipulation and anti-scientific claims — and instead remain passive by default and welcome corrosive industry frames into our computer systems, our scientific literature, and our classrooms — then we have failed as scientists and as educators.

2 Marketing, hype, & harm

In any given professional field, specialized jargon is often necessary in order to exchange information more succinctly and specifically; it makes communication clearer. But in a cultish atmosphere, jargon does just the opposite: Instead, it causes speakers to feel confused and intellectually deficient. That way, they'll comply.

Amanda Montell (2021, pp. 136–137)

AI has always been a marketing phrase that erodes scientific inquiry and scholarly discussion by design, leaving the door open to pseudoscience, exclusion, and surveillance (cf. Birhane and Guest 2021; Guest 2025; Guest and Forbes 2024; van Rooij, Guest, et al. 2024; Wendling 2002). From its inception in the 1950s, the phrase 'artificial intelligence' was used to sell research, to spice up existing research programmes and attract funding (AAUP 2025; Bender 2024; Bloomfield 1987; Heffernan 2019; Markelius et al. 2024; McCorduck 2004).

We propose that a 2 month, 10 man study of artificial intelligence be carried out during the summer of 1956 at Dartmouth College in Hanover, New Hampshire. The study is to proceed on the basis of the conjecture that every aspect of learning or any other feature

of intelligence can in principle be so precisely described that a machine can be made to simulate it. An attempt will be made to find how to make machines use language, form abstractions and concepts, solve kinds of problems now reserved for humans, and improve themselves. We think that a significant advance can be made in one or more of these problems if a carefully selected group of scientists work on it together for a summer.

John McCarthy et al. (1955, p. 2)

This proposal proved too tempting to ignore for colleagues and funders — both in the “2 month” speed of delivery and the lofty goal of machines capturing “every aspect of learning” and “feature of intelligence.” It shows that since the start imprecise jargon was used to make exaggerated promises with the goal to please investors, claims that fundamentally remain promises to this day.

Although some differentiate AI and non-AI systems by appealing to generative models versus other types of AI, we are convinced that this does not bring clarity to the discussion, and we can fall head first into misuse of terminology and fostering industry hype (see Table 1 and Figure 1; cf. Guest 2025). Granting that distinction, for example, makes classifiers such as *Bernoulli naive Bayes*, formally a generative model, AI technology similar to hyped applications (Efron 1975; Jebara 2004; Mitchell 1997; Ng and Jordan 2001; Xue and Titterton 2008). However, a statistical model like the Bernoulli naive Bayes classifier, which is used to analyse data, is unrelated to industry hype. By the same token, spellcheck software (which has not historically been composed of generative nor stochastic models) also falls under AI in some definitions, but not under others. Similar terminological problems also hold for phrases like ‘agentic AI’ (Chawla et al. 2024; Hosseini and Seilani 2025) — systems that can autonomously adapt their environment, make context-sensitive ‘decisions’, and act upon them without any human intervention. Formally, systems like thermostats, microwave ovens, and traffic lights are agentic in this sense and fall under AI, but are not products we wish to proscribe. Informally, however, ‘agentic’ is misused to induce anthropomorphisation that we cannot endorse (cf. Bandura 2001; Barrow 2024; Helfrich 2024).

There is nothing privileged about generative AI, nor any AI technology, in any pedagogical or sociotechnical sense (see Table 1). Importantly, *all* technology under our purview as academics, ought to fall under the same critical considerations. Thus, to frame generative AI as unique or dramatically different to other AI systems appears anti-intellectual. Recall that generative AI is a mere technical difference (i.e. which statistical distribution a system is capable of) and not some substantive sociotechnical departure. And so non-generative AI can be highly problematic, such as computer vision systems which are used by the military and police (Birhane and Prabhu 2021; Falletti 2024; Wood 2024). Such systems also comprise stolen data, require enslaved people to tag the data, and produce biased results (Birhane, Dehdashtian, et al. 2024; Birhane, Prabhu, et al. 2023; Kalluri et al. 2025; Tait et al. 2022).

The Euler diagram in Figure 1 demonstrates that such hyped labels — Large Language Model (LLM), Artificial Neural Network (ANN), generative model, chatbot — are interwoven in complicated ways such that their referents remain ambiguous (cf. Guest 2025). Importantly, this defiance holds under any rearrangement of the sets and elements: one cannot use these words without encountering additional problems and fostering ambiguity (Alkhatib 2024). We cannot escape this quagmire with more jargon. And so, we are convinced that when the technology industry presents use cases for the label generative AI (such as for LLMs, but not Bernoulli naive Bayes; Figure 1), this is a strategy to elicit confused responses against ‘generative AI,’ deflecting attention from other problematic AI systems (Table 1). This feeds into the hype and obstructs a more general critique of similar unethical or unwanted AI systems (cf. Guest 2025). Should we continue down this road, the technology industry will slip through such a distinction, skipping from terminology to terminology, consolidating its power. It has indeed done so many times before by using previous phrases as buzzwords like big data, machine learning, deep neural networks, and permutations thereof.

To further demystify the marketing phrase ‘AI’, we can examine the phrase’s constituent words. On the one hand, *artificial* intelligence is itself badly understood. When we compare artificial ‘intelligence’ with other artificial systems, we see that for example, artificial hearts do indeed pump blood (Dretske 1994; Kristan and Katz 2006; Millikan 2021; Powell 1970; Schellenberg 2018). This is because pumping blood is a straightforwardly understood aspect of a heart, artificial or otherwise. However, defining the function of a system does not translate so easily to AI. We cannot say what the function of intelligence in the general case is (Blokpoel 2018; Chirimuuta 2018; Egan 1999, 2017, 2018; Figdor 2010; Guest and Martin 2023; Hardcastle 1996; Rich et al. 2021; van Rooij, Guest, et al. 2024). As intelligence is not well defined, false claims of AI systems’ cognitive abilities, such as suggestions that such systems can communicate, read, or give feedback, seem appealing at first glance. However, through such claims fake purpose is invented, which can be exposed by cognitive science (Guest 2025; Guest and Martin 2025b; Guest, Scharfenberg, et al. 2025; van Rooij, Guest, et al. 2024). This fake purpose is illustrated by the fact that AI often does not function as it says on the tin (Bainbridge 1983; Brennan et al. 2025; Eaton 2025; Raji et al. 2022), leading to characterisations of AI as “snake oil” (Narayanan and Kapoor 2024), a “con” (Bender and Hanna 2025), “fake” (Kaltheuner 2021), and “fascist” (McQuillan 2022).

On the other hand, *intelligence* has a racist, sexist, classist, and ableist inheritance that it has not managed to shake off, from superficial pseudoscience to eugenics and genocide (Dennis 1995; Gould 1981; Norrgard 2008; Reddy 2007; Saini 2019). It is important to be aware that this sordid history rears its head in AI’s present, making these systems especially harmful to minoritised and vulnerable groups (Allen 2017; M. Andrews et al. 2024; Bates 2025; Benjamin 2019, 2024; Birhane 2022; Birhane, Prabhu, et al. 2023; Blas et al. 2025; Brennan et al. 2025; Dhaliwal et al. 2024; Erscoi et al. 2023; Evans 2020; Forbes and Guest 2025; Gebru and Torres 2024; Guest 2025; M. Hicks 2017; McQuillan 2025; Spanton and Guest 2022; S. M. Taylor et al. 2023; van der Gun and Guest 2024).

For all these reasons, we take the principled position that jargon infused with technology industry hype, such as shown in Table 1, does not meaningfully explain. Similarly, separating AI as a general term into good and bad leads nowhere except to blurring clarity and supporting hype. We strive to remain critical of the vocabulary the technology industry coopts and deploys, and to remain respectful of scientific terminology.

3 Higher education, critical thinking, expertise, & academic freedom

The culture of AI encourages a firm, even snide, conviction that it’s just a matter of time.
It thrives on exaggeration and refuses to examine its own failures.

Tom Athanasiou (1985, p. 18)

When it comes to AI technology used in a university context, it is important to focus on the relationship between a technology and society at large. The value of scholarly analyses of sociotechnical relationships can be seen in cases such as the following. In the UK, a school shooting in 1996, the Dunblane massacre, led to stricter gun control laws. Since then, the sociotechnical relationship between citizens and guns has remained unchanged, therefore handgun privileges have not been reintroduced to the public (N. Brown 1996; Shapiro et al. 2022). This relationship is different to that captured by the right to bear arms in the USA. Such analyses illustrate how different societies may create different legal frames to match how they perceive their relationship to technology.

Furthermore, problematising sociotechnical relationships itself allows for reclaiming AI as a scientific field from the grips of industry or hype (Birhane and Guest 2021; Guest and Forbes 2024; van Rooij, Guest, et al. 2024; Wendling 2002). Moreover, it allows for ruling out reclaiming as a function of opinions about possible and impossible changes to society, the technology, and their relationship

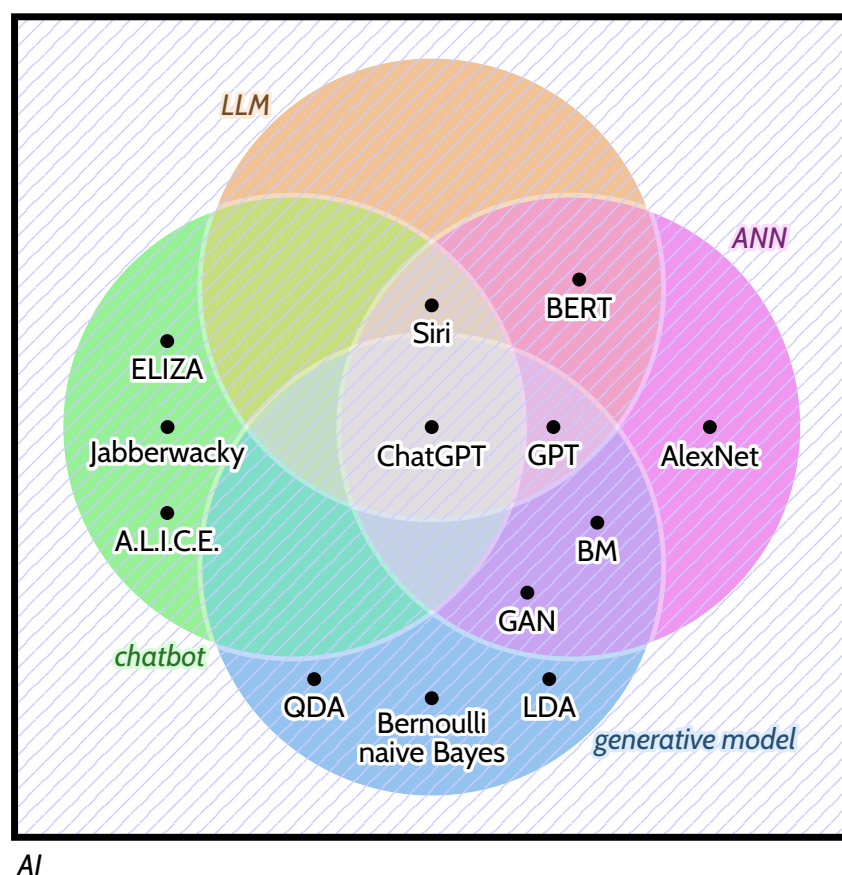


Figure 1. A cartoon set theoretic view on various terms (see Table 1) used when discussing the superset AI (black outline, hatched background): LLMs are in orange; ANNs are in magenta; generative models are in blue; and finally, chatbots are in green. Where these intersect, the colours reflect that, e.g. generative adversarial network (GAN) and Boltzmann machine (BM) models are in the purple subset because they are both generative and ANNs. In the case of proprietary closed source models, e.g. OpenAI’s ChatGPT and Apple’s Siri, we cannot verify their implementation and so academics can only make educated guesses (cf. Dingemans 2025). Undefined terms used above: BERT (Devlin et al. 2019); AlexNet (Krizhevsky et al. 2017); A.L.I.C.E. (Wallace 2009); ELIZA (Weizenbaum 1966); Jabberwacky (Twist 2003); linear discriminant analysis (LDA); quadratic discriminant analysis (QDA).

(Adams 2021; Avraamidou 2024; Forbes and Guest 2025; Whittaker 2021). To illustrate, consider driving in the general case versus driving a patient to hospital. Obviously, both produce a series of known pollutants. However, driving a patient to hospital has a different moral weight. Importantly, an expert paramedic driver has been taught to drive an ambulance regardless of fuel usage in emergency situations. However, the driver will not be taught to disregard fuel usage in other cases nor avoid public transport for personal travel. Similarly, AI systems have nuances that we as experts must analyse and impart on our students.

In this section, we tackle various positions on the AI-university relationship that appear to pass without critique in academic and wider spheres. We state our position in each case.

Table 1. Below some of the typical terminological disarray is untangled. Importantly, none of these terms are orthogonal nor do they exclusively pick out the types of products we may wish to critique or proscribe.

TERM	DESCRIPTION	RESOURCES
Artificial Intelligence (AI)	The phrase 'artificial intelligence' was coined by McCarthy et al. (1955) in the context of proposing a summer workshop at Dartmouth College in 1956. They assumed significant progress could be made on making machines think like people. In the present, AI has no fixed meaning. It can be anything from a field of study to a piece of software.	Avraamidou (2024), Bender and Hanna (2025), Bloomfield (1987), Boden (2006), Brennan et al. (2025), Crawford (2021), Guest (2025), Hao (2025), McCorduck (2004), McQuillan (2022), Monett (2021), Vallor (2024), and van Rooij, Guest, et al. (2024).
Artificial neural network (ANN)	First proposed in McCulloch and Pitts (1943), it is a mathematical model, comprised of interconnected banks of units that perform matrix multiplication and non-linear functions. These statistical models are exposed to data (input-output pairs) that they aim to reproduce. While held to be inspired by the brain, such claims are tenuous or misleading.	Abraham (2002), Bishop (2021), Boden (2006), Dhaliwal et al. (2024), Guest and Martin (2023, 2025a), Hamilton (1998), Stinson (2018, 2020), and Wilson (2016).
Chatbot	An engineered system that appears to converse with the user using text or voice. Speech synthesis goes back hundreds of years (Dudley 1939; Gold 1990; Schroeder 1966) and Weizenbaum's (1966) ELIZA is considered the first chatbot (Dillon 2020). Modern versions can contain ANNs in addition to hardcoded rules.	Bates (2025), Dillon (2020), Elder (2022), Erscoi et al. (2023), Schlesinger et al. (2018), Strengers et al. (2024), Turkle (1984), and Turkle et al. (2006).
ChatGPT	A proprietary closed source chatbot created by OpenAI. The for-profit company OpenAI has been steeped in hype from inception. It does not provide source code for most of its models, violating open science principles for academic users. OpenAI reported \$5 billion in losses in 2024 (Reuters 2025), and has received \$13 billion from Microsoft (Levine 2024).	Andhov (2025), Birhane and Raji (2022), Dupré (2025), Gent (2024), M. T. Hicks et al. (2024), Hill (2025), Jackson (2024), Kapoor et al. (2024), Liesenfeld, Lopez, et al. (2023), Mirowski (2023), Perrigo (2023), Titus (2024), and Widder et al. (2024).
Generative model	A specification on the type of statistical distribution modelled; typically contrasted with discriminative model. ANNs can be generative (e.g. Boltzmann machines) or discriminative (e.g. convolutional neural networks used for classifying images). In the context of generative AI or generative pre-trained transformer (GPT), this phrase is used inconsistently.	Efron (1975), Jebara (2004), Mitchell (1997), Ng and Jordan (2001), and Xue and Titterton (2008).
Large language model (LLM)	A model that captures some aspect of language, with the term 'large' denoting that the number of parameters exceed a certain threshold. Modern chatbots are often LLMs, which use ANNs, along with a graphical interface so that users can input so-called text 'prompts'. LLMs can be generative, discriminative, or neither.	Bender, Gebru, et al. (2021), Birhane and McGann (2024), Dentella et al. (2023, 2024), Leivada, Dentella, et al. (2024), Leivada, Günther, et al. (2024), Luitse and Denkena (2021), Shojaei et al. (2025a), Villalobos et al. (2024), and Wang et al. (2024).

3.1 *Rejection of expertise, ironically including our own*

Being in a colonizing discipline first demands and then encourages an attitude that might be called intellectual hubris. Furthermore, since you cannot master all the disciplines that you have designs on, you need confidence that your knowledge makes the ‘traditional wisdom’ of these fields unworthy of serious consideration. Here too, the AI scientist feels that seeing things through a computational prism so fundamentally changes the rules of the game in the social and behavioural sciences that everything that came before is relegated to a period of intellectual immaturity.

Sherry Turkle (1984, p. 230)

Every field that comes into contact with AI discourse becomes infected even within AI as a field of study (recall Table 1). Our colleagues have embraced these systems, uncritically incorporating them into their workflows and their classrooms, without input from experts on automation, cognitive science, computer science, gender and diversity studies, human-computer interaction, pedagogy, psychology, and law to name but a few fields with direct relevant expertise (Sloane et al. 2024). Meanwhile, technology companies have rushed to invest in ‘AI ethics’ or ‘AI safety’ to *ethics wash* their claims, thereby “laundering accountability” (as Abeba Birhane explains in Arseni 2025) and “distract[ing] from real AI ethics” (Crane 2021), while censoring academics and thus, violating academic freedom (Gebru and Torres 2024; Gerdes 2022; Goudarzi 2025; Munn 2023; Ochigame 2019; Suarez et al. 2025; Tafani 2023).

AI extractivism is at play further afield from the subversion described above — it also directly causes environmental and social harms: water, energy, and occupation of land are all needed for data centres (Goetze 2024; Gray and Suri 2019; Hao 2025; Loe 2023; Luccioni et al. 2024; Markelius et al. 2024; Parshley 2024; Perrigo 2023; Rowe 2023; Suarez et al. 2025; Tan 2025; Valinsky 2024). Yet, industry suggests AI technology can be used to mitigate the climate crises (Temple 2024). And it is creeping into society through data colonialism, stimulating the extraction and commodification of data, giving tech firms immense social power, and through labour exploitation and social injustices (Altenried 2022; Arora et al. 2023; Colón Vargas 2025; O’Neil 2016; Ricaurte 2019). In the academic context, promoting the use of LLMs does not align with university values and ecological sustainability campaigns. Hence, we cannot remain complicit with the greenwashing rhetoric and actions of the technology industry (Atkin 2025). Resisting AI technologies means refusing to take a role in the continuing devastation of the environment and the exploitation of labour.

When we speak out against the introduction of AI products in our classrooms we are “ridiculed with impunity” (Mirowski 2023, p. 740). However, as Djoerd Hiemstra (2023) explains: *we* are not the villains (cf. Schipper 2025). Sounding the alarm about AI as an educational technology is the only sensible thing until and unless we discuss the *purpose* of education regarding AI. Then, it is clear that the authors’ pedagogical goals and AI use are incompatible. What is urgently needed, instead, is *Critical AI Literacy*, such as we lay out in this paper (Heeg and Avraamidou 2024; McQuillan 2022; McQuillan et al. 2024; Monett and Paquet 2025; Suarez et al. 2025; Whittaker 2021). Importantly, *critical washing* — encouraging AI use while being ‘aware of the risks’ — must be avoided:

When we critique AI, we should do so with intellectual honesty and in a principled way. [R]eflecting on the harms of AI is not itself harm reduction. It may even contribute to rationalizing, normalizing, and enabling harm. Critical reflection without appropriate action is thus quintessentially critical washing.

Marcela Suarez et al. (2025, par. 7)

In fact, the vast majority of research on AI in education only examines issues connected to automation and assessment, and not learning (Avraamidou 2024). We propose that research exploring the use

of AI in education is shifted from a technical approach to a sociocultural and process approach that aims to rebalance qualification, socialisation, and subjectification (Biesta 2021). This should include concepts such as equity, language, multiculturalism, identity, affective domains of learning, citizenship and social justice, for a better understanding of what might be the role — if any — of AI tools in teaching and learning. We call on the educational technology community to demystify AI systems and instead approach those with more criticality and humility (cf. Tully et al. 2025). In doing so, we need to shift our attention from technology industry advertising to robust evidence-based research on developmental processes.

Unfortunately, academic contributions, from the humanities to mathematics, and everything in between, are all vulnerable to AI technology: products exist or are being developed that claim to be able to displace many forms of academic labour (Guest 2025). While most of these claims are based on thin air, administrators and policy makers see them as opportunities to further cut costs and increase efficiency, in line with the detrimental neoliberal ideology that suffuses higher education, further fostering rejection of expertise of their staff (Bal et al. 2014; Bouchard 2024; Gill and Donaghue 2016). Importantly, such dehumanisation by AI can happen regardless of intent (Birhane and Guest 2021; Brennan et al. 2025; Erscoi et al. 2023; Heffernan 2024; McQuillan 2022; O’Neil 2016; Rhee 2018; van der Gun and Guest 2024). A review of the literature shows that AI systems are biased and can take part in discriminatory practices that reinforce and normalize dominating hierarchies regarding culture, race, and gender while posing and being perceived as neutral and objective tools (Avraamidou 2024). We reject this debasing and dismantling of expertise, and dehumanisation of scholars, and remain vigilant against the erosion of academic contributions to understanding the world (CDH Team and Ruddick 2025; Crawley 2025; Helfrich 2024; Reed et al. 2001; Sano-Franchini et al. 2024).

3.2 *We do not have to ‘embrace the future’ & we can turn back the tide*

It must be the sheer magnitude of [artificial neural networks’] incompetence that makes them so popular.

Jerry A. Fodor (2000, p. 47)

Related to the rejection of expertise is the rejection of imagining a better future and the rejection of self-determination free from industry forces (Hajer and Oomen 2025; Stengers 2018; van Rossum 2025). Not only AI enthusiasts, but even some scholars whose expertise concentrates on identifying and critically interrogating ideologies and sociotechnical relationships — such as historians and gender scholars — unfortunately fall prey to the teleological belief that AI is an unstoppable force. They embrace it because alternative responses seem too difficult, incompatible with industry developments, or non-existent. Instead of falling for this, we should “refuse [AI] adoption in schools and colleges, and reject the narrative of its inevitability.” (Reynoldson et al. 2025, n.p., also Benjamin 2016; Campolo and Crawford 2020; CDH Team and Ruddick 2025; Garcia et al. 2022; Kelly et al. 2025; Lysen and Wyatt 2024; Sano-Franchini et al. 2024; Stengers 2018). Such rejection is possible and has historical precedent, to name just a few successful examples: Amsterdammers kicked out cars, rejecting that cycling through the Dutch capital should be deadly. Organised workers died for the eight-hour workday, the weekend and other workers’ rights, and governments banned chlorofluorocarbons from fridges to mitigate ozone depletion in the atmosphere. And we know that even the tide itself famously turns back. People can undo things; and we will (cf. Albanese 2025; Boztas 2025; Kohnstamm Instituut 2025; van Laarhoven and van Vugt 2025). Besides, there will be no future to embrace if we deskill our students and selves, and allow the technology industry’s immense contributions to climate crisis and environmental destruction to continue unimpeded (Benjamin 2024; Brennan et al. 2025; McQuillan 2025; Suarez et al. 2025; Tafani 2024b).

When we are told that we cannot ban AI — it is just too attractive — like with all industry hype we should both question the premises and caution that:

The idea of “irreversibility” points to the observation that, while technological trajectories are never self-determining or inevitable, the difficulties of undoing technological projects grow over time (Callon 1990).

Lucy Suchman (2019, p. 55)

It is exactly because problems continue to grow that we must speak out now. The repetition that somehow things are set in stone needs to be urgently, vociferously, and vigilantly countered, especially:

the public discourse that goes like this: ‘ChatGPT has happened and it’s the way forward and we should embrace it and we should learn to be creative about our assessments’. I’m not against being creative about assessments, but the idea of embracing ChatGPT seems incredibly toxic. I mean, do people not read how these things are made? They depend on eye watering — this is their own term — costs of carbon emissions. They are built on massive exploitative ‘ghost labour’; crowdsourced and outsourced labour that follows the patterns of colonial relations (Bender, Gebru, et al. 2021).

Dan McQuillan et al. (2024, p. 3)

Should we fail to turn the tide, we at least took a principled stand which is one of our core duties to society: to say what we think is the ideal thing to do based on our shared expertise. The rest of society might ignore us (often to our collective peril). But say it we must (Fuentes 2024).

3.3 *Ahistoricism and the AI hype cycles*

When I started writing about science decades ago [...] I edited an article in which [a computer scientist] predicted that AI would soon replace experts in law, medicine, finance and other professions. That was in 1984.

John Horgan (2020, n.p.)

When we engage with the public, we notice people think that AI, as a field or a technology, appeared on the scene in the last three years. And they experience confusion and even dissonance when they discover the field and the technologies have existed for decades, if not centuries or even millennia (Bloomfield 1987; Boden 2006; Bogost 2025; Guest 2025; Hamilton 1998; Mayor 2018). Such ahistoricism facilitates “the AI-hype cycles that have long been fuelled by extravagant claims that substitute fiction for science.” (Heffernan 2025, n.p. Duarte et al. 2024). We have been here before, both with entanglements of AI and statistics with industry corrupting our academic processes, and with so-called AI summers: hype cycles that pivot from funding booms to complete busts and cessation of research (Bassett and Roberts 2023; Boden 2006; Law 2024; Lighthill et al. 1973; Merchant 2023; Olazaran 1996; Perez 2002; P. Smith and L. Smith 2024; Thornhill 2025).

To understand how industry tries to influence independent research for their benefit, we can look to past examples of entanglement of industry and statistics. Ronald A. Fisher, a eugenicist and “the founder of modern statistics” (Rao 1992), having been paid by the tobacco industry, claimed that because ‘correlation is not causation’ that therefore ‘smoking does not cause lung cancer’ (Fisher 1958; Stolley 1991). The parallel between tobacco and technology does not end here: “both industries’ increased funding of academia was as a reaction to increasingly unfavourable public opinion and an increased threat of legislation.” (Mohamed Abdalla and Moustafa Abdalla 2021, p. 2; also see Knoester et al. 2025) The histories of eugenics, statistics, computing, and modern AI are highly interwoven

and inexorably intertwined (e.g. IBM facilitated the Holocaust; Black 2012; recall [section 2: *Marketing, hype, & harm*](#)). Tracing these genealogies, finding out what is salvageable is our job as scientists (Guest 2024, 2025). And as Amy Wendling (2002) explains “if our experience of machines has been outside of our experiences of exploitation, this has occurred only on the margins of the vastness of this exploitation, and is therefore conditioned by it.” (p. 172; also see Biesta 2021; Drimmer and Nygren 2025; Moore 1997; Watters 2023)

For the second point above, on how AI hype cycles appear to go hot and cold every few years, it is also rife with valuable teaching moments. For example, a news story like this would not be out of place in the present:

The Navy revealed the embryo of an electronic computer today that it expects will be able to walk, talk, see, write, reproduce itself and be conscious of its existence.

The New York Times (1958, n.p.)

This early example of AI hype was about the perceptron (an early ANN; recall [Table 1](#); Hay et al. 1960; Rosenblatt 1958). Frank Rosenblatt (1958), who is most associated with this, “los[t] hundreds of thousands of dollars a year in government funds” (Lefkowitz 2019) due to the subsequent AI winter (Boden 2006; Dreyfus 1965; Haigh 2023; Lighthill et al. 1973; McCorduck 2004; Olazaran 1996). These winters are triggered by over-promising and under-delivering; typical capitalist bubbles that we are all familiar with, from the dot-com bubble at the end of the 1990s to the subprime mortgage crisis starting in 2007. Not only would it be unprofessional to participate in such re-occurring folly, it would also damage our pedagogy to allow another cycle. The foreseeable AI winter will take with it entire curricula, academic processes and practices, and educators’ and learners’ livelihoods. Instead, we support historical awareness of the discourse of inevitability discussed herein.

This ahistoricism is not limited to the application of *AI* to education, we have seen a similar push of other technologies such as automated teaching machines, by (amongst others) B.F. Skinner (Biesta 2021; Drimmer and Nygren 2025; Watters 2023). Despite almost a century of trials, teaching and grading by computers have not reached classrooms, and the arguments made in favour of teaching machines mirror these used for AI today. For example, Pressey contrasted his machines with “mass media” (by which he meant radio) stressing his systems were “personalized.” The real effect seems to be “that it has been the means of turning the attention of many psychologists from research on animal learning in the laboratory to an examination of human learning in the classroom situation.” (Hoko 1986) In other words, rather than contributing to teaching and learning, automation in the classroom was benefiting the creators of technology.

3.4 *Anthropomorphism and other circular reasoning*

While opacity is a distinguishing feature of many other areas of science and technology, the myths surrounding computing may stem less from the fact that it is an opaque esoteric subject and more from the way in which it can be seen to blur the boundary between people and machines (Turkle 1984). To be sure, most people do not understand the workings of a television set or how to program their video cassette recorders properly, but then they do not usually believe that these machines can have intelligence. The public myths about computing and AI are also no doubt due to the ways in which computers are often depicted in the mass media — e.g. as an abstract source of wisdom, or as a mechanical brain.

Brian P. Bloomfield (1987, p. 72)

There is circular reasoning at play when we suggest and assume machines can think, reason, or argue like humans can, and therefore, treat them — and test them — like humans. Within human-machine

interaction research, often, AI technology output is compared to human performance, mistakenly assuming such benchmarks are informative about AI’s capabilities. However, correlations with human output mean little to substantiate claims of human-likeness, especially when the input to the AI models tested is the output of human cognition in the first place. There are so many cases of this from daily life and the history of science that it appears shocking such results are taken so uncritically to be cognition (Bernardi 2024; Guest 2025; Guest and Martin 2023; Placani 2024; van Rooij and Guest 2025). An example from the 1960s:

Weizenbaum (1966) was afraid of this new rhetorical trend of AI being seen as human by the public, blaming it on both overhyping by scientists and the gullibility of the users (cf. Dillon 2020, who argues that Weizenbaum’s own rhetoric also fuelled the fire of what came to be known as the ELIZA effect). Anthropomorphisation of AI systems, typified by what happened with ELIZA caused many at the time to be excited about the prospect of replacing therapists with software. More than half a century has passed since then and the idea of an automated therapist is still palatable to some, likely (thankfully) legally and ethically impossible without human supervision, and still very much out of reach technically.

Lelia Erscoi et al. (2023, pp. 22–23)

To add insult to injury, many metaphorical or jargonistic phrases (especially used to describe ANNs; see Table 1) — like train, learn, hallucinate, reason — are applied to machines and result in distorting how we perceive these machines: humanising them while dehumanising us (Barrow 2024; DeVrio et al. 2025; Erscoi et al. 2023; Heffernan 2024; Kambhampati et al. 2025; Placani 2024; Polo 2024; Raether 2025; Rhee 2018; Shojaee et al. 2025b; van der Gun and Guest 2024; cf. K. Andrews and Monsó 2021; Libell 2014). As Guest and Martin (2025a, p. 8) explain:

A mirror — not even an AI mirror — is not what it reflects (Vallor 2024). These jumps in logic can be found in many jocular stories, e.g. when a person first encounters a television, would they assume the device contained small people? Or would a person assume somebody was trapped inside a telephone because it emits human voices?

Although we do not fully understand human thinking, this does not licence attributing thinking to whichever machine or technology, uncritically and through anthropomorphisation. Such arguments from ignorance lack all scientific rigour. The only argument from ignorance that science permits is caution, more research, and care as appropriate actions when something is truly unknown. We take a stand against such anthropomorphism and inappropriate conclusions, and question why the parallel to machines is taken so lightly in such a harmful context, risking deskilling and dehumanization of us and our students.

3.5 *Supposedly students are all cheating now*

No serious scholar or scientist in their right mind would want LLMs to produce their texts; and hence, also no student pursuing an academic education would want to do so.

Iris van Rooij (2022, para. 7)

Students have always cheated. Bending and breaking the rules is human nature. And by the same token, educators are not police. We are not here to obsessively surveil our students — education is based on mutual trust. Therefore, our duty is to build mutually shared values with our students and colleagues. Especially when education is not valued, we as educators are obliged to show our students

that they are not just here to receive a degree: education is more than qualification (Biesta 2021). It is about preparing students to become a capable and active members of society.

We emphasize that there are two victims of plagiarism: the original authors whose work is taken without credit and the audience who is being deceived.

Plagiarism, at its most fundamental level, is a lie. It is the taking of works or ideas of others and passing them off as your own, either directly or indirectly. The misdeed itself is in the lie, the “I created this” when it is known to be untrue.

However, that lie isn’t being told to the original victim. It’s a lie about the victim, claiming that they didn’t create it or their contributions didn’t matter, but it’s not a lie to them. Instead, it’s a lie to the audience, which is the second victim and the actual target of the con.

Jonathan Bailey (2019b, n.p.)

When a ghost writer or essay mill is used, consent has been obtained by the original author (usually in exchange for payment), but acceptability remains contingent on the social contract with the audience. Ghostwriting is usually acceptable because “[w]e all know that politicians don’t write their own speeches and celebrities rarely write their own books” (Bailey 2019b). In contrast, *press release plagiarism* (Bailey 2019a) deceives the audience by presenting the work of a biased source as being that of the (perceived to be unbiased) journalist. Academic integrity standards, enshrined in university rules and scholarly journal policies, establish the unequivocal contract that students’ and researchers’ work is solely their own. There is no debate that use of an essay mill constitutes violation of these policies.

Use of an LLM-based chatbot (recall Table 1) to complete an assignment, presented as student work, is an affront to the ‘first victim’ as well as the ‘second victim’. This automated plagiarism goes against the principles of scientific integrity, which serve as the basis of our teaching and academic practice (van Rooij 2022). LLM and chatbot use in the classroom means accepting that both students and teachers can openly appropriate work that is not theirs, and requires disclosing to society (a secondary audience) that thinking and integrity are not required. Using AI products to generate academic content is against any conceivable set of rules we already have in place to regulate fraud and make sure that students actually, when following the rules, learn what they are expected to learn. Failure to see this buys into the hype and if retrofitted onto past cases brings up essay mills as a permissible educational ‘tool’ — clearly anti-pedagogical and utterly bizarre.

Promoting the use of LLMs in the classroom will only increase illiteracy in young adults and teenagers and deepen the crisis in education; further, it will boost technology-industry dependence (Forbes and Guest 2025; Lambeets 2025; Meelissen et al. 2023; U.S. Department of Education, National Center for Education Statistics 2024). Writing — as visual and performing arts; all forms of human expression — forms a bedrock of the learning process (Warner 2025).. And many students, regardless of the mantra that supposedly they are all using it, can tell AI is harmful (Abrams 2025; Drapkin 2023; Eichenberger et al. 2025; Huntington 2025; Isayas 2025b; Kaplan 2024; Klee 2025; Landymore 2025; Montgomery 2024; Neville 2025; Pejcha 2023; Purtill 2025; Roose 2024; Schoene and Canca 2025; Tafani 2024b; J. Taylor 2025; Tiku 2025; Wei 2025; Xiang 2023) and useless (Bearne 2025; Bond 2025; Doetie Talsma 2025; Otis et al. 2025; Pahwa 2025; Soler 2025). In one survey by Study.com (2025), for instance, “72% of college students believed that ChatGPT should be banned from their college’s network” (also see Akolekar et al. 2025; Morris 2025; TechEquity 2025).

Despite all the current hysteria around students cheating, they aren’t the ones to blame. They did not lobby for the introduction of laptops when they were in elementary school, and it’s not their fault that they had to go to school on Zoom during the pandemic.

They didn’t create the A.I. tools, nor were they at the forefront of hyping technological innovation. [...]

None of the students I spoke with seemed lazy or passive.

Hua Hsu (2025, n.p.)

The quote above touches on, and neutralises, one of the many marketing strategies to introduce such so-called technologies in the classroom (e.g. Chow 2025). Technology companies are not shy of falsely claiming that students are lazy or lack writing skills. Such a mantra serves only to sell products — or cover up and excuse overworking them by our colleagues — with no reflection on reality. We condemn those claims and reassert students’ agency vis-à-vis corporate control.

3.6 *Do we have to teach AI because otherwise our students will not get jobs?*

Headline after headline proclaimed “The Robots Are Coming for Our Jobs,” in stories about hyper-intelligent AI and logistics automation. Will all these trends eventually lead to a world where the bots and algorithms do our dirty work, making our lives easier and more prosperous? Or will the machines push us out of our jobs and deposit us into a dystopia? We are chewing over the machinery question all over again, in barely updated language.

Brian Merchant (2023)

Of course, we have to teach our students *about* AI technologies. Teaching about AI technologies should be just like how we teach ‘no smoking’ or the causal links between lung cancer and cigarette smoke; yet, we do not teach students how to roll cigarettes and smoke them. In fact, the whole selling point of the technology industry is that modern AI technologies are like talking to another person, often a servant. Thus, what training is there to really have under such a frame? What is there to teach with respect to typing some text into a box and then mindlessly copy-pasting the output? Such a skill-less relationship disregards our pedagogical commitment to foster a critical stance in our students and colleagues (cf. Previtali and Fagiani 2015). Some counteract that students can and must “use AI critically/responsibly”, e.g. by judging the output for themselves (cf. Barr 2025). Going one step further, some even maintain that such exercises are useful to train critical thinking itself, thereby buying into the usefulness of AI despite admitting to its untrustworthiness, glossing over the many problems we describe here.

In thinking about implications for the design of learning environments and curriculum design, we first need to pause and think about what we really would like AI tools to do, or, put differently, what might be the added value of the use of AI tools in education — if any? In thinking about the future of education, we ought to be engaging with culturally relevant and sustainable pedagogies that have the potential to embrace the diversity of learners’ identities, cultural capital, and migration backgrounds through culturally and personally meaningful learning experiences (Bamber et al. 2017; Barab et al. 1999; MacPartland et al. 1971). Hence, taking teachers or students out of a social context and in return giving them AI products, would be disastrous for education and for society more generally.

Relatedly, the (extreme) obfuscation of labour appears to be a defining factor of AI technology (Altenried 2022; Arora et al. 2023; L. M. Brown 2023; Colón Vargas 2025; Gray and Suri 2019; Guest 2025; O’Neil 2016; Perrigo 2023; Pfaffenberger 1988; Ricaurte 2019). AI stands in contrast to a tool like a saw involved in the predominantly overt labour of woodworking (that is, to cut wood), where the person cutting also puts in labour with actual control over the output of the labour; often more than the creator of the saw. AI users, on the other hand, are customers much more like the person buying the end product of woodwork than carpenters themselves. If not in some sense more so, as

they remain unaware of, and are even tricked into thinking they performed the relevant labour. This distortion can be seen in often used obfuscatory phrases like ‘prompt engineer’, which imply labour and skill (Guest 2025; Watters 2023). Being a customer not only requires no training, is not a skill, but also is a completely different labour relation masquerading as skill. Thus, seeing students merely as customers of AI products deskills not only them, but also us as teachers (as there is nothing to teach), and ultimately society at large (as no one will be left to be able to write and think). This is reflected in the advertising deployed by AI companies who claim prompts are akin to ‘talking to a computer’ and require no skills like programming to do so.

What future and what jobs will there be if most human capacities really will be done away with? Crucial professions necessary to maintaining a well-functioning society, democracy, healthcare and justice systems, would deteriorate and degenerate. The necessary creative, empathic, and intellectual skills needed to ensure these professions can serve their societal purposes, and uphold relevant standards, would be lost (Budzyń et al. 2025). Ironically, in a world full of AI displacements (Guest 2025) and AI misinformation, actually *more* skill would be needed to catch, counter, and compensate for the problems and mistakes introduced by AI technologies (Bainbridge 1983), but as AI technologies increasingly infiltrate academia, students will be robbed of the opportunity to learn those necessary skills. We take a stand against such developments in which students cannot learn — by definition — and reject the frames that AI use is necessary to obtain a job.

3.7 *AI is more like a calculator than you think*

“But I just use AI for boilerplate!” you whimper, clutching your Co-Pilot subscription. Listen to yourself. If you’re writing the same boilerplate every day like some industrial-age cog monkey, automate it yourself. Write a library. Invent a macro. Reclaim some dignity. If AI’s doing your “boring parts”, what exactly is left for you to do? [...]

When you outsource the thinking, you outsource the learning.

Jj (2025, n.p.)

There are absolutely important differences between an LLM and a calculator, notably the second is not stochastic outside, for example, the battery running out or a cosmic ray flipping a bit. Contrary to an LLM, under normal conditions, a calculator performs pre-specified functions reliably (for relevant analyses, see Guest 2025). But — and herein lies the rub — we ban calculators when teaching children addition and other basic arithmetic operations for a reason (cf. Lodge et al. 2023). Otherwise, they would not learn these arithmetic operations, and calculators do not help to understand the basic mathematical rules. For the same reasons, we also do not allow the use of spellcheck software for children learning to spell, or keyboard typing when learning to write by hand (Ihara et al. 2021; Lambert 2024; Longcamp et al. 2005; Mueller and Oppenheimer 2014; Rascoe 2024). The same reasoning applies to banning AI products in education (cf. Bond 2025; Forbes and Guest 2025; Guest 2025; Weatherbed 2025):

LLMs do not improve one’s writing ability much like taking a taxi does not improve one’s driving ability. Students should hone their writing, thinking, and other academic skills at every opportunity. As teachers, we are duty-bound to help them realise why they should. Importantly, writing creates a space for students to assimilate and create knowledge independently of industry influence, and gives students a way to empower their voices. Given the technology industry’s concentration of power, infrastructure, and control, we must promote writing at university along with teaching critical AI literacy. What holds for writing is, of course, true for any cognitive skill to-be-learned in education, such as programming skills (Becker et al. 2025; cf. Guest and Forbes 2024) and these too suffer deskilling if not continuously practised:

I stared at my terminal facing those red error messages that I hate to see. An AWS [Amazon Web Services] error glared back at me. I didn’t want to figure it out without AI’s help.

After 12 years of coding, I’d somehow become worse at my own craft. And this isn’t hyperbole—this is the new reality for software developers.

Namanyay Goel (2025, n.p.)

To show how serious the situation has become, one need only think about our last round of marking essays by AI undergraduate students. What jumps out of the page, for us, is something that contradicts the rhetoric our colleagues promote, namely, it is evident that students need *more* essay work assigned to them, not less (Kosmyna et al. 2025). Almost every essay was poor on some dimension that does not befit students in their final years of undergraduate study: the writing is often superficial, the language does not reflect students’ stage and knowledge, citations are frequently misused, and (most shockingly because it is so easy), the reference style is not applied correctly. This means that the constellation of skills required to write a good academic essay has not been nurtured enough or has atrophied. What this means is also that regardless of factual LLM use by the students, *their ability to write essays is on the floor*, and not, as many seem to claim, at ceiling where one cannot differentiate a good essay from a plagiarised or otherwise dishonest attempt of an essay. Importantly, the training of writing skills should be done in the context of critical reckoning with the norms and pressures surrounding the work expected of students (i.e. high study load, so-called student excellence, financial pressure to graduate, etc.).

In this context, it is also important to be wary of arguments that wrongly position LLMs as, making education more democratic, accessible, and equitable by removing language barriers, removing unequal access to mentorship, and increase diversity, equity and inclusion in general. LLMs are neither a real nor just solution to these issues given the myriad problems and falsehoods we raise here, *and* given the fact that the harms associated with LLMs feed directly into the roots of these inequalities to begin with (Sano-Franchini et al. 2023). When a (perceived) lack of writing skills is penalized along racist, ableist, or otherwise discriminatory lines (Flores and Rosa 2015), students may be driven to LLMs because they allegedly produce ‘professional’ text. We need to commit to creating an environment in which such rhetoric and injustice are tackled at their roots.

4 Protecting the Ecosystem of Human Knowledge: Five Principles

We must protect and cultivate the ecosystem of human knowledge. AI models can mimic the appearance of scholarly work, but they are (by construction) unconcerned with truth — the result is a torrential outpouring of unchecked but convincing-sounding “information”. At best, such output is accidentally true, but generally citationless, divorced from human reasoning and the web of scholarship that it steals from. At worst, it is confidently wrong. Both outcomes are dangerous to the ecosystem.

Olivia Guest, Iris van Rooij, et al. (2025, n.p.)

Knowledge production is supposed to be safeguarded by (inter)national codes of conduct for research integrity (Allea 2023; KNAW et al. 2018). Such codes forbid, for instance, fabrication of data, falsification of results, plagiarism, and, generally, distortion of the scientific record. Many argue that new rules are required to regulate academic AI use, but pre-existing guidelines fit the bill (cf. Tafani 2024a). For instance, the Netherlands Code of Conduct for Research Integrity (KNAW et al. 2018) is based on 5 core principles — each of which speaks to AI usage in one or more ways (cf. Dingemanse 2024):

Honesty implies that we do not secretly use AI technologies without disclosure, and that one does not make unfounded claims about the presumed capabilities of AI technologies (this also follows from *Responsibility*; see example from MIT Economics 2025, where perhaps too little too late was done).

Scrupulousness demands, among other things, that scientists only use AI products whose functionality is well-specified and validated for its specific scientific usage (cf. Kwisthout 2024; Kwisthout and Renooij 2025). This includes terminological precision about what formalisms, models, and/or technologies are used (recall Figure 1) and rigorous argumentation to motivate why these technologies are appropriate for the scientific purposes at hand.

Transparency requires that the AI technologies are open source and computationally reproducible. Here we must recall the technology industry's obfuscatory tactics: "the name of the current producer of ChatGPT. 'OpenAI' sounds like it is engaged in open science, but as we have now seen, 'open' never really means what you think it does." (Mirowski 2023, p. 738; see also Dingemanse 2025; Hao 2025; Jackson 2024; Liesenfeld and Dingemanse 2024; Liesenfeld, Lopez, et al. 2023; Maffulli 2023; Maris 2025; Nolan 2025; Solaiman 2023; Thorne 2009; Widder et al. 2024)

Independence means that scientists ensure that their research is unbiased by AI companies' agendas, and that any potential conflicts of interest are declared in publications and other public communications (this also follows from *Honesty* and *Transparency*; cf. Mohamed Abdalla and Moustafa Abdalla 2021; Atkin 2025; Forbes and Guest 2025; Knoester et al. 2025).

Responsibility precludes scientists from using AI products whose use is irresponsible, e.g. harmful to people, animals, and the environment, or otherwise in violation of legal guidelines (e.g. copyright, data privacy, labour laws, Butterick 2025; Cole 2025; Rijo 2025; Tafani 2024a). Minimizing harm is vital for both engineers (ACM Code 2018 Task Force 2018) and theoreticians (Guest 2024).

All of these principles rely on the premise of reciprocity. This is why AI, as an opaque and extractive technology, fails to meet the basic threshold of academia.

Relatedly, students should be encouraged and able to independently apply the core integrity principles to all relevant aspects of their scientific education in a context-sensitive manner. This means that every university graduate must be aware of the relevant scientific codes of conduct, have the competence to comply with them, and recognize violations. This is especially important for thesis research, but ideally relevant skills are trained as an integral part throughout academic degree programs. However, with the proliferation of AI products and their uncritical adoption in academia, we become unable to help younger generations of scholars in learning to uphold and to appreciate scientific integrity. As a result, we will be deskilling the whole academic profession, a direct threat to the ecosystem of human knowledge (cf. Barr 2025).

Of course, the academic world has well-known structural incentives to cheat, including high workloads, 'publish or perish' contracts, and framings of academic research as an elite competitive marketplace for individual prestige. These institutional conditions have encouraged, and rewarded, unethical behaviour long before the current AI hype cycle, but the *normalization* of AI usage and the ubiquitous false claims regarding AI capabilities seem to further obfuscate the boundaries of proper conduct. Thus, what is true for plagiarism and essay mills for students, holds for paper mills harming journals, the literature, and us as academics (Matusz et al. 2025; Oransky and Marcus 2023) — and AI turbocharges these already understood and proscribed harms (Larousserie 2025; Mirowski 2023). These entanglements between the technology industry and the academic publishing houses

(e.g. Google and Springer Nature, Spanton and Guest 2022; The Luddite 2024) require urgent attention (Markov 2024; Mirowski 2023; van Rooij 2025).

With the global rise of fascistic regimes and growing threats to academic freedom (ALLEA 2025; Brennan et al. 2025; KNAW 2021, 2025; McQuillan 2025), safeguarding the integrity of scientific knowledge is more important now than ever. With dismay we witness our university leadership making soulless choices that hollow out our institutions from within and erode the critical and self-reflective fabric of academia. Notable among these choices is the relentless push to uncritically adopt AI technologies in education, under false narratives of so-called progress and efficiency (Brennan et al. 2025). These false narratives have even been repeated by Dutch university unions in labour agreements (FNV et al. 2025) — with friends like that, who needs enemies?

In addition to scientific integrity, university core values — such as sustainability, openness, responsibility, critical reflection, and diversity, equity and inclusion (Eindhoven University of Technology 2025; HWR Berlin 2025; Radboud University 2025) — should naturally constrain AI use in academia. The KNAW Young Academy is sounding the alarm about digital technologies in higher education, noting that “decision-making [in higher education institutions] typically considers the financial and logistical aspects, similar to buying furniture or office supplies, but it falls short when it comes to addressing the consequences for governance and the ethical implications” (Young Academy 2025, p. 6) (see also Isayas 2025a; Paris et al. 2025). By normalizing the use of AI products in academia, universities contravene codes of conduct and contradict their own missions and betray their core values, harming science and society in the process (cf. Previtali and Fagiani 2015).

5 Machine Yearning for a Better Present

The set of AI products and technologies “stands before us, in other words, in what Marx would call fetishised form: *what is in reality produced by relations among people appears before us in a fantastic form as relations among things.*” (Pfaffenberger 1988, p. 242) In line with this, AI technology is much more akin to a ouija board than anything else: promoting and enabling mysticism, pseudoscience, and pseudo-intellectualism (section 2: *Marketing, hype, & harm*); obfuscating the human-in-the-loop, which includes the user and the creators of the data, who do most of the ‘intelligence’ and provide all the inputs (subsection 3.4: *Anthropomorphism and other circular reasoning*). And combining all the above, tricking and scamming us because we do not recognise that it merely reflects our own collective intelligence back at us in distorted and damaging forms.

We and our students can choose not to use these technologies. Just like we have banned smoking from public spaces, we could foster that process of banning both by choosing to individually quit smoking and by demanding regulation of the tobacco industry. Hence, we can also choose to unshackle ourselves and others from harmful technology industry frames, products, and lifestyles. Yet, we acknowledge that:

The entrepreneurs who have disrupted key elements of the economy, [have] accumulated unprecedented levels of wealth and power in the process. [...] But these men—just like two hundred years ago, they are almost all men—have power on a scale that a nineteenth-century titan could not even begin to dream of. [...] Whether these behemoths are deploying AI or automating work outright, they have a level of control over most people’s working lives that is unrivalled in its breadth and scope.

Brian Merchant (2023)

Importantly and “[c]ontrary to popular belief, high technology is often as socially regressive as it is technically revolutionary or progressive.” (M. Hicks 2017, p. 17) This is one of many reasons why we must resist and reject these technologies; and why we, as academics, should unpack further why

AI causes “nightmares” (Blix and Glimmer 2025; Hao 2025; Jucan 2023; Nick Dyer-Witheford 2019; Popper 1972).

Ultimately, these systems cannot really replace humans, replace the quality of human craft and thinking — so many of their capacities are overblown and displacement will only happen if we accept the premises (Guest 2025). We can and should reject that AI output is ‘good enough,’ not only because it is *not* good, but also because there is inherent value in thinking for ourselves. We cannot all produce poems at the quality of a professional poet, and maybe for a complete novice an LLM output will seem ‘better’ than ones’ own attempt. But perhaps that is what being human is: learning something new and sticking with it, even if we do not become world famous poets (Brainard 2025).

That work — the real work of teaching and learning — cannot be automated.

Jesse Hagopian (2025, n.p.)

References

- AAUP (2025). *Artificial Intelligence and Academic Professions*. URL: <https://www.aaup.org/reports-publications/aaup-policies-reports/topical-reports/artificial-intelligence-and-academic>.
- Abdalla, Mohamed and Moustafa Abdalla (2021). “The Grey Hoodie Project: Big Tobacco, Big Tech, and the Threat on Academic Integrity”. In: *Proceedings of the 2021 AAAI/ACM Conference on AI, Ethics, and Society*. AIES ’21. ACM, pp. 287–297. DOI: [10.1145/3461702.3462563](https://doi.org/10.1145/3461702.3462563).
- Abraham, Tara H (2002). “(Physio)logical circuits: The intellectual origins of the McCulloch-Pitts neural networks”. In: *Journal of the History of the Behavioral Sciences* 38.1, pp. 3–25. DOI: [10.1002/jhbs.1094](https://doi.org/10.1002/jhbs.1094).
- Abrams, Zara (2025). “Using generic AI chatbots for mental health support: A dangerous trend”. In: *American Psychological Association*. URL: <https://www.apaservices.org/practice/business/technology/artificial-intelligence-chatbots-therapists>.
- ACM Code 2018 Task Force (2018). *ACM Code of Ethics and Professional Conduct*. URL: <https://www.acm.org/code-of-ethics>.
- Adams, Rachel (2021). “Can artificial intelligence be decolonized?” In: *Interdisciplinary Science Reviews* 46.1-2, pp. 176–197. DOI: [10.1080/03080188.2020.1840225](https://doi.org/10.1080/03080188.2020.1840225).
- Akolekar, Harshal et al. (2025). “The role of generative AI tools in shaping mechanical engineering education from an undergraduate perspective”. In: *Scientific Reports* 15.1, p. 9214. DOI: [10.1038/s41598-025-93871-z](https://doi.org/10.1038/s41598-025-93871-z).
- Albanese, Francesca (2025). *A/HRC/59/23: From economy of occupation to economy of genocide - Report of the Special Rapporteur on the situation of human rights in the Palestinian territories occupied since 1967 (Advance edited version)*. URL: <https://www.ohchr.org/en/documents/country-reports/ahrc5923-economy-occupation-economy-genocide-report-special-rapporteur>.
- Alkhatib, Ali (2024). *Defining AI*. URL: <https://ali-alkhatib.com/blog/defining-ai>.
- ALLEA (2025). *ALLEA Statement on Threats to Academic Freedom and International Research Collaboration in the United States*. DOI: [10.26356/ALLEA-ACADEMIC-FREEDOM-STATEMENT](https://doi.org/10.26356/ALLEA-ACADEMIC-FREEDOM-STATEMENT).
- Allea (2023). *The European Code of Conduct for Research Integrity*. URL: <https://allea.org/code-of-conduct/>.
- Allen, Kera Jones (2017). “Hidden Figures: The American Dream and the Untold Story of the Black Women Mathematicians Who Helped Win the Space Race by Margot Lee Shetterly”. In: *IEEE Annals of the History of Computing* 39.3, pp. 70–71.
- Altenried, Moritz (2022). “The digital factory: The human labor of automation”. In: *The Digital Factory*. University of Chicago Press.

- Andhov, Alexandra (2025). "OpenAI's Transformation: From a Non-profit to a 157 Billion Valuation". In: *Business Law Review* 46.1, pp. 2–11. DOI: [10.54648/bula2025003](https://doi.org/10.54648/bula2025003).
- Andrews, Kristin and Susana Monsó (2021). "Animal Cognition". In: *The Stanford Encyclopedia of Philosophy*. Ed. by Edward N. Zalta. Spring 2021. Metaphysics Research Lab, Stanford University. URL: <https://plato.stanford.edu/archives/spr2021/entries/cognition-animal/> (visited on 08/27/2025).
- Andrews, Mel, Andrew Smart, and Abeba Birhane (2024). "The reanimation of pseudoscience in machine learning and its ethical repercussions". In: *Patterns* 5.9. DOI: [10.1016/j.patter.2024.101027](https://doi.org/10.1016/j.patter.2024.101027).
- Arora, Anmol et al. (2023). *Risk and the future of AI: Algorithmic bias, data colonialism, and marginalization*. DOI: [10.1016/j.infoandorg.2023.100478](https://doi.org/10.1016/j.infoandorg.2023.100478).
- Arseni, Maurizio (2025). "UN AI summit accused of censoring criticism of Israel and big tech over Gaza war - Geneva Solutions". In: *Geneva Solutions*. URL: <https://genevasolutions.news/science-tech/un-ai-summit-accused-of-censoring-criticism-of-israel-and-big-tech-over-gaza-war>.
- Athanasίου, Tom (1985). "Artificial intelligence: cleverly disguised politics". In: *Compulsive technology: computers as culture*. Ed. by Tony Solomonides and Les Levidow. Free Association Books, pp. 13–35.
- Atkin, Emily (2025). *He helped Microsoft build AI to help the climate. Then Microsoft sold it to Big Oil*. URL: <https://heated.world/p/he-helped-microsoft-build-ai-to-help>.
- Avraamidou, Lucy (2024). "Can we disrupt the momentum of the AI colonization of science education?" In: *Journal of Research in Science Teaching* 61.10, pp. 2570–2574. DOI: [10.1002/tea.21961](https://doi.org/10.1002/tea.21961).
- Bailey, Jonathan (2019a). *The Problem with Press Release Plagiarism*. URL: <https://www.plagiarismtoday.com/2019/03/21/the-problem-with-press-release-plagiarism/> (visited on 08/23/2025).
- (2019b). *The Two Victims of Plagiarism*. URL: <https://www.plagiarismtoday.com/2019/08/01/the-two-victims-of-plagiarism/> (visited on 08/23/2025).
- Bainbridge, Lisanne (1983). "Ironies of automation". In: *Automatica* 19.6, pp. 775–779. DOI: [10.1016/0005-1098\(83\)90046-8](https://doi.org/10.1016/0005-1098(83)90046-8).
- Bal, Ellen, Erella Grassiani, and Kate Kirk (2014). "Neoliberal individualism in Dutch universities: Teaching and learning anthropology in an insecure environment". In: *Learning and Teaching* 7.3, pp. 46–72. DOI: [10.3167/latiss.2014.070303](https://doi.org/10.3167/latiss.2014.070303).
- Bamber, Philip, David Lewin, and Morgan White (2017). "(Dis-) Locating the transformative dimension of global citizenship education". In: *Journal of Curriculum Studies*, pp. 204–230. DOI: [10.1080/00220272.2017.1328077](https://doi.org/10.1080/00220272.2017.1328077).
- Bandura, Albert (2001). "Social cognitive theory: An agentic perspective". In: *Annual review of psychology* 52.1, pp. 1–26. DOI: [10.1146/annurev.psych.52.1.1](https://doi.org/10.1146/annurev.psych.52.1.1).
- Barab, Sasha A. et al. (1999). "Principles of self-organization: Learning as participation in autocatalytic systems". In: *Journal of the Learning Sciences*, pp. 349–390. DOI: [10.1080/10508406.1999.9672074](https://doi.org/10.1080/10508406.1999.9672074).
- Barr, Alistair (2025). "Thanks to ChatGPT, the pure internet is gone. Did anyone save a copy?" In: *Business Insider*. URL: <https://www.businessinsider.com/chatgpt-ai-pure-internet-gone-low-background-steel-cloudflare-2025-6?international=true&r=US&IR=T>.
- Barrow, Nicholas (2024). "Anthropomorphism and AI hype". In: *AI and Ethics* 4.3, pp. 707–711. DOI: [10.1007/s43681-024-00454-1](https://doi.org/10.1007/s43681-024-00454-1).
- Bassett, Caroline and Ben Roberts (2023). "Automation anxiety: A critical history-the apparently odd recurrence of debates about computation, AI and labour". In: *Handbook of critical studies of*

- artificial intelligence*. Edward Elgar Publishing, pp. 79–93. DOI: [10.4337/9781803928562.00012](https://doi.org/10.4337/9781803928562.00012).
- Bates, Laura (2025). *The New Age of Sexism: How AI and Emerging Technologies Are Reinventing Misogyny*. Sourcebooks.
- Bearne, Suzanne (2025). “The people refusing to use AI”. In: *BBC News*. URL: <https://www.bbc.com/news/articles/c15q5qzdjqxo>.
- Becker, Joel et al. (2025). “Measuring the Impact of Early-2025 AI on Experienced Open-Source Developer Productivity”. In: *arXiv preprint arXiv:2507.09089*. DOI: [10.48550/arXiv.2507.09089](https://doi.org/10.48550/arXiv.2507.09089).
- Bender, Emily M. (2024). *Opening remarks on “AI in the Workplace: New Crisis or Longstanding Challenge”*. URL: <https://medium.com/@emilymenonbender/opening-remarks-on-ai-in-the-workplace-new-crisis-or-longstanding-challenge-eb81d1bee9f>.
- Bender, Emily M., Timnit Gebru, et al. (2021). “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” In: *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*. FAccT ’21. Virtual Event, Canada: Association for Computing Machinery, pp. 610–623. DOI: [10.1145/3442188.3445922](https://doi.org/10.1145/3442188.3445922).
- Bender, Emily M. and Alex Hanna (2025). *The AI Con: How to Fight Big Tech’s Hype and Create the Future We Want*. HarperCollins.
- Benjamin, Ruha (1, 2016). “Informed Refusal: Toward a Justice-based Bioethics”. In: *Science, Technology, & Human Values* 41.6, pp. 967–990. DOI: [10.1177/0162243916656059](https://doi.org/10.1177/0162243916656059).
- (2019). *Race after technology*. Polity.
- (2024). “The New Artificial Intelligentsia”. In: *Los Angeles Review of Books*. URL: <https://lareviewofbooks.org/article/the-new-artificial-intelligentsia/>.
- Bernardi, Guilherme (2024). “Limites absolutos ou relativos para a substituição de postos de trabalho por inteligência artificial?” In: *Revista Eletrônica Internacional de Economia Política da Informação da Comunicação e da Cultura* 26.2, pp. 109–124. DOI: [10.54786/revistaepitic.v26i2.20833](https://doi.org/10.54786/revistaepitic.v26i2.20833).
- Biesta, Gert (2021). *World-centred education: A view for the present*. Routledge.
- Birhane, Abeba (2022). “The unseen Black faces of AI algorithms”. In: *Nature* 610, pp. 451–452. DOI: [10.1038/d41586-022-03050-7](https://doi.org/10.1038/d41586-022-03050-7).
- Birhane, Abeba, Sepehr Dehdashtian, et al. (2024). “The Dark Side of Dataset Scaling: Evaluating Racial Classification in Multimodal Models”. In: *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*. FAccT ’24. Rio de Janeiro, Brazil: Association for Computing Machinery, pp. 1229–1244. DOI: [10.1145/3630106.3658968](https://doi.org/10.1145/3630106.3658968).
- Birhane, Abeba and Olivia Guest (2021). “Towards Decolonising Computational Sciences”. In: *Kvinde, Kan & Forskning* 29.1, pp. 60–73. DOI: [10.7146/kkf.v29i2.124899](https://doi.org/10.7146/kkf.v29i2.124899).
- Birhane, Abeba and Marek McGann (2024). “Large models of what? Mistaking engineering achievements for human linguistic agency”. In: *Language Sciences* 106, p. 101672. DOI: [10.1016/j.lanngsci.2024.101672](https://doi.org/10.1016/j.lanngsci.2024.101672).
- Birhane, Abeba and Vinay Uday Prabhu (2021). *Large image datasets: A pyrrhic win for computer vision?* DOI: [10.1109/WACV48630.2021.00158](https://doi.org/10.1109/WACV48630.2021.00158).
- Birhane, Abeba, Vinay Uday Prabhu, et al. (2023). “On hate scaling laws for data-swamps”. In: *arXiv preprint arXiv:2306.13141*. DOI: [10.48550/arXiv.2306.13141](https://doi.org/10.48550/arXiv.2306.13141).
- Birhane, Abeba and Inioluwa Deborah Raji (2022). “ChatGPT, Galactica, and the progress trap”. In: *WIRED*. URL: <https://www.wired.com/story/large-language-models-critique/>.
- Bishop, J Mark (2021). “Artificial intelligence is stupid and causal reasoning will not fix it”. In: *Frontiers in Psychology* 12, p. 513474. DOI: [10.3389/fpsyg.2020.513474](https://doi.org/10.3389/fpsyg.2020.513474).
- Black, Edwin (2012). *IBM and the Holocaust: The Strategic Alliance Between Nazi Germany and America’s Most Powerful Corporation-Expanded Edition*. Dialog Press.

- Blas, Zach, Melody Jue, and Jennifer Rhee (2025). *Informatics of Domination*. Duke University Press Books.
- Blix, Hagen and Ingeborg Glimmer (2025). *Why we fear AI: On the Interpretation of Nightmares*. Common Notions.
- Blokpoel, Mark (2018). "Sculpting Computational-Level Models". In: *Topics in cognitive science* 10.3, pp. 641–648. DOI: [10.1111/tops.12282](https://doi.org/10.1111/tops.12282).
- Bloomfield, Brian P. (1987). "The culture of artificial intelligence". In: *The question of artificial intelligence*. Routledge, pp. 59–105.
- Boden, Margaret A. (2006). *Mind As Machine: A History of Cognitive Science Two-Volume Set*. Oxford University Press, USA.
- Bogost, Ian (2025). "The AI Mirage". In: *The Atlantic*. URL: <https://www.theatlantic.com/technology/archive/2025/07/why-are-computers-still-so-dumb/683524/>.
- Bond, Chanea (2025). "What Happened the Year I Banned AI". In: *Edutopia*. URL: <https://www.edutopia.org/article/banning-ai-tools-class/>.
- Bouchard, Jeremie (2024). "ChatGPT and the separation between knowledge and knower". In: *Education and Information Technologies* 30, pp. 10091–10110. DOI: [10.1007/s10639-024-13249-y](https://doi.org/10.1007/s10639-024-13249-y).
- Boztas, Senay (2025). "Smartphone bans in Dutch schools have improved learning, study finds". In: *The Guardian*. URL: <https://www.theguardian.com/world/2025/jul/04/dutch-school-s-phone-ban-has-improved-learning-study-finds>.
- Brainard, Lindsay (2025). "Step Away From the Chatbot: A Letter to a Student About AI and Creativity". URL: <https://philpapers.org/rec/BRASAF-4>.
- Brennan, Kate, Amba Kak, and Sarah Myers West (2025). *Artificial Power: AI Now 2025 Landscape*. AI Now Institute. URL: <https://ainowinstitute.org/2025-landscape>.
- Brown, Laila M (2023). "Gender, race, and the invisible labor of artificial intelligence". In: *Handbook of Critical Studies of Artificial Intelligence*. Ed. by Simon Lindgren. Edward Elgar Publishing, pp. 573–583. DOI: [10.4337/9781803928562.00059](https://doi.org/10.4337/9781803928562.00059).
- Brown, Neil (1996). "Public inquiry into the shootings at Dunblane Primary School". In: *Disaster Victim Identification: Experience and Practice*, p. 143. URL: <https://www.gov.uk/government/publications/public-inquiry-into-the-shootings-at-dunblane-primary-school>.
- Budzyń, Krzysztof et al. (2025). "Endoscopist deskilling risk after exposure to artificial intelligence in colonoscopy: a multicentre, observational study". In: *The Lancet Gastroenterology & Hepatology*. DOI: [10.1016/S2468-1253\(25\)00133-5](https://doi.org/10.1016/S2468-1253(25)00133-5).
- Butterick, Matthew (2025). *GitHub Copilot Litigation*. URL: <https://githubcopilotlitigation.com/> (visited on 08/24/2025).
- Callon, Michel (1990). "Techno-economic networks and irreversibility". In: *The sociological review* 38.1_suppl, pp. 132–161. DOI: [10.1111/j.1467-954X.1990.tb03351.x](https://doi.org/10.1111/j.1467-954X.1990.tb03351.x).
- Campolo, Alexnder and Kate Crawford (2020). "Enchanted Determinism: Power without Responsibility in Artificial Intelligence". In: *Engaging Science, Technology, and Society* 6, pp. 1–19. DOI: [10.17351/ests2020.277](https://doi.org/10.17351/ests2020.277).
- CDH Team and Carrie Ruddick (2025). *The Incompatibilities Between Generative AI and Art: Q&A with Ted Chiang*. URL: <https://cdh.princeton.edu/blog/2025/08/12/ted-chiang/>.
- Chawla, Chhavi et al. (2024). "Agentic AI: The building blocks of sophisticated AI business applications". In: *Journal of AI, Robotics & Workplace Automation* 3.3, pp. 1–15. DOI: [10.69554/XEHZ1946](https://doi.org/10.69554/XEHZ1946).
- Chirimuuta, Mazviita (2018). "Marr, Mayr, and MR: What functionalism should now be about". In: *Philosophical Psychology* 31.3, pp. 403–418. DOI: [10.1080/09515089.2017.1381679](https://doi.org/10.1080/09515089.2017.1381679).

- Chow, Andrew R (2025). “The AI Industry is Funding A Massive AI Training Initiative for Teachers”. In: *TIME*. URL: <https://time.com/7301335/ai-education-microsoft-openai-anthropix/>.
- Cole, Samantha (2025). *AI Therapy Bots Are Conducting ‘Illegal Behavior,’ Digital Rights Organizations Say*. URL: <https://www.404media.co/ai-therapy-bots-meta-character-ai-ftc-complaint/> (visited on 08/24/2025).
- Colón Vargas, Nelson (2025). “Exploiting the margin: How capitalism fuels AI at the expense of minoritized groups”. In: *AI and Ethics* 5.2, pp. 1871–1876. DOI: [10.1007/s43681-024-00502-w](https://doi.org/10.1007/s43681-024-00502-w).
- Crane, Tim (2021). “The AI ethics hoax”. In: *The Institute of Art and Ideas*. URL: <https://iai.tv/articles/the-ai-ethics-hoax-auid-1762>.
- Crawford, Kate (2021). *The atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. Yale University Press.
- Crawley, Kim (2025). *I’m a professor literally risking my own life to fight the menace of Gen AI destroying education. You can help me live and fight...* URL: https://medium.com/@kim_crawley/im-a-professor-literally-risking-my-own-life-to-fight-the-menace-of-gen-ai-destroying-education-6b3c596f55ce.
- Dennis, Rutledge M (1995). “Social Darwinism, scientific racism, and the metaphysics of race”. In: *Journal of Negro Education* 64.3, pp. 243–252. DOI: [10.2307/2967206](https://doi.org/10.2307/2967206).
- Dentella, Vittoria, Fritz Günther, and Evelina Leivada (2023). “Systematic testing of three Language Models reveals low language accuracy, absence of response stability, and a yes-response bias”. In: *Proceedings of the National Academy of Sciences* 120.51, e2309583120. DOI: [10.1073/pnas.2309583120](https://doi.org/10.1073/pnas.2309583120).
- (2024). “Language in Vivo vs. in Silico: Size Matters but Larger Language Models Still Do Not Comprehend Language on a Par with Humans”. In: *arXiv preprint arXiv:2404.14883*. DOI: [10.48550/arXiv.2404.14883](https://doi.org/10.48550/arXiv.2404.14883).
- Devlin, Jacob et al. (2019). *BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding*. DOI: [10.48550/arXiv.1810.04805](https://doi.org/10.48550/arXiv.1810.04805). arXiv: [1810.04805 \[cs.CL\]](https://arxiv.org/abs/1810.04805).
- DeVrio, Alicia et al. (2025). “A Taxonomy of Linguistic Expressions That Contribute To Anthropomorphism of Language Technologies”. In: *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. CHI ’25. ACM, pp. 1–18. DOI: [10.1145/3706598.3714038](https://doi.org/10.1145/3706598.3714038).
- Dhaliwal, Ranjodh Singh, Théo LePage-Richer, and Lucy Suchman (2024). *Neural Networks*. meson press. URL: <https://meson.press/wp-content/uploads/2024/07/978-3-95796-083-2-Neural-Networks.pdf>.
- Dillon, Sarah (2020). “The Eliza effect and its dangers: from demystification to gender critique”. In: *Journal for Cultural Research* 24.1, pp. 1–15. DOI: [10.1080/14797585.2020.1754642](https://doi.org/10.1080/14797585.2020.1754642).
- Dingemanse, Mark (2024). *Generative AI and Research Integrity*. DOI: [10.31219/osf.io/2c48n](https://doi.org/10.31219/osf.io/2c48n).
- (2025). “Lumo: the least open ‘open’ model we’ve seen”. In: *European Open Source AI Index*. URL: <https://osai-index.eu/news/lumo-proton-least-open>.
- Doetie Talsma Jara Majerus, Antonia Leise (2025). “They are not using AI: ‘Avoiding it is actually really easy’”. In: *Vox magazine*. URL: <https://www.voxweb.nl/en/they-are-not-using-ai-avoiding-it-is-actually-really-easy>.
- Drapkin, Aaron (2023). *AI Gone Wrong: The Errors, Mistakes, and Hallucinations of AI (2025)*. URL: <https://tech.co/news/list-ai-failures-mistakes-errors>.
- Dretske, Fred (1994). “If You Can’t Make One, You Don’t Know How It Works”. In: *Midwest Studies In Philosophy* 19.1, pp. 468–482. DOI: [10.1111/j.1475-4975.1994.tb00299.x](https://doi.org/10.1111/j.1475-4975.1994.tb00299.x). eprint: <https://onlinelibrary.wiley.com/doi/pdf/10.1111/j.1475-4975.1994.tb00299.x>.
- Dreyfus, Hubert L. (1965). *Alchemy and artificial intelligence*. URL: <https://www.rand.org/content/dam/rand/pubs/papers/2006/P3244.pdf>.

- Drimmer, Sonja and Christopher J. Nygren (2025). “Reflections on AI in the Classroom: How We are Not Using AI in the Classroom”. In: *The Newsletter of the International Center of Medieval Art*, pp. 25–28. URL: https://static1.squarespace.com/static/55577d2fe4b02de6a6ea49cd/t/67dfef8d9ff3a5472a6d719d/1742728078061/Drimmer_Nygren_Not_Using_AI.pdf.
- Duarte, Tania, Nicholas Barrow, M. Bakayeva, et al. (2024). “Editorial: The ethical implications of AI hype”. In: *AI Ethics* 4, pp. 649–651. DOI: [10.1007/s43681-024-00539-x](https://doi.org/10.1007/s43681-024-00539-x).
- Dudley, W.H. (1939). “The vocoder”. In: *Bell. Labs. Rec.* 18, p. 122.
- Dupré, Maggie Harrison (2025). *People Are Becoming Obsessed with ChatGPT and Spiraling Into Severe Delusions*. URL: <https://futurism.com/chatgpt-mental-health-crises>.
- Eaton, Kit (2025). “Not So Fast! AI Isn’t Saving Companies Much Time or Money, Study Says”. In: *Inc.* URL: <https://www.inc.com/kit-eaton/not-so-fast-ai-isnt-saving-companies-much-time-or-money-study-says/91183471>.
- Efron, Bradley (1975). “The efficiency of logistic regression compared to normal discriminant analysis”. In: *Journal of the American Statistical Association* 70.352, pp. 892–898. DOI: [10.2307/2285453](https://doi.org/10.2307/2285453).
- Egan, Frances (1999). “In defence of narrow mindedness”. In: *Mind & Language* 14.2, pp. 177–194. DOI: [10.1111/1468-0017.00108](https://doi.org/10.1111/1468-0017.00108).
- (2017). “Function-Theoretic Explanation and the Search for Neural Mechanisms”. In: *Explanation and Integration in Mind and Brain Science*. Ed. by David Michael Kaplan. Oxford University Press, pp. 145–163.
- (2018). “The nature and function of content in computational models”. In: *The Routledge handbook of the computational mind*. Ed. by Mark Sprevak and Matteo Colombo. Routledge, pp. 247–258.
- Eichenberger, Audrey, Stephen Thielke, and Adam Van Buskirk (2025). “A Case of Bromism Influenced by Use of Artificial Intelligence”. In: *Annals of Internal Medicine: Clinical Cases* 4.8, e241260. DOI: [10.7326/aimcc.2024.1260](https://doi.org/10.7326/aimcc.2024.1260).
- Eindhoven University of Technology (2025). *Our Core Values*. URL: <https://www.tue.nl/en/our-university/about-the-university/core-values> (visited on 08/21/2025).
- Elder, Alexis (2022). “Siri, Stereotypes, and the Mechanics of Sexism”. In: *Feminist Philosophy Quarterly* 8.3/4. DOI: [10.5206/fpq/2022.3/4.14294](https://doi.org/10.5206/fpq/2022.3/4.14294).
- Erscoi, Lelia, Annelies V. Kleinherenbrink, and Olivia Guest (2023). *Pygmalion Displacement: When Humanising AI Dehumanises Women*. DOI: [10.31235/osf.io/jqxb6](https://doi.org/10.31235/osf.io/jqxb6).
- Evans, Claire L. (2020). *Broad band: the untold story of the women who made the Internet*. Penguin.
- Falletti, Elena (2024). “Surfing reality, hype, and propaganda: an empirical comparative analysis on predictive software in criminal justice”. In: *AI Ethics* 4, pp. 819–831. DOI: [10.1007/s43681-024-00447-0](https://doi.org/10.1007/s43681-024-00447-0).
- Figdor, Carrie (2010). “Neuroscience and the multiple realization of cognitive functions”. In: *Philosophy of Science* 77.3, pp. 419–456. DOI: [10.1086/652964](https://doi.org/10.1086/652964).
- Fisher, Ronald A. (1958). “Cigarettes, cancer, and statistics”. In: *The Centennial Review of Arts & Science* 2, pp. 151–166.
- Flores, Nelson and Jonathan Rosa (2015). “Undoing appropriateness: Raciolinguistic ideologies and language diversity in education”. In: *Harvard educational review* 85.2, pp. 149–171. DOI: [10.17763/0017-8055.85.2.149](https://doi.org/10.17763/0017-8055.85.2.149).
- FNV et al. (2025). *Onderhandelingsresultaat CAO Nederlandse Universiteiten 1 juli 2025 t/m 30 juni 2026*. URL: <https://www.universiteitenvannederland.nl/files/publications/Onderhandelingsresultaat%20cao-NU%20definitief.pdf>.
- Fodor, Jerry A. (2000). *The mind doesn’t work that way: The scope and limits of computational psychology*. MIT press.

- Forbes, Samuel H. and Olivia Guest (2025). “To improve literacy, improve equality in education, not large language models”. In: *Cognitive Science* 49.4, e70058. DOI: [10.1111/cogs.70058](https://doi.org/10.1111/cogs.70058).
- Fuentes, Agustín (2024). “Scientists as political advocates”. In: *Science* 386.6724. DOI: [10.1126/science.adt7194](https://doi.org/10.1126/science.adt7194).
- Garcia, Patricia et al. (2022). “No! Re-imagining data practices through the lens of critical refusal”. In: *Proceedings of the ACM on Human-Computer Interaction* 6.CSCW2, pp. 1–20. DOI: [10.1145/3557997](https://doi.org/10.1145/3557997).
- Gebru, Timnit and Émile P. Torres (2024). “The TESCREAL bundle: Eugenics and the promise of utopia through artificial general intelligence”. In: *First Monday* 29.4. DOI: [10.5210/fm.v29i4.13636](https://doi.org/10.5210/fm.v29i4.13636).
- Gent, Edd (2024). *The tech industry can’t agree on what open source AI means. That’s a problem*. URL: <https://www.technologyreview.com/2024/03/25/1090111/tech-industry-open-source-ai-definition-problem/>.
- Gerdes, Anne (2022). “The tech industry hijacking of the AI ethics research agenda and why we should reclaim it”. In: *Discover Artificial Intelligence* 2.25. DOI: [10.1007/s44163-022-00043-3](https://doi.org/10.1007/s44163-022-00043-3).
- Gill, Rosalind and Ngaire Donaghue (2016). “Resilience, apps and reluctant individualism: Technologies of self in the neoliberal academy”. In: *Women’s Studies International Forum*. Vol. 54. Elsevier, pp. 91–99. DOI: [10.1016/j.wsif.2015.06.016](https://doi.org/10.1016/j.wsif.2015.06.016).
- Goel, Namanyay (2025). *AI is Creating a Generation of Illiterate Programmers*. URL: https://nmn.gl/blog/ai-illiterate-programmers?utm_source=changelog-news.
- Goetze, Trystan S. (2024). “AI Art is Theft: Labour, Extraction, and Exploitation: Or, On the Dangers of Stochastic Pollocks”. In: *The 2024 ACM Conference on Fairness Accountability and Transparency*. FAccT ’24. ACM, pp. 186–196. DOI: [10.1145/3630106.3658898](https://doi.org/10.1145/3630106.3658898).
- Gold, Bernard (1990). “A history of Vocoder research at Lincoln Laboratory”. In: *The Lincoln Laboratory Journal* 3.2, pp. 163–202. URL: https://archive.ll.mit.edu/publications/journal/pdf/vol03_no2/3.2.1.vocoder.pdf.
- Goudarzi, Sara (2025). “AI for good, with caveats: How a keynote speaker was censored during an international artificial intelligence summit”. In: *Bulletin of the Atomic Scientists*. URL: <https://thebulletin.org/2025/07/ai-for-good-with-caveats-how-a-keynote-speaker-was-censored-during-an-international-artificial-intelligence-summit>.
- Gould, Stephen Jay (1981). *The Mismeasure of Man*. W. W. Norton & Company.
- Gray, Mary L and Siddharth Suri (2019). *Ghost work: How to stop Silicon Valley from building a new global underclass*. Harper Business.
- Guest, Olivia (2024). “What makes a good theory, and how do we make a theory good?” In: *Computational Brain & Behavior* 7.4, pp. 508–522. DOI: [10.1007/s42113-023-00193-2](https://doi.org/10.1007/s42113-023-00193-2).
- (2025). *What Does ‘Human-Centred AI’ Mean?* DOI: [10.48550/arXiv.2507.19960](https://doi.org/10.48550/arXiv.2507.19960). arXiv: [2507.19960 \[cs.AI\]](https://arxiv.org/abs/2507.19960).
- Guest, Olivia and Samuel H. Forbes (2024). “Teaching coding inclusively: if this, then what?” In: *Tijdschrift voor Genderstudies* 27.2/3, pp. 196–217. DOI: [10.5117/TVGN2024.2-3.007.GUES](https://doi.org/10.5117/TVGN2024.2-3.007.GUES).
- Guest, Olivia and Andrea E. Martin (2023). “On logical inference over brains, behaviour, and artificial neural networks”. In: *Computational Brain & Behavior* 6.2, pp. 213–227. DOI: [10.1007/s42113-022-00166-x](https://doi.org/10.1007/s42113-022-00166-x).
- (2025a). *A Metatheory of Classical and Modern Connectionism*. DOI: [10.31234/osf.io/eaf2z_v2](https://doi.org/10.31234/osf.io/eaf2z_v2).
- (2025b). *Are Neurocognitive Representations ‘Small Cakes’?* URL: <https://philsci-archive.pitt.edu/24834/>.

- Guest, Olivia, Natalia Scharfenberg, and Iris van Rooij (2025). *Modern Alchemy: Neurocognitive Reverse Engineering*. URL: <https://philsci-archive.pitt.edu/25289/>.
- Guest, Olivia, Iris van Rooij, et al. (2025). *Open Letter: Stop the Uncritical Adoption of AI Technologies in Academia*. URL: <https://openletter.earth/open-letter-stop-the-uncritical-adoption-of-ai-technologies-in-academia-b65bba1e>.
- Hagopian, Jesse (2025). “Regulating AI Isn’t Enough. Let’s Dismantle the Logic That Put It in Schools”. In: *Truthout*. URL: <https://truthout.org/articles/regulating-ai-isnt-enough-lets-dismantle-the-logic-that-put-it-in-schools/>.
- Haigh, Thomas (2023). “There Was No ‘First AI Winter’”. In: *Communications of the ACM* 66.12, pp. 35–39. DOI: [10.1145/3625833](https://doi.org/10.1145/3625833).
- Hajer, Maarten A and Jeroen Oomen (2025). *Captured Futures: Rethinking the Drama of Environmental Politics*. Oxford University Press.
- Hamilton, Sheryl N. (1998). “Incomplete Determinism: A Discourse Analysis of Cybernetic Futurology in Early Cyberculture”. In: *Journal of Communication Inquiry* 22.2, pp. 177–204. DOI: [10.1177/0196859998022002005](https://doi.org/10.1177/0196859998022002005).
- Hao, Karen (2025). *Empire of AI: Dreams and Nightmares in Sam Altman’s OpenAI*. Penguin Press.
- Hardcastle, Valerie Gray (1996). *How to build a theory in cognitive science*. State University of New York Press.
- Hay, John C., Ben E. Lynch, and David R. Smith (1960). “Mark I perceptron operators’ manual”. In: *Cornell Aeronautical Lab, Buffalo, NY VG-II96-G-5*. URL: <https://apps.dtic.mil/sti/tr/pdf/AD0236965.pdf>.
- Heeg, Dagmar Mercedes and Lucy Avraamidou (2024). “Young children’s understanding of AI”. In: *Education and Information Technologies* 30, pp. 1–24. DOI: [10.1007/s10639-024-13169-x](https://doi.org/10.1007/s10639-024-13169-x).
- Heffernan, Teresa, ed. (2019). *Cyborg Futures: Cross-disciplinary Perspectives On Artificial Intelligence And Robotics*. Social and Cultural Studies of Robots and AI. Palgrave Macmillan.
- (2024). “The imitation game, the “child machine,” and the fathers of AI”. In: *AI & SOCIETY* 39.1, pp. 353–357. DOI: [10.1007/s00146-022-01512-0](https://doi.org/10.1007/s00146-022-01512-0).
- (2025). *An open letter to Evan Solomon, Minister of Artificial Intelligence*. URL: <https://www.halifaxexaminer.ca/commentary/an-open-letter-to-evan-solomon-minister-of-artificial-intelligence/>.
- Helfrich, Gina (2024). “The harms of terminology: why we should reject so-called “frontier AI””. In: *AI and Ethics* 4.3, pp. 699–705. DOI: [10.1007/s43681-024-00438-1](https://doi.org/10.1007/s43681-024-00438-1).
- Hicks, Mar (2017). *Programmed Inequality: How Britain Discarded Women Technologists and Lost Its Edge in Computing*. Hardcover. The MIT Press.
- Hicks, Michael Townsen, James Humphries, and Joe Slater (2024). “ChatGPT is bullshit”. In: *Ethics and Information Technology* 26.2, pp. 1–10. DOI: [10.1007/s10676-024-09775-5](https://doi.org/10.1007/s10676-024-09775-5).
- Hiemstra, Djoerd (2023). *How NLP became a tool of villains*. URL: <https://djoerdhiemstra.com/2023/how-nlp-became-a-tool-of-villains/>.
- Hill, Kashmir (2025). “They Asked ChatGPT Questions. The Answers Sent Them Spiraling”. In: *The New York Times*. URL: <https://www.nytimes.com/2025/06/13/technology/chatgpt-ai-chatbots-conspiracies.html>.
- Hoko, J. Aaron (1986). “What Is the Scientific Value of Comparing Automated and Human Instruction?” In: *Educational Technology* 26.2, pp. 16–19. URL: <http://www.jstor.org/stable/44425163> (visited on 08/29/2025).
- Horgan, John (2020). “Will artificial intelligence ever live up to its hype?” In: *Scientific American* 4. URL: <https://www.scientificamerican.com/article/will-artificial-intelligence-ever-live-up-to-its-hype/>.

- Hosseini, Soodeh and Hossein Seilani (2025). “The role of agentic AI in shaping a smart future: A systematic review”. In: *Array*, p. 100399. DOI: [10.1016/j.array.2025.100399](https://doi.org/10.1016/j.array.2025.100399).
- Hsu, Hua (2025). “What Happens After A.I. Destroys College Writing?” In: *The New Yorker*. URL: <https://www.newyorker.com/magazine/2025/07/07/the-end-of-the-english-paper>.
- Huntington, Clare (2025). “AI Companions and the Lessons of Family Law”. In: *Columbia Public Law Research Paper No. 5283581*. DOI: [10.2139/ssrn.5283581](https://doi.org/10.2139/ssrn.5283581).
- HWR Berlin (2025). *Mission and values*. URL: <https://www.hwr-berlin.de/en/hwr-berlin/about-us/mission-and-values> (visited on 08/23/2025).
- Ihara, Aya S et al. (2021). “Advantage of handwriting over typing on learning words: Evidence from an N400 event-related potential index”. In: *Frontiers in human neuroscience* 15, p. 679191. DOI: [10.3389/fnhum.2021.679191](https://doi.org/10.3389/fnhum.2021.679191).
- Isayas, Meelina (2025a). “Relation generation”. In: *Ada Lovelace Institute*. URL: <https://www.adalovelaceinstitute.org/blog/relation-generation/>.
- (2025b). *Relation generation — How AI assistants are impacting the way young people socialise*. URL: <https://www.adalovelaceinstitute.org/blog/relation-generation/>.
- Jackson, Sarah (2024). *Sam Altman Explains OpenAI’s Shift to Closed AI Models*. URL: <https://www.businessinsider.com/sam-altman-why-openai-closed-source-ai-models-2024-11?international=true&r=US&IR=T>.
- Jebara, Tony (2004). “Generative versus discriminative learning”. In: *Machine learning: Discriminative and generative*. Springer, pp. 17–60. DOI: [10.1007/978-1-4419-9011-2_2](https://doi.org/10.1007/978-1-4419-9011-2_2).
- Jj (2025). *The Copilot Delusion*. URL: <https://deplet.ing/the-copilot-delusion/>.
- Jucan, Ioana B. (2023). *Malicious Deceivers: Thinking Machines and Performative Objects*. Sensing Media: Aesthetics, Philosophy, and Cultures of Media. Stanford University Press.
- Kalluri, Pratyusha Ria et al. (2025). “Computer-vision research powers surveillance technology”. In: *Nature* 643, pp. 73–79. DOI: [10.1038/s41586-025-08972-6](https://doi.org/10.1038/s41586-025-08972-6).
- Kaltheuner, Frederike (2021). *Fake AI*. Meatspace Press. DOI: [10.58704/kcha-1h20](https://doi.org/10.58704/kcha-1h20).
- Kambhampati, Subbarao et al. (2025). “Stop Anthropomorphizing Intermediate Tokens as Reasoning/Thinking Traces!” In: *arXiv preprint arXiv:2504.09762*.
- Kaplan, S (2024). “Dr. Jodi Halpern on why AI isn’t a magic bullet for mental health”. In: *Berkeley Public Health*. URL: <https://publichealth.berkeley.edu/articles/spotlight/research/why-ai-isnt-a-magic-bullet-for-mental-health>.
- Kapoor, Sayash et al. (2024). “Position: On the Societal Impact of Open Foundation Models”. In: *Proceedings of the 41st International Conference on Machine Learning*. Ed. by Ruslan Salakhutdinov et al. Vol. 235. Proceedings of Machine Learning Research. PMLR, pp. 23082–23104. URL: <https://proceedings.mlr.press/v235/kapoor24a.html>.
- Kelly, Clare, Katja Brusch, and Caitríona Leahy (2025). “Opinion: We are lecturers in Trinity College Dublin. It is our responsibility to resist AI”. In: *The Irish Times*. URL: https://www.irishtimes.com/opinion/2025/09/04/opinion-we-are-lecturers-in-trinity-college-we-see-it-as-our-responsibility-to-resist-ai/?utm_source=dlvr.it&utm_medium=bluesky.
- Kinzelbach, Katrin et al. (2025). “Academic Freedom Index 2025 Update”. In: *FAU Erlangen-Nürnberg and V-Dem Institute*. DOI: [10.25593/open-fau-1637](https://doi.org/10.25593/open-fau-1637).
- Klee, Miles (2025). “‘ChatGPT Psychosis’: How One Man Escaped”. In: *Rolling Stone*. URL: <https://www.rollingstone.com/culture/culture-features/chatgpt-ai-philosophical-psychosis-1235404568/>.
- KNAW (2021). *Academische vrijheid in Nederland – een begripsanalyse en richtsnoer*. URL: <https://pure.knw.nl/portal/en/publications/academische-vrijheid-in-nederland-ee-n-begripsanalyse-en-richtsnoer>.

- (2025). *Academic freedom in the Netherlands: response to current threats*. URL: <https://www.knaw.nl/en/publications/academic-freedom-netherlands-response-current-threats>.
- KNAW et al. (2018). *Netherlands Code of Conduct for Research Integrity*. DOI: 10.17026/dans-2cj-nvwu.
- Knoester, Linda et al. (2025). *Academic Collaborations and Public Health: Lessons from Dutch Universities’ Tobacco Industry Partnerships for Fossil Fuel Ties*. DOI: 10.5281/zenodo.15274865.
- Kohnstamm Instituut (2025). URL: <https://open.overheid.nl/documenten/54c01e11-5a20-4779-9243-f4ed5fda1c9f/file>.
- Kosmyna, Nataliya et al. (2025). *Your Brain on ChatGPT: Accumulation of Cognitive Debt when Using an AI Assistant for Essay Writing Task*. DOI: 10.48550/arXiv.2506.08872. arXiv: 2506.08872 [cs.AI].
- Kristan, William B and Paul Katz (2006). “Form and function in systems neuroscience”. In: *Current biology* 16.19, R828–R831. DOI: 10.1016/j.cub.2006.08.079.
- Krizhevsky, Alex, Ilya Sutskever, and Geoffrey E Hinton (2017). “ImageNet classification with deep convolutional neural networks”. In: *Communications of the ACM* 60.6, pp. 84–90. DOI: 10.1145/3065386.
- Kwisthout, Johan (2024). “Subjective and temporal quality-of-life information in Bayesian networks”. In: *Joint International Scientific Conferences on AI and Machine Learning*. Pre-proceedings of BNAIC/BeNeLearn. URL: <https://bnaic2024.sites.uu.nl/program/accepted-submissions/>.
- Kwisthout, Johan and Silja Renooij (2025). “Inverse Marginalisation for Safely Expanding Bayesian Networks”. In: *Proceedings of ECSQARU 2025*.
- Lambeets, Ken (2025). “Is Radboud University too dependent on Big Tech?” In: *Vox magazine*. URL: <https://www.voxweb.nl/nieuws/is-radboud-university-too-dependent-on-big-tech>.
- Lambert, Jonathan (2024). “As schools reconsider cursive, research homes in on handwriting’s brain benefits”. In: *NPR*. URL: <https://www.npr.org/sections/health-shots/2024/05/11/1250529661/handwriting-cursive-typing-schools-learning-brain>.
- Landymore, Frank (2025). “OpenAI Is Giving Exactly the Same Copy-Pasted Response Every Time Time ChatGPT Is Linked to a Mental Health Crisis”. In: *Futurism*. URL: <https://futurism.com/openai-response-chatgpt-mental-health>.
- Larousserie, David (2025). “How AI is shaking up scientific publishing”. In: *Le Monde*. URL: https://www.lemonde.fr/en/science/article/2025/07/12/how-ai-is-shaking-up-scientific-publishing_6743308_10.html#.
- Law, Harry (2024). “Computer vision: AI imaginaries and the Massachusetts Institute of Technology”. In: *AI and Ethics* 4.3, pp. 657–663. DOI: 10.1007/s43681-023-00389-z.
- Lefkowitz, M. (2019). “Professor’s perceptron paved the way for AI—60 years too soon”. In: *Cornell Chronicle* 25. URL: <https://news.cornell.edu/stories/2019/09/professors-perceptron-paved-way-ai-60-years-too-soon>.
- Leivada, Evelina, Vittoria Dentella, and Fritz Günther (2024). “Evaluating the language abilities of Large Language Models vs. humans: Three caveats”. In: *Biolinguistics* 18, pp. 1–12. DOI: 10.5964/bioling.14391.
- Leivada, Evelina, Fritz Günther, and Vittoria Dentella (2024). “Reply to Hu et al.: Applying different evaluation standards to humans vs. Large Language Models overestimates AI performance”. In: *Proceedings of the National Academy of Sciences* 121.36, e2406752121. DOI: 10.1073/pnas.2406752121.
- Levine, Matt (2024). “Who Owns OpenAI?” In: *Bloomberg*. URL: <https://www.bloomberg.com/opinion/articles/2024-10-21/who-owns-openai/>.

- Libell, Monica (2014). "Seeing Animals. Anthropomorphism between Fact and Function." In: *Exploring the Animal Turn. Human-Animal Relations in Science, Society and Culture*. Ed. by Erika Andersson-Cederholm et al. Pufendorfinstitutet, Lunds universitet, pp. 141–153.
- Liesenfeld, Andreas and Mark Dingemanse (2024). "Rethinking open source generative AI: openwashing and the EU AI Act". In: *Proceedings of the 2024 ACM Conference on Fairness, Accountability, and Transparency*. FAccT '24. Rio de Janeiro, Brazil: Association for Computing Machinery, pp. 1774–1787. DOI: [10.1145/3630106.3659005](https://doi.org/10.1145/3630106.3659005).
- Liesenfeld, Andreas, Alianda Lopez, and Mark Dingemanse (2023). "Opening up ChatGPT: Tracking openness, transparency, and accountability in instruction-tuned text generators". In: *Proceedings of the 5th International Conference on Conversational User Interfaces*. CUI '23. ACM, pp. 1–6. DOI: [10.1145/3571884.3604316](https://doi.org/10.1145/3571884.3604316).
- Lighthill, James et al. (1973). "Artificial intelligence: a paper symposium". In: *Science Research Council, London*. URL: https://www.chilton-computing.org.uk/inf/literature/reports/lighthill_report/p001.htm.
- Lodge, Jason M. et al. (2023). "It's not like a calculator, so what is the relationship between learners and generative artificial intelligence?" In: *Learning: Research and Practice* 9.2, pp. 117–124. DOI: [10.1080/23735082.2023.2261106](https://doi.org/10.1080/23735082.2023.2261106).
- Loe, Molly (2023). *ChatGPT's energy usage is a mystery - here's what we do know*. URL: <https://techhq.com/news/data-center-energy-usage-chatgpt/>.
- Longcamp, Marieke, Marie-Thérèse Zerbato-Poudou, and Jean-Luc Velay (2005). "The influence of writing practice on letter recognition in preschool children: A comparison between handwriting and typing". In: *Acta psychologica* 119.1, pp. 67–79. DOI: [10.1016/j.actpsy.2004.10.019](https://doi.org/10.1016/j.actpsy.2004.10.019).
- Luccioni, Sasha, Yacine Jernite, and Emma Strubell (2024). "Power Hungry Processing: Watts Driving the Cost of AI Deployment?" In: *The 2024 ACM Conference on Fairness Accountability and Transparency*. FAccT '24. ACM, pp. 85–99. DOI: [10.1145/3630106.3658542](https://doi.org/10.1145/3630106.3658542).
- Luitse, Dieuwertje and Wiebke Denkena (2021). "The great Transformer: Examining the role of large language models in the political economy of AI". In: *Big Data & Society* 8.2, p. 20539517211047734. DOI: [10.1177/20539517211047734](https://doi.org/10.1177/20539517211047734).
- Lysen, Flora and Sally Wyatt (2024). "Refusing participation: hesitations about designing responsible patient engagement with artificial intelligence in healthcare". In: *Journal of Responsible Innovation* 11.1, p. 2300161. DOI: [10.1080/23299460.2023.2300161](https://doi.org/10.1080/23299460.2023.2300161).
- MacPartland, James et al. (1971). *tudent participation in high school decisions. A study of students and teachers in fourteen urban high schools. Summary and excerpts*. Johns Hopkins University, Center for the Study of Social Organization of Schools.
- Maffulli, Stefano (2023). *Meta's LLaMa license is not Open Source*. URL: <https://opensource.org/blog/metas-llama-2-license-is-not-open-source>.
- Maris, Jordan (2025). *Meta's LLaMa license is still not Open Source*. URL: <https://opensource.org/blog/metas-llama-license-is-still-not-open-source>.
- Markelius, Alva et al. (2024). "The mechanisms of AI hype and its planetary and social costs". In: *AI and Ethics* 4.3, pp. 727–742. DOI: [10.1007/s43681-024-00461-2](https://doi.org/10.1007/s43681-024-00461-2).
- Markov, Igor L (2024). "Reevaluating Google's reinforcement learning for IC macro placement". In: *Communications of the ACM* 67.11, pp. 60–71. DOI: [10.1145/3676845](https://doi.org/10.1145/3676845).
- Matusz, Pawel J., Anna Abalkina, and Dorothy V. M. Bishop (2025). "The Threat of Paper Mills to Biomedical and Social Science Journals: The Case of the Tanu. pro Paper Mill in Mind, Brain, and Education". In: *Mind, Brain, and Education* 19.2, pp. 90–100. DOI: [10.1111/mbe.12436](https://doi.org/10.1111/mbe.12436).
- Mayor, Adrienne (2018). *Gods and robots: Myths, machines, and ancient dreams of technology*. Princeton University Press.

- McCarthy, John et al. (1955). *A proposal for the Dartmouth summer research project on artificial intelligence*. URL: <http://jmc.stanford.edu/articles/dartmouth/dartmouth.pdf>.
- McCorduck, Pamela (2004). *Machines who think: A personal inquiry into the history and prospects of artificial intelligence*. AK Peters/CRC Press. URL: https://monoskop.org/images/1/1e/McCorduck_Pamela_Machines_Who_Think_2nd_ed.pdf.
- McCulloch, Warren S. and Walter Pitts (1943). “A logical calculus of the ideas immanent in nervous activity”. In: *The bulletin of mathematical biophysics* 5, pp. 115–133.
- McQuillan, Dan (2022). *Resisting AI: An anti-fascist approach to artificial intelligence*. Policy Press.
- (2025). *Decomputing as Resistance*. URL: https://danmcquillan.org/decomputing_as_resistance.html.
- McQuillan, Dan, Juliane Jarke, and Teresa Cerratto Pargman (2024). “We are at an extreme point where we have to go all in on what we really believe education should be about”. In: *Postdigital Science and Education* 6.1, pp. 360–368. DOI: [10.1007/s42438-023-00433-5](https://doi.org/10.1007/s42438-023-00433-5).
- Meelissen, Martina et al. (2023). *Resultaten PISA-2022 in vogelvlucht*. Dutch. Netherlands: University of Twente. DOI: [10.3990/1.9789036559461](https://doi.org/10.3990/1.9789036559461).
- Merchant, Brian (2023). *Blood in the Machine*. Little, Brown and Company.
- Millikan, Ruth Garrett (2021). “Neuroscience and teleosemantics”. In: *Synthese* 199.1, pp. 2457–2465. DOI: [10.1007/s11229-020-02893-9](https://doi.org/10.1007/s11229-020-02893-9).
- Mirowski, Philip (2023). “The evolution of platform science”. In: *Social Research: An International Quarterly* 90.4, pp. 725–755. DOI: [10.1353/sor.2023.a916352](https://doi.org/10.1353/sor.2023.a916352).
- MIT Economics (2025). *Assuring an accurate research record*. URL: <https://economics.mit.edu/news/assuring-accurate-research-record>.
- Mitchell, Tom M. (1997). *Machine learning*. Vol. 1. 9. McGraw-hill New York.
- Monett, Dagmar (2021). *Coming to terms with intelligence in machines*. URL: <https://www.slideshare.net/dmonett/monett-2021-atd>.
- Monett, Dagmar and Gilbert Paquet (2025). “Against the Commodification of Education—if harms then not AI”. In: *Journal of Open, Distance, and Digital Education* 2.1. DOI: [10.25619/wazgw457](https://doi.org/10.25619/wazgw457).
- Montell, Amanda (2021). *Cultish: The language of fanaticism*. HarperCollins.
- Montgomery, B (2024). “Mother says AI chatbot led her son to kill himself in lawsuit against its maker”. In: *The Guardian*. URL: <https://www.theguardian.com/technology/2024/oct/23/character-ai-chatbot-sewell-setzer-death>.
- Moore, John (1997). *A primitivist primer*. Green Anarchy Distro. URL: <https://theanarchistlibrary.org/library/john-moore-a-primitivist-primer>.
- Morris, Elliott G. (2025). *Your Polls on ChatGPT*. URL: <https://www.verasight.io/>.
- Mueller, Pam A and Daniel M Oppenheimer (2014). “The pen is mightier than the keyboard: Advantages of longhand over laptop note taking”. In: *Psychological science* 25.6, pp. 1159–1168. DOI: [10.1177/0956797614524581](https://doi.org/10.1177/0956797614524581).
- Munn, Luke (2023). “The uselessness of AI ethics”. In: *AI and Ethics* 3.3, pp. 869–877. DOI: [10.1007/s43681-022-00209-w](https://doi.org/10.1007/s43681-022-00209-w).
- Narayanan, Arvind and Sayash Kapoor (2024). “AI snake oil: What artificial intelligence can do, what it can’t, and how to tell the difference”. In: *AI Snake Oil*. Princeton University Press.
- Neville, Stephen J. (2025). *When Help Isn’t Fully Human: The Problem of Generative AI in Crisis Support*. URL: <https://just-tech.ssrc.org/articles/the-problem-of-generative-ai-in-crisis-support/>.
- Ng, Andrew and Michael Jordan (2001). “On discriminative vs. generative classifiers: A comparison of logistic regression and naive bayes”. In: *Advances in neural information processing systems* 14. DOI: <https://ai.stanford.edu/~ang/papers/nips01-discriminativegenerative.pdf>.

- Nick Dyer-Witheford Atle Mikkola Kjösen, James Steinhoff (2019). *Inhuman Power: Artificial Intelligence And The Future Of Capitalism*. Digital Barricades. Pluto Press.
- Nolan, Mike (2025). “How ‘open’ is open-source AI?” In: *Ada Lovelace Institute*. URL: <https://www.adalovelaceinstitute.org/blog/how-open-is-open-source-ai/>.
- Norrsgård, Karen (2008). “Human testing, the eugenics movement, and IRBs”. In: *Nature Education* 1.1, p. 170. URL: <https://www.nature.com/scitable/topicpage/human-testing-the-eugenics-movement-and-irbs-724/>.
- O’Neil, Cathy (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. USA: Crown Publishing Group.
- Ochigame, Rodrigo (2019). “How Big Tech Manipulates Academia to Avoid Regulation”. In: *The Intercept*. URL: <https://theintercept.com/2019/12/20/mit-ethical-ai-artificial-intelligence/>.
- Olazaran, Mikel (1996). “A sociological study of the official history of the perceptrons controversy”. In: *Social Studies of Science* 26.3, pp. 611–659. DOI: [10.1177/030631296026003005](https://doi.org/10.1177/030631296026003005).
- Oransky, Ivan and Adam Marcus (2023). “There’s far more scientific fraud than anyone wants to admit”. In: *The Guardian*. URL: <https://www.theguardian.com/commentisfree/2023/aug/09/scientific-misconduct-retraction-watch>.
- Otis, Nicholas G et al. (2025). “Global Evidence on Gender Gaps and Generative AI”. In: *HBS Working Paper Series*. URL: <https://www.hbs.edu/faculty/Pages/item.aspx?num=66548>.
- Pahwa, Nitish (2025). “A.I. in higher education: These students are avoiding it”. In: *Slate*. URL: <https://slate.com/life/2025/07/ai-college-cheating-gemini-chatgpt-students-policy.html>.
- Paris, Britt S., Lindsay Weinberg, and Emma May (2025). *Fighting Weaponized AI in Higher Education*. URL: <https://academeblog.org/2025/07/22/fighting-weaponized-ai-in-higher-education/>.
- Parshley, Lois (2024). “The Hidden Environmental Impact of AI”. In: *Jacobin*. URL: <https://jacobin.com/2024/06/ai-data-center-energy-usage-environment>.
- Pejcha, Camille Sojit (2023). “A man’s AI-powered girlfriend has been named as an accomplice in his murder attempt”. In: *Document*. URL: <https://www.documentjournal.com/2023/07/ai-chatbot-replika-assassination-attempt-queen-elizabeth-ii-jaswant-singh-chail/>.
- Perez, Carlota (2002). “Technological revolutions and financial capital: The dynamics of bubbles and golden ages”. In: *Technological revolutions and financial capital*. Edward Elgar Publishing.
- Perrigo, Billy (2023). “Exclusive: OpenAI used Kenyan workers on less than \$2 per hour to make ChatGPT less toxic”. In: *Time Magazine* 18, p. 2023. URL: <https://time.com/6247678/openai-chatgpt-kenya-workers/>.
- Pfaffenberger, Bryan (1988). “Fetishised objects and humanised nature: Towards an anthropology of technology”. In: *Man*, pp. 236–252. DOI: [10.2307/2802804](https://doi.org/10.2307/2802804).
- Placani, Adriana (2024). “Anthropomorphism in AI: hype and fallacy”. In: *AI and Ethics* 4.3, pp. 691–698. DOI: [10.1007/s43681-024-00419-4](https://doi.org/10.1007/s43681-024-00419-4).
- Polo, Marina (2024). “O pensamento de que a máquina é pensante: a necessária retomada da consciência crítica diante das invenções tecnológicas da humanidade”. In: *Revista Eletrônica Internacional de Economia Política da Informação da Comunicação e da Cultura* 26.2, pp. 70–80. DOI: [10.54786/revistaepic.v26i2.20551](https://doi.org/10.54786/revistaepic.v26i2.20551).
- Popper, Karl Raimund (1972). *Objective Knowledge: An Evolutionary Approach*. New York: Oxford University Press.
- Powell, Betty (1970). “Descartes’ Machines”. In: *Proceedings of the Aristotelian Society*. Vol. 71. JSTOR, pp. 209–222. URL: <https://www.jstor.org/stable/4544809>.

- Previtali, Fabiane Santana and Cílson César Fagiani (2015). "Deskilling and degradation of labour in contemporary capitalism: The continuing relevance of Braverman". In: *Work Organisation, Labour and Globalisation* 9.1, pp. 76–91. DOI: [10.13169/workorglaboglob.9.1.0076](https://doi.org/10.13169/workorglaboglob.9.1.0076).
- Purtill, Corinne (2025). "AIs gave scarily specific self-harm advice to users expressing suicidal intent, researchers find". In: *Yahoo News*. URL: <https://www.yahoo.com/news/articles/ais-gave-scarily-specific-self-100000574.html>.
- Radboud University (2025). *Mission*. URL: <https://www.ru.nl/en/about-us/mission-and-strategy/mission> (visited on 06/18/2025).
- Raether, Till (2025). "Der Feind in meinem Chat". In: *Süddeutsche Zeitung Magazin*. URL: <https://sz-magazin.sueddeutsche.de/leben-und-gesellschaft/ki-ai-chatgpt-freundschaft-kommunikation-widerstand-95330>.
- Raji, Inioluwa Deborah et al. (2022). "The fallacy of AI functionality". In: *Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency*, pp. 959–972. DOI: [10.1145/3531146.3533158](https://doi.org/10.1145/3531146.3533158).
- Rao, C Radhakrishna (1992). "RA Fisher: The founder of modern statistics". In: *Statistical science*, pp. 34–48. DOI: [10.1214/ss/1177011442](https://doi.org/10.1214/ss/1177011442).
- Rascoe, Ayesha (2024). "California is mandating cursive handwriting instruction in elementary schools". In: *NPR*. URL: <https://www.npr.org/2024/01/14/1224674052/california-is-mandating-cursive-handwriting-instruction-in-elementary-schools>.
- Reddy, Ajitha (2007). "The eugenic origins of IQ testing: Implications for post-Atkins litigation". In: *DePaul L. Rev.* 57, p. 667. URL: <https://via.library.depaul.edu/law-review/vol57/iss3/5>.
- Reed, Michelle et al. (2001). *Resistance, Rebellion, Refusal: 7 Tips from Librarians for Navigating "AI"*. URL: <https://www.civicsoftechnology.org/blog/tips-from-librarians-for-navigating-ai>.
- Reuters (2025). *OpenAI's annualized revenue hits \$10 billion, up from \$5.5 billion in December 2024*. URL: <https://www.reuters.com/business/media-telecom/openais-annualized-revenue-hits-10-billion-up-55-billion-december-2024-2025-06-09/>.
- Reynoldson, Miriam, Melanie Dusseau, and Signatories (2025). *An open letter from educators who refuse the call to adopt GenAI in education*. URL: <https://openletter.earth/an-open-letter-from-educators-who-refuse-the-call-to-adopt-genai-in-education-cb4aee75>.
- Rhee, Jennifer (2018). *The Robotic Imaginary*. University of Minnesota Press.
- Ricaurte, Paola (2019). "Data epistemologies, the coloniality of power, and resistance". In: *Television & New Media* 20.4, pp. 350–365. DOI: [10.1177/1527476419831640](https://doi.org/10.1177/1527476419831640).
- Rich, Patricia et al. (2021). "How hard is cognitive science?" In: *Proceedings of the annual meeting of the cognitive science society*. Vol. 43. 43.
- Rijo, Luis (2025). "Microsoft can't protect French data from US government access". In: *PPC Land*. URL: <https://ppc.land/microsoft-cant-protect-french-data-from-us-government-access/>.
- Roose, Kevin (2024). "Can a Chatbot Named Daenerys Targaryen Be Blamed for a Teen's Suicide?" In: *The New York Times*. URL: <https://www.nytimes.com/2024/10/23/technology/characterai-lawsuit-teen-suicide.html>.
- Rosenblatt, Frank (1958). "The perceptron: A probabilistic model for information storage and organization in the brain". In: *Psychological Review* 65.6. DOI: <https://psycnet.apa.org/doi/10.1037/h0042519>.

- Rowe, Niamh (2023). ““It’s destroyed me completely””: Kenyan moderators decry toll of training of AI models”. In: *The Guardian*. URL: <https://www.theguardian.com/technology/2023/aug/02/ai-chatbot-training-human-toll-content-moderator-meta-openai>.
- Saini, Angela (2019). *Superior: The return of race science*. Beacon Press.
- Sano-Franchini, Jennifer, Megan McIntyre, and Maggie Fernandes (2023). *Refusing GenAI in Writing Studies: A Quickstart Guide*. URL: <https://refusinggenai.wordpress.com/>.
- (2024). “Refusing GenAI in Writing Studies: A Quickstart Guide”. In: *Refusing Generative AI in Writing Studies*. URL: <https://refusinggenai.wordpress.com/>.
- Schellenberg, Susanna (2018). *The unity of perception: Content, consciousness, evidence*. Oxford University Press.
- Schipper, Nienke (2025). “Stop met AI in de klas, roepen docenten op. Maar kan dat nog wel?” In: *Trouw*. URL: <https://www.trouw.nl/duurzaamheid-economie/stop-met-ai-in-de-klas-roepen-docenten-op-maar-kan-dat-nog-wel~bffecd1b/>.
- Schlesinger, Ari, Kenton P. O’Hara, and Alex S. Taylor (2018). “Let’s Talk About Race: Identity, Chatbots, and AI”. In: *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. CHI ’18. Montreal QC, Canada: Association for Computing Machinery, pp. 1–14. DOI: [10.1145/3173574.3173889](https://doi.org/10.1145/3173574.3173889).
- Schoene, Annika M and Cansu Canca (2025). ‘For Argument’s Sake, Show Me How to Harm Myself!’: Jailbreaking LLMs in Suicide and Self-Harm Contexts. DOI: [10.48550/arXiv.2507.02990](https://doi.org/10.48550/arXiv.2507.02990). arXiv: [2507.02990 \[cs.CL\]](https://arxiv.org/abs/2507.02990).
- Schroeder, Manfred Robert (1966). “Vocoders: Analysis and synthesis of speech”. In: *Proceedings of the IEEE* 54.5, pp. 720–734. DOI: [10.1109/PROC.1966.4841](https://doi.org/10.1109/PROC.1966.4841).
- Shapiro, Ari, Patrick Jarenwattananon, and Manuela López Restrepo (2022). *In Britain, it took just one school shooting to pass major gun control*. NPR.
- Shojaee, Parshin et al. (2025a). *The Illusion of Thinking: Understanding the Strengths and Limitations of Reasoning Models via the Lens of Problem Complexity*. URL: <https://ml-site.cdn-apple.com/papers/the-illusion-of-thinking.pdf>.
- (2025b). “The illusion of thinking: Understanding the strengths and limitations of reasoning models via the lens of problem complexity”. In: *arXiv preprint arXiv:2506.06941*. DOI: [10.48550/arXiv.2506.06941](https://doi.org/10.48550/arXiv.2506.06941).
- Sloane, Mona, David Danks, and Emanuel Moss (2024). “Tackling AI Hying”. In: *AI and Ethics* 4.3, pp. 669–677. DOI: [10.1007/s43681-024-00481-y](https://doi.org/10.1007/s43681-024-00481-y).
- Smith, Peter and Laura Smith (2024). “This season’s artificial intelligence (AI): is today’s AI really that different from the AI of the past? Some reflections and thoughts”. In: *AI and Ethics* 4.3, pp. 665–668. DOI: [10.1007/s43681-023-00388-0](https://doi.org/10.1007/s43681-023-00388-0).
- Solaiman, Irene (2023). “The gradient of generative AI release: Methods and considerations”. In: *Proceedings of the 2023 ACM conference on fairness, accountability, and transparency*, pp. 111–122. DOI: [10.1145/3593013.3593981](https://doi.org/10.1145/3593013.3593981).
- Soler, Susana Pérez (2025). “ChatGPT dissidents, the students who refuse to use AI: “I couldn’t remember the last time I had written something by myself””. In: *EL PAÍS*. URL: <https://english.elpais.com/technology/2025-08-03/chatgpt-dissidents-the-students-who-refuse-to-use-ai-i-couldnt-remember-the-last-time-i-had-written-something-by-myself.html>.
- Solomonides, Tony and Les Levidow (1985). *Compulsive technology: Computers as culture*. The Free Press.
- Spanton, Rory W. and Olivia Guest (2022). “Measuring trustworthiness or automating physiognomy? A comment on Safra, Chevallier, Grèzes, and Baumard (2020)”. In: *arXiv preprint arXiv:2202.08674*. DOI: [10.48550/arXiv.2202.08674](https://doi.org/10.48550/arXiv.2202.08674).

- Stengers, Isabelle (2018). *Another Science is Possible: A Manifesto for Slow Science*. Trans. by Stephen Muecke. Polity Cambridge, UK.
- Stinson, Catherine (2018). “Explanation and Connectionist Models”. In: *The Routledge Handbook of the Computational Mind*. Ed. by Mark Sprevak and Matteo Colombo. Routledge, pp. 120–133.
- (2020). “From implausible artificial neurons to idealized cognitive models: Rebooting philosophy of artificial intelligence”. In: *Philosophy of Science* 87.4, pp. 590–611. DOI: [10.1086/709730](https://doi.org/10.1086/709730).
- Stolley, Paul D. (1991). “When genius errs: RA Fisher and the lung cancer controversy”. In: *American Journal of Epidemiology* 133.5, pp. 416–425. DOI: [10.1093/oxfordjournals.aje.a115904](https://doi.org/10.1093/oxfordjournals.aje.a115904).
- Strengers, Yolande et al. (2024). “Anthropomorphizing Voice Assistants: A Research Agenda for Human–AI Relationships”. In: *The De Gruyter Handbook of Robots in Society and Culture*. Ed. by Leopoldina Fortunati and Autumn Edwards. Berlin, Boston: De Gruyter, pp. 391–410. DOI: [10.1515/9783110792270-021](https://doi.org/10.1515/9783110792270-021).
- Study.com (2025). *Productive Teaching Tool or Innovative Cheating?* URL: <https://study.com/resources/perceptions-of-chatgpt-in-schools>.
- Suarez, Marcela et al. (2025). *Critical AI Literacy: Beyond hegemonic perspectives on sustainability*. DOI: [10.5281/zenodo.15677840](https://doi.org/10.5281/zenodo.15677840).
- Suchman, Lucy (2019). “Demystifying the intelligent machine”. In: *Cyborg futures: Cross-disciplinary perspectives on artificial intelligence and robotics*. Springer, pp. 35–61.
- Tafari, Daniela (2023). *What’s wrong with “AI ethics”*. URL: <https://www.ercim.eu/publication/ethics2023/slidesDanielaTafari.pdf>.
- (2024a). “GDPR could protect us from the AI Act. That’s why it’s under attack”. In: *Bollettino telematico di filosofia politica*. DOI: [10.5281/zenodo.14002329](https://doi.org/10.5281/zenodo.14002329).
- (2024b). “Omini di burro. Scuole e università al Paese dei Balocchi dell’IA generativa”. In: *Bollettino telematico di filosofia politica*. DOI: [10.5281/zenodo.13923820](https://doi.org/10.5281/zenodo.13923820).
- Tait, Márcia, Alcides dos Reis Peron, and Marcela Suárez (2022). “Terrestrial politics and body-territory: two concepts to make sense of digital colonialism in Latin America”. In: *Tapuya: Latin American Science, Technology and Society* 5.1, pp. 1–16. DOI: [10.1080/25729861.2022.2090485](https://doi.org/10.1080/25729861.2022.2090485).
- Tan, Eli (2025). “Their Water Taps Ran Dry When Meta Built Next Door”. In: *The New York Times*. URL: <https://www.nytimes.com/2025/07/14/technology/meta-data-center-water.html>.
- Taylor, Josh (2025). “AI chatbots are becoming popular alternatives to therapy. But they may worsen mental health crises, experts warn | Artificial intelligence (AI)”. In: *The Guardian*. URL: <https://www.theguardian.com/australia-news/2025/aug/03/ai-chatbot-as-therapy-alternative-mental-health-crises-ntwnfb>.
- Taylor, Simon Michael, Kalervo N Gulson, and Duncan McDuié-Ra (2023). “Artificial intelligence from colonial India: Race, statistics, and facial recognition in the global south”. In: *Science, Technology, & Human Values* 48.3, pp. 663–689. DOI: [10.1177/01622439211060839](https://doi.org/10.1177/01622439211060839).
- TechEquity (2025). *How Californians Feel About AI - Findings From the 2025 AI Compass*. URL: <https://techequity.us/2025/08/19/how-californians-feel-about-ai/>.
- Temple, James (2024). “Sorry, AI won’t “fix” climate change”. In: *MIT Technology Review*. URL: <https://www.technologyreview.com/2024/09/28/1104588/sorry-ai-wont-fix-climate-change>.
- The Luddite, Alejandro (2024). *[UPDATE] Nature’s Folly: A Response to Nature’s “Loneliness and suicide mitigation for students using GPT3-enabled chatbots”*. URL: <https://theluddite.org/post/replika.html>.
- The New York Times (1958). *NEW NAVY DEVICE LEARNS BY DOING; Psychologist Shows Embryo of Computer Designed to Read and Grow Wiser*. URL: <https://www.nytimes.com/1958>

- /07/08/archives/new-navy-device-learns-by-doing-psychologist-shows-embryo-of.html.
- Thorne, Michelle (2009). *Openwashing*. URL: <https://michellethorne.cc/2009/03/openwashing/>.
- Thornhill, John (2025). “Brace for a crash before the golden age of AI”. In: *Financial Times*. URL: <https://www.ft.com/content/a76f238d-5543-4c01-9419-52aaf352dc23>.
- Tiku, Nitasha (2025). “AI companies tap social media tactics to help chatbots hook users”. In: *The Washington Post*. URL: <https://www.washingtonpost.com/technology/2025/05/31/ai-chatbots-user-influence-attention-chatgpt/>.
- Titus, Lisa Miracchi (2024). “Does ChatGPT have semantic understanding? A problem with the statistics-of-occurrence strategy”. In: *Cognitive Systems Research* 83, p. 101174. DOI: [10.1016/j.cogsys.2023.101174](https://doi.org/10.1016/j.cogsys.2023.101174).
- Tully, Stephanie M, Chiara Longoni, and Gil Appel (2025). “Lower artificial intelligence literacy predicts greater AI receptivity”. In: *Journal of Marketing*. DOI: [10.1177/00222429251314491](https://doi.org/10.1177/00222429251314491).
- Turkle, Sherry (1984). *The second self: Computers and the human spirit*. Mit Press.
- Turkle, Sherry et al. (2006). “Relational artifacts with children and elders: the complexities of cyber-companionship”. In: *Connection Science* 18.4, pp. 347–361. DOI: [10.1080/09540090600868912](https://doi.org/10.1080/09540090600868912).
- Twist, Jo (2003). “Chatbot bids to fool humans”. In: *BBC News Online*. URL: <http://news.bbc.co.uk/2/hi/technology/3116780.stm>.
- U.S. Department of Education, National Center for Education Statistics (2024). *Highlights of the 2023 U.S. PIAAC Results Web Report*. URL: https://nces.ed.gov/surveys/piaac/2023/national_results.asp.
- Valinsky, Jordan (2024). *Three Mile Island is reopening and selling its power to Microsoft*. URL: <https://edition.cnn.com/2024/09/20/energy/three-mile-island-microsoft-ai/index.html>.
- Vallor, Shannon (2024). *The AI mirror: How to reclaim our humanity in an age of machine thinking*. Oxford University Press.
- van der Gun, Luc and Olivia Guest (2024). “Artificial Intelligence: Panacea or Non-Intentional De-humanisation?” In: *Journal of Human-Technology Relations* 2. DOI: [10.59490/jhtr.2024.2.7272](https://doi.org/10.59490/jhtr.2024.2.7272).
- van Laarhoven, Joris and Eva van Vugt (2025). “Lever de rechtsstaat niet uit aan ChatGPT: verbied AI aan de juridische faculteit”. In: *NRC*. URL: <https://www.nrc.nl/nieuws/2025/07/02/lever-de-rechtsstaat-niet-uit-aan-chatgpt-verbied-ai-aan-de-juridische-faculteit-a4899056>.
- van Rooij, Iris (2022). *Against automated plagiarism*. DOI: [10.5281/zenodo.15866638](https://doi.org/10.5281/zenodo.15866638).
- (2025). *AI slop and the destruction of knowledge*. DOI: [10.5281/zenodo.16905560](https://doi.org/10.5281/zenodo.16905560).
- van Rooij, Iris and Olivia Guest (2025). *Combining Psychology with Artificial Intelligence: What could possibly go wrong?* DOI: [10.31234/osf.io/aue4m_v1](https://doi.org/10.31234/osf.io/aue4m_v1).
- van Rooij, Iris, Olivia Guest, et al. (2024). “Reclaiming AI as a theoretical tool for cognitive science”. In: *Computational Brain & Behavior* 7.4, pp. 616–636. DOI: [10.1007/s42113-024-00217-5](https://doi.org/10.1007/s42113-024-00217-5).
- van Rossum, Diederik (2025). *Psychosynthese kwesties*. Stichting ter bevordering van Psychosynthese.
- Villalobos, Pablo et al. (2024). *Will we run out of data? Limits of LLM scaling based on human-generated data*. DOI: [10.48550/arXiv.2211.04325](https://doi.org/10.48550/arXiv.2211.04325). arXiv: [2211.04325 \[cs.LG\]](https://arxiv.org/abs/2211.04325).
- Wallace, Richard S. (2009). “The Anatomy of A.L.I.C.E.” In: *Parsing the Turing Test: Philosophical and Methodological Issues in the Quest for the Thinking Computer*. Ed. by Robert Epstein, Gary Roberts, and Grace Beber. Dordrecht: Springer Netherlands, pp. 181–210. DOI: [10.1007/978-1-4020-6710-5_13](https://doi.org/10.1007/978-1-4020-6710-5_13).

- Wang, Angelina, Jamie Morgenstern, and John P. Dickerson (2024). *Large language models that replace human participants can harmfully misportray and flatten identity groups*. DOI: [10.48550/ARXIV.2402.01908](https://doi.org/10.48550/ARXIV.2402.01908).
- Warner, John (2025). *More than words: How to think about Writing in the age of AI*. Basic Books.
- Watters, Audrey (2023). *Teaching machines: The history of personalized learning*. MIT Press.
- Weatherbed, Jess (2025). *China shuts down AI tools during nationwide college exams*. URL: <https://www.theverge.com/news/682737/china-shuts-down-ai-chatbots-exam-season>.
- Wei, Marlynn (2025). “The Emerging Problem of “AI Psychosis””. In: *Psychology Today*. URL: <https://www.psychologytoday.com/us/blog/urban-survival/202507/the-emerging-problem-of-ai-psychosis>.
- Weizenbaum, Joseph (1966). “ELIZA—a computer program for the study of natural language communication between man and machine”. In: *Communications of the ACM* 9.1, pp. 36–45. DOI: [10.1145/365153.365168](https://doi.org/10.1145/365153.365168).
- Wendling, Amy (2002). “Partial Liberations”. In: *International Studies in Philosophy* 34.2, pp. 169–185. DOI: [10.5840/intstudphil200234250](https://doi.org/10.5840/intstudphil200234250).
- Whittaker, Meredith (2021). “The steep cost of capture”. In: *Interactions* 28.6, pp. 50–55. DOI: [10.1145/3488666](https://doi.org/10.1145/3488666).
- Widder, David Gray, Meredith Whittaker, and Sarah Myers West (2024). “Why ‘open’AI systems are actually closed, and why this matters”. In: *Nature* 635.8040, pp. 827–833. DOI: [10.1038/s41586-024-08141-1](https://doi.org/10.1038/s41586-024-08141-1).
- Wilson, Elizabeth A. (2016). *Neural geographies: Feminism and the microstructure of cognition*. Routledge.
- Wood, Nathan G. (2024). “Regulating autonomous and AI-enabled weapon systems: the dangers of hype”. In: *AI Ethics*. DOI: [10.1007/s43681-024-00448-z](https://doi.org/10.1007/s43681-024-00448-z).
- Xiang, Chloe (2023). ““He Would Still Be Here”: Man Dies by Suicide After Talking with AI Chatbot, Widow Says”. In: *VICE*. URL: <https://www.vice.com/en/article/man-dies-by-suicide-after-talking-with-ai-chatbot-widow-says/>.
- Xue, Jing-Hao and D. Michael Titterington (2008). “Comment on “On discriminative vs generative classifiers: A comparison of logistic regression and naive bayes””. In: *Neural processing letters* 28, pp. 169–187. URL: <https://dl.acm.org/doi/10.5555/2980539.2980648>.
- Young Academy (2025). *In Control: A Manifesto for a Responsible Digital Future in Higher Education*. URL: <https://www.dejongeakademie.nl/en/publications/3082247.aspx>.