

What does machine-generated text do to us? Consequences of a careless use of AI tools

Abstract

AI tools are spreading worldwide and are used in many aspects of our lives. History has shown us that technological progress comes with opportunities and dangers. It therefore seems timely to consider the long-term consequences of this latest technological revolution. In this paper, I will discuss in light of empirical findings four potential negative consequences of careless use of generative AI tools that concern issues like the pollution of our textual corpora, homogenization of language and thinking, deskilling with respect to cognitive abilities, and a weakened human influence on many decision-making processes. Furthermore, I shall offer some considerations on how we may prevent the described consequences of this new technology from dominating the development of our societies and how new social practices could help us achieve this. By doing this, I do not advocate that interactions with the new AI tools should be avoided; instead, I aim to contribute to discussions on how we can live a good life with smart machines.

Keywords

LLM, careless use of AI tools, deskilling, homogenization, technological impact assessment

The rise of generative AI (GenAI) technology has evoked intense public debates and scientific discourses, revealing at least two irreconcilable camps: one warning about the consequences of this new technology, and the other embracing technological progress without reservations. My initial enthusiasm concerning the breathtaking possibilities that come with GenAI led me to fine-tune one of the old models (GPT-3) with the corpus of a famous person (*reference & name of the person deleted for anonymous review*). Years later, my skepticism has grown, and I somehow find myself moving towards the camp of the warning people. Still, I think from a scientific perspective, the further development of GenAI is fascinating, but the widespread careless use of those tools is something that worries me. Perhaps humanity will find a way to develop new cultural techniques and thereby avoid the dangers I anticipate. I am open to being proven wrong by future developments.

In this paper, I shall examine our social practices that guide our ascriptions of trustworthiness and, in particular, how these practices nowadays seem to suggest trusting machine-generated text without verifying its contents, despite our knowledge that all GenAI outputs based on a transformer architecture inherit limitations in reliability. Using GenAI tools without verifying their outputs is what I refer to as careless use of AI (Strasser 2025). Careless, because it would be appropriate to verify the outputs of AI tools before accepting their results. Here, I would like to point out that even OpenAI researchers recently published a paper that provided a comprehensive mathematical framework explaining why AI systems must generate plausible but false information even when trained on perfect data and come to the conclusion that “[...] ‘hallucinations’ persist even in state-of-the-art systems and undermine trust” (Kalai et al. 2025).

At the moment, I have the impression that the capacity of smart machines to process much more data than any individual human could ever do seems to persuade us to believe that all the patterns such machines discover are relevant to us. In other words, because GenAI tools can do a lot of things we cannot do, we seem to be prone to neglect their limitations with respect to reliability. Moreover, since we experience that many outputs are truly impressive, this seems to be another factor that makes us more inclined to think of ‘trustworthy’ AI. Because using an AI tool can improve your performance in many tasks without requiring you to undergo an effortful and time-consuming training.

I admit that it is too early to speak of a fully developed cultural technique with respect to the use of GenAI technology. Nevertheless, I shall present some considerations that point to factors contributing to an increase in careless use of GenAI, thereby shaping the way we use GenAI. To this end, I will report on several findings from recent studies that provide initial indications of the possible consequences of

an over-reliance on GenAI tools and argue for the claim that a careless use of GenAI tools can come with far-reaching consequences that we should view critically. This concerns issues like the pollution of textual corpora, homogenization of language and thinking, deskilling with respect to cognitive abilities, and a weakened human influence on many decision-making processes.

After presenting some considerations about factors that in my view speak for the likeliness of an increasing careless use of GenAI tools (*1. What motivates a careless use of GenAI*), I will provide an overview of first empirical studies dealing with the consequences of AI use to argue, in light of empirical findings, for the following four theses: I start with the claim that an ongoing careless use of GenAI tools will contribute to a pollution of our textual corpora, in which we collect our knowledge and with which we decide questions of social and moral relevance and that this has the potential to destroy established trust relations and may contribute to an epistemological crisis (*2. A new kind of pollution*). Then, I will investigate the extent to which an incorporation of GenAI tools can lead to further decreasing in individuality and diversity and argue that through homogenization, we might become more similar, not only in how we express ourselves in written text but also in what we think (*3. Homogenization through the use of AI tools for writing*). Exploring the potential effects that come with cognitive offloading, I argue for the claim that some of our cognitive skills that contribute to our capacities concerning critical thinking, autonomy, moral agency, and other abilities might become objects of deskilling. This suggests that we may become more dependent on AI tools, which is a critical consequence given that GenAI tools have significant limitations in terms of reliability. (*4. Deskilling*). Last but not least, I will demonstrate that the human influence on many decision-making processes could be weakened due to our tendency to follow recommendations of AI tools without revisiting the reasons that speak for certain decisions (*5. Weakened human influence*). In the last part, I will examine how we may prevent the described consequences of this new technology from dominating the development of our societies and how new social practices could help us achieve this. By doing this, I do not aim to advocate in any sense that interactions with new technology like smart GenAI machines would be something we should not do. My aim is rather to contribute to discussions on how we can live a good life with smart machines.

1. What motivates a careless use of Gen AI

One of the widespread AI-narratives is that one will lose the competition in economic and intellectual domains without GenAI. This comes together with the promise that the use of GenAI will make us more efficient and faster (e.g., projects, such as setting up an advertisement campaign, can now be finalized much faster with the use of GenAI). Furthermore, people are fascinated by the fact that GenAI puts them in the position to create output for which they lack the necessary expertise (e.g., coding, translating, writing long text, etc.). Everybody can make the experience that the output of an AI tool can be better than anything one would have been able to do without the tool. This seems to come along with the promise that we may become less dependent on others' expertise. That means instead of asking or hiring a human with the necessary expertise, we are tempted by the option to use one of the many GenAI tools. In addition, we are confronted with the fact that certain progress in science can only be achieved if GenAI tools are involved because only GenAI tools have the capacity to process the huge number of possibilities. One example of this is the success of *AlphaFold* (Jumper et al. 2021). However, one should be aware that the use of *AlphaFold* was accompanied by the ability of the experts to verify the proposed outputs of this algorithm. Nevertheless, the seemingly breathtaking speed of this technological development, with all the success stories, seems to culminate in the feeling that problems with reliability will be fixed sooner or later.

It is not surprising that in view of a world that entails so much complexity, produces so much data and sources of knowledge, we are welcoming a tool that seems to make our lives easier. Our attempts to make sense of the world, to come to decisions that are based on relevant information, to gain new knowledge, and to agree on social and moral norms that can guide our behavior seem to become more

and more difficult. Taking decision-making as an example, we are overwhelmed by the amount of information that might be relevant. We have no chance to process all of them, we cannot filter them to find the relevant ones, and this means we have to decide under uncertainty or more or less blindly trust in the expertise of others, whereas we may find ourselves in the position of having difficulties in evaluating their expertise. Since we increasingly find ourselves in situations where we cannot come to a decision that considers all relevant information, we are tempted by the promises of GenAI tools to outsource some of the effortful tasks that are no longer manageable by us to GenAI tools.

Under the assumption that we believe those tools are helpful, we feel the pressure to use them in order not to suffer from disadvantages in competition with all the other users. For example, using GenAI tools can make it easier and faster for students to achieve good results (Fauzi et al. 2023). Furthermore, we might in general prefer a less effortful way of handling effortful cognitive tasks and use our skills for other things. Cognitive offloading may indeed put us in a position in which we can solve more difficult tasks, as reducing the cognitive load leaves space to engage in other cognitive challenges. Especially in the economic domain, it seems as if using GenAI tools that can overtake formerly human work can lead to immense advantages, as this promises to be not only faster but also significantly cheaper. And one might even suppose that it could also be more skillful, since machines can do things that only humans with expertise can do.

Due to the ongoing distribution of labor, which is leading to a wide range of specializations, we must acknowledge that we lack expertise in many domains. At the same time, we find it challenging to evaluate the expertise of others. Of course, we have developed strategies to recognize expertise; for example, we tend to trust in scientific procedures, in the effects of the education of specialists, and so on. However, in view of the diversity of recommendations that so-called experts deliver, we may become uncertain about those strategies. In this situation, we seem prone to rely on a technology that apparently offers easy access to expertise in many domains. With respect to many problems that concern our way of living, we increasingly feel dependent on the expertise of others, and we might prefer to gain expertise ourselves with the help of AI tools instead of trusting in others. No matter whether we want to make a decision concerning our well-being, health, finances, or politics, we are often confronted with a huge diversity of experts that we might consider. In the absence of a guiding role from a religion or another authority that provides direction on what and who to believe, we may feel overwhelmed. This is where we are prone to see GenAI as a means to manage the vast amount of information that needs evaluation and critical testing before making a decision.

These general observations are intended to explain why I suspect that we will have to deal with a lot of careless use of GenAI tools in the future. In the following sections, I will elaborate on four consequences, namely, the pollution of textual corpora, homogenization of language and thinking, deskilling with respect to cognitive abilities, and a weakened human influence on many decision-making processes, that I anticipate here.

2. A new kind of pollution

To argue that a widespread careless use of GenAI tools will contribute to polluting our textual corpora, from which we collect knowledge and with which we decide questions of social and moral relevance, and how this can destroy established trust relations, I start with a simple example. It is well-known that LLMs tend to hallucinate. The term ‘hallucination’ is used to refer to mistakes in machine-generated texts that are semantically or syntactically plausible but are, in fact, incorrect or nonsensical.¹ Contrary to the optimistic statements made by companies offering LLM-based chatbots, it is scientifically questionable whether future LLMs will be able to refrain from hallucinations; I state that making untenable hypotheses is a feature of the architecture used and not a bug (Banerjee et al. 2025; Kalai et

¹ I consider the choice of this term, which was popularized by Google AI researchers (Agarwal et al. 2018), to be unfortunate and would suggest to rather describe this kind of erroneous output as untenable hypotheses instead.

al. 2025). Of course, it cannot be ruled out that future hybrid systems, based on different architectures, will be able to filter out hallucinations.

Evaluating the performances of current smart machines, we see that LLMs produce hallucinations. For example, it is well-known that LLMs frequently hallucinate references (Chelli et al. 2024; Mugaanyi et al. 2024). Since many scholars nowadays use GenAI tools like *DeepResearch* (Open-AI 2025) when writing papers, and at the same time, they do not verify all the suggestions they get from their tools, it is a matter of fact that hallucinated references are already entailed in some bibliographies of published papers. Even misleading persistent identifiers or links to records in bibliographic databases or library catalogs can be easily fabricated by LLMs. In an ideal world, authors would check all their references, and journals would take action against such violations of our scientific practice – namely, citing non-existent papers.² Furthermore, we now have an increasing number of preprints that are not checked by third parties. All this leads to the situation in which we find bibliographies that entail never-written papers and books (DeGeurin 2024; Tramèr 2025). In parallel, there are freely accessible web search engines like *Google Scholar* that index the full text or metadata of scholarly literature across an array of publishing formats and disciplines. Such web search engines do not check for the existence of the referenced papers, and therefore, they become polluted by hallucinated references, which other researchers then use in good faith. Sure, there were false references before LLMs came along; however, we should be aware that LLMs make it more likely that plausible-sounding references will find their way into our search engines. So, even if we don't use GenAI tools ourselves, we should approach established search engines with a new sense of skepticism in the future. It remains to be seen what measures the academic community will establish to prevent hallucinated references from finding their way into peer-reviewed articles.

In a similar way, other faulty outputs of LLMs can pollute formerly trustworthy collections of text and will also be used for the training of future LLMs. This assumption is, for example, supported by a recent study that indicates the increasing distribution of GenAI-generated text in scholarly papers (Liang et al. 2025), as well as by a database that collects suspected undeclared AI usage in the academic literature (Glynn n.d.; Jacobs 2025). Another case concerns, for example, legal judgments based on non-existent previous judgments (Charlotin n.d.; Maruf 2023). Such observations can be made in many domains and demonstrate that over-reliance on LLMs can have disruptive consequences (Hopster 2021).³

The consequences this will have for the suitability of published texts, especially those distributed on the Internet, in terms of their use for knowledge acquisition, cannot yet be fully assessed. In view of the presumable huge quantity by which texts will be polluted by machine-generated text, I think that this has the potential to contribute to an epistemological crisis because it can question formerly trustworthy sources. Moreover, the increasing indistinguishability between human-created and machine-generated text already presents a severe obstacle when we try to take action. I will now turn to the question of whether increasing GenAI use can contribute to the homogenization of written text.

3. Homogenization through the use of AI tools for writing

Homogenization is, of course, something we can observe since the invention of the printing press, the publications of dictionaries, and the education of language grammar in schools, by which spoken and written language has become more standardized (Sasaki 2017). A recent major contributor to language standardization has been the mass media.

² One case that became public knowledge is an article published in PlosOne, which was retracted 45 days after publication because 18 of the article's 76 references could not be identified (PLOS ONE Editors 2024).

³ Another example of the unreliability of LLM-generated text is a report in *The Guardian* about Amazon selling mushroom-picking guides in which poisonous mushrooms were described as edible (Milmo 2023).

In this section, I shall report from several studies that investigate whether the incorporation of GenAI tools may be an additional factor that can lead to a decrease in individuality and diversity. To this end, I will focus on studies that investigate the collective level. That means that the question here is not whether AI tools can help an individual to come up with more ideas, but whether texts that are created with the help of AI tools do have similarities with each other. In other words, LLM-driven homogenization of creative outputs can be described as an effect of algorithmic monoculture (Bommasani et al. 2022; Kleinberg and Raghavan 2021), as multiple actors rely on the same kinds of tools to generate text. Since LLMs are trained to reproduce statistically likely results, they may support homogenization. This can concern the conversational tone, the language use across genres, as well as the shaping of individual language production (Levent and Shroff 2023). Those considerations can be related to the ideas Marshall McLuhan developed in his work with respect to the influence of media (McLuhan 1964). Especially, a careless use of AI tools may foster the influence that machine-generated text will have on us. Another factor that might contribute to this is that people may develop a general trust in such machines due to diverse trustworthiness cues, which make it more likely that suggestions by LLMs are taken up without close examination. Before the rise of LLMs, there was a study that indicated that even single-word suggestions given by smartphone keyboards encourage predictable writing (Arnold et al. 2020).

Nowadays, we can find more and more studies investigating homogenization effects due to the use of LLMs. A recent study by researchers of Cornell University with 118 participants from India and the US that explored cultural bias in AI models could show that AI suggestions homogenize writing toward Western styles and thereby diminish cultural nuances (Agarwal et al. 2025). The participants completed short writing tasks designed to elicit cultural values and artifacts. One group was allowed to use AI tools, while the control group completed the same tasks without AI. Using a similarity metric to check for a hypothesized increased similarity of texts created with the help of AI, this study could show that AI made writing more homogeneous, with respect to common phrasing, style, structure, and content. Since the AI tools influenced Indian and American participants to write more similarly, it also contributed to a decrease in cultural differences in writing.

Similar findings, although with fewer participants (only 54 participants distributed in three groups: brain-only group, search-engine group, and LLM group), can be reported from the widely discussed MIT study titled ‘Your brain on ChatGPT’ (Kosmyna et al. 2025); they

found that the brain-only group exhibited strong variability in how participants approached essay writing across most topics. In contrast, the LLM group produced statistically homogeneous essays within each topic, showing significantly less deviation compared to the other groups. (Kosmyna et al. 2025, p.130)

Comparing the influence of two distinct creativity support tools (CSTs), one LLM-based and the other a version of an Oblique Strategies deck,⁴ Anderson and colleagues came up with similar results with respect to the collective level (Anderson et al. 2024). Investigating 33 participants, this study found that LLM-based CSTs have a stronger homogenization effect on human-in-the-loop divergent ideation processes at the group level. In addition, this study also evaluated the individual level and found that ChatGPT users exhibited greater fluency, flexibility, and elaboration than users of the alternative CST, even though ChatGPT didn’t help the users to develop truly original ideas.

In another study with 38 participants, the effects of two models, the baseline GPT-3 and the further developed InstructGPT model, were compared using a short-form argumentative essay writing task (Padmakumar and He 2024). This study found an increased homogenization effect from LLM assistance at the lexical and content levels for the further developed instruction-tuned LLM. This may suggest that further developed models could increase a homogenization effect.

⁴ This is a card-based method for promoting creativity originally created by the artists Brian Eno and Peter Schmidt in 1975 (“Oblique Strategies” 2025).

In a recent study with 293 participants evaluating GPT-4, a homogenization effect with respect to collective diversity of novel content was found in the context of short-form fictional narrative writing. However, it is worth noting that this study also demonstrated that the use of GenAI increased individual creativity (Doshi and Hauser 2024).

Further research on opinionated LLM-based tools could show that there are influences on the user's opinions in argumentative writing (Jakesch et al. 2023). Such influences are evident even in cases where the user disliked the tool's suggestions (Bhat et al. 2023). It seems as if the LLM-generated text has the potential to shape the ideas of LLM's users even when they do not incorporate LLM-generated text directly into their writing (Roemmele 2021).

A somewhat more positive perspective on the influence of LLMs is suggested by a study by Farhana Shahid and colleagues (2025). They investigated the effects of using LLMs to refine comments on threads related to Islamophobia and homophobia, and found that the resulting comments were more constructive, more positive, less toxic, and retained the original intent of the authors. Nevertheless, they also observed that LLMs often distorted people's original views; this was especially obvious when their views were on a spectrum instead of being outright polarizing.

Investigating the assumed high-quality writing capabilities of LLMs, Chakrabarty and colleagues compared stories generated by three top-performing LLMs (GPT3.5, GPT4, and Claude V1.3) with stories written by expert human writers. The results of this study indicate that the human-written stories substantially outperform the LLM-generated stories on all dimensions of creativity measured by the Torrance Test of Creative Thinking (Torrance 1966),⁵ including the originality dimension (Chakrabarty et al. 2024).

In an article titled 'A.I. Is Homogenizing Our Thoughts' and published in the *New Yorker*, Kyle Chayka presents a nice overview of some of the above-mentioned studies based on several interviews with the involved researchers (Chayka 2025). For example, he cites Nataliya Kosmyna, the first author of the above mentioned MIT-study who said that, *'the output was very, very similar for all of these different people, coming in on different days, talking about high-level personal, societal topics, and it was skewed in some specific directions, [...] no divergent opinions being generated, average everything everywhere all at once—that's kind of what we're looking at here'*. Chayka points out that AI is a technology of averages because LLMs are trained to spot patterns across vast tracts of data. This is, from his point of view, a reason why the answers tended toward consensus, with respect to the quality of the writing and the underlying ideas. Furthermore, he reported that one of the researchers of the above-mentioned study from Cornell University (Agarwal et al. 2025), Aditya Vashistha, compared the AI to *'a teacher who is sitting behind you every time you are writing, saying, 'This is the better version.'* Vashistha's co-author, Mor Naaman, told Chayke that AI suggestions *'work covertly, sometimes very powerfully, to change not only what you write but what you think'*; therefore, a widespread use of GenAI tools may result, over time, in a shift in what people think.

In this context, it is worth mentioning that even though neither humans nor detection software can reliably distinguish between AI-generated text and human-created text with certainty (Brown et al. 2020; Clark et al. 2021; Dugan et al. 2020; Gao et al. 2023; Porter and Machery 2024; Schwitzgebel et al. 2023; Weber-Wulff et al. 2023), there are promising studies indicating the use of GenAI tools by checking differences in word frequencies. For example, typical ChatGPT words like 'delve' (Hern 2024; Shapira 2024), 'realm', 'intricate', 'showcasing', and 'pivotal' are increasingly used in scientific publications (Liang et al. 2025). Since the widespread use of GenAI tools will, of course, also influence word frequencies in human-written text, it remains to be seen how this research will develop.

⁵ Since evaluating quality is a complicated endeavor, the authors of this study did empirically validate TTCW as an evaluation protocol for creativity with respect to fictional short stories by building first a benchmark consisting of 48 short stories: 12 stories written by professionals, and 36 by three top performing LLMs (ChatGPT, GPT4, and Claude 1.3) with 1400 words on average per story before recruiting a set of 10 experts.

Another factor that is likely to foster further homogenization effects is the fact that future models will presumably suffer from self-reinforcing loops. That means they will be trained on a large amount of machine-generated text, given its widespread use (Marr 2024).

In sum, one can interpret recent results in studying homogenization in the context of GenAI tools as an indicator that we are now confronted with a new language homogenization driver that may have a similar impact on language evolution as the invention and adoption of print and mass media have had in earlier times.

4. Deskilling

Especially in educational domains, there is a growing concern that students using GenAI in a careless way – what may lead to an over-reliance on Gen AI tools – will eventually miss the opportunity to develop and maintain important cognitive skills like critical thinking, understanding complex and long texts, and making their minds and thoughts understandable through writing. This is despite the undeniable advantages GenAI tools can add in streamlining research processes, enhancing academic efficiency, and playing a positive role in tutoring. Unfortunately, we do not yet have many studies published that investigate whether deskilling effects are already in place.

There is no question that various kinds of cognitive skills are necessary for academic success, professional competence, and informed citizenship. Cognitive processes enable us to succeed in problem-solving, decision-making, and reflective thinking, and they are crucial for navigating in our complex and dynamic environments. With the rise of GenAI, we now have tools at our disposal with which we can outsource various cognitive tasks, we can use them to get summaries of certain text corpora, to get an overview of relevant literature, to formulate an argument, an abstract, or a conclusion.

Even though the use of AI assistants can reduce our mental effort by automating certain cognitive tasks and thus freeing up mental resources for other things, I believe it is not unreasonable to fear that excessive dependence on these tools and the outsourcing of many cognitive tasks to external aids could lead to a decline in reflective thinking and hinder the development of important cognitive abilities and expertise. In the following, I will report from the first paradigmatic results from different kinds of studies, including experimental investigations, questionnaire-based explorations, and a literature review that investigated this question.

Even though there are still far too few experimental studies to date and many have only examined a small number of participants (e.g., the MIT study ‘Your brain on ChatGPT’ (Kosmyna et al. 2025) examined only 54 people), they may nevertheless be considered an initial weak indication of the extent to which the use of AI tools can bring about a change concerning our cognitive skills. The MIT study aimed to investigate the cognitive cost of using an LLM in the educational context of writing an essay by using electroencephalography (EEG) to record participants' brain activity, together with *Natural Language Processing* analysis, and interviews with the participants. To this end, they compared three groups, two that used a designated tool – the LLM group and the search engine group – and one that used no tool – the brain-only group – while writing an essay. The EEG analysis indicated that LLM, search engine, and brain-only groups had significantly different neural connectivity patterns, reflecting divergent cognitive strategies, and that the brain connectivity systematically scaled down (brain-only group > search engine group > LLM group). In a second step, with only 18 remaining participants, for which it is clear that a larger participant sample is needed to confirm the indicated result, they changed the assignment of the tools. That means that the participants who started with using an LLM now were assigned to use no tool. Evaluating the collected EEG measures, the study revealed that LLM-to-brain participants exhibited weaker neural connectivity and under-engagement of alpha and beta networks. Interestingly, they also demonstrated a decline in their ability to quote from the essays. Therefore, the authors of this study come to the conclusion that their findings support an educational model that delays AI integration until learners have engaged in sufficient self-driven cognitive effort. In addition, they

emphasize the urgent need for longitudinal studies in order to understand the long-term impact of the LLMs on the human brain. Future research will show whether the conclusions the authors of this study draw from their preliminary results can be supported.

A study that examined a larger number of participants was conducted by Michael Gerlich with 666 participants to confirm two hypotheses, namely whether a higher AI tool usage is associated with reduced critical thinking skills and whether cognitive offloading mediates the relationship between AI tool usage and critical thinking skills. According to Gerlich, “students with strong critical thinking skills tend to perform better academically, as they can understand complex concepts, analyse texts, and construct well-reasoned arguments” (Gerlich 2025, p.3). In this study, the author developed a structured questionnaire based on validated scales and existing methods to measure AI tool usage, cognitive offloading, and critical thinking skills to investigate the impact of AI tool usage on cognitive skills. The focus of this study was on critical thinking, as the authors aimed to investigate whether cognitive offloading is a potential mediating factor for a decrease in critical thinking skills. The findings of this study indicate that higher usage of AI tools is associated with reduced critical thinking skills, and cognitive offloading plays a significant role in this relationship.

A systematic literature review examined 14 papers⁶ that covered the implications of students’ over-reliance on GenAI tools (Zhai et al. 2024). The result of this study underscored a significant impact of overdependence on essential cognitive abilities, including decision-making, critical thinking, and analytical reasoning. The authors come to the conclusion that we can observe a trend towards a potential erosion of critical cognitive skills due to challenges such as misinformation, algorithmic biases, plagiarism, privacy breaches, and transparency issues that are features of the use of GenAI tools.

Even though self-reports are a questionable method to evaluate the consequences of AI use, it is interesting to see that a study that used an online questionnaire to examine self-reports of 300 college students in South Korea delivered the result that the top five negative effects of AI dependency mentioned were increased laziness, the spread of misinformation, decreased creativity, and reduced critical and independent thinking (Zhang et al. 2024).

Approaching the subject from a theoretical perspective, I think it is obvious that one has to distinguish between the performance and the skill acquisition (learning). Even if using an AI tool in certain domains may have the potential to improve performance, it can at the same time hinder the learning of the skills needed to succeed with similar tasks without AI assistance, and it may also have detrimental effects on existing cognitive skills. One could object that the degradation of certain human skills may not be critical if the reliance on AI can lead to optimal performance. However, one should be aware that especially GenAI tools have limitations with respect to reliability; their outputs are in need of verification by experts. That’s why using AI tools and accepting deskilling as a trade-off can be especially problematic because the users will then also lack the ability to evaluate the performance of the AI systems and thereby will be vulnerable to system errors.

Focusing on AI assistants used in radiology (Hosny et al. 2018), Brooke Macnamara and colleagues (2024) assume that future radiologist trainees trained with AI assistance may not develop the same visual detection skills as former trainees that have been trained without AI assistance because such AI-learning aids are rather designed to prepare trainees for work with AI assistants and therefore are not focusing on developing learners’ cognitive skills independent of AI. They argue that this kind of dependence on AI tools is especially problematic in the domain of medicine as one needs human expertise to handle, for example, new viruses, unique injuries from accidents, or individual context and significant variations in a patient’s physiology and conclude that

⁶ This were the 14 papers that were included in this literature review: (Abd-alrazaq et al. 2023; Dergaa et al. 2023; Duhaylungsod and Chavez 2023; Gao et al. 2023; Grassini 2023; Kim et al. 2023; Koos and Wachsmann 2023; Lee et al. 2023; Malik et al. 2023; Marzuki et al. 2023; Pokkakillath and Suleri 2023; Santiago et al. 2023; Semrl et al. 2023; Watts et al. 2023).

if AI assistants can be developed such that they are near perfect in all (or nearly all) circumstances in the future, then developing and maintaining human skills for that task may be unnecessary. But, in fields where problems rapidly evolve or novel events are likely to occur, such as in medicine, or if the AI tool is biased, offline, or making errors, then developing and maintaining human skills will continue to be advantageous. (Macnamara et al. 2024).

Furthermore, I think that we should keep in mind that human expertise will also be needed to create future AI systems.

5. Weakened human influence

Instead of approaching the larger question of whether the human influence on what is framing our attempts to make sense of our world and to agree on social and moral norms could be weakened, I will limit myself to considerations that focus on decision-making. To this end, I start with the observation that a lot of decision-making is already outsourced to various kinds of AIs. To avoid misunderstandings, I want to emphasize that I don't think that the increase in outsourcing of decision-making is already caused by potential effects of deskilling; instead, I assume that it may be more due to human laziness, the social pressure to make so many decisions, and the general trend that recommends the use of AI. And this was already the case before the advent of GenAI.

The topic of an increasing human dependency on AI or other technologies in almost every walk of life has been a subject of many debates. While I do not aim to ignore that technological progress did contribute to improving living standards and made life easier, I focus here on critical impacts on humans that may turn out to be problematic in the long run. It is outside of the scope of this paper to investigate the whole range of technological progress that has led to automation in many domains. Instead, I will focus on tools that offer us to outsource cognitive effort and pose the question of whether the increasing use of such tools may contribute to weakening the human influence in making choices.

Like John Danaher (2018), who presented a more positive view on the consequences of the usage of AI tools, I suggest that one can view AI assistance as a form of algorithmic cognitive outsourcing, which makes AI assistance a special kind of automation. While in earlier times automation mainly concerned physical, non-cognitive elements of human tasks, nowadays automation entails more and more cognitive elements. The interesting question now is whether cognitive outsourcing is reducing our autonomy.

I will argue that if we replace many of our own choices with AI choices, our autonomous role might be minimized because we will have no access to the rationales and reasons that speak in favor of this recommendation if we just follow AI recommendations. Furthermore, we could be worried that we could gradually be nudged into the conviction that those recommendations present our preferences without we are being in the loop of making them. Of course, all of our choices are influenced by our cultural environment, but if we withdraw ourselves from participating in the evaluation and forming of preferences, we may become more vulnerable with respect to external manipulation. Still, one could object that the influence of religion or other cultural institutions has always influenced our choices; however, this may not be an argument that we should welcome additional factors that contribute to the impact of external influences, especially if mainly some global players like Google, Amazon, Facebook, and Apple control this technology.

In sum, one may say that if we let ourselves be reduced to mere implementers of AI suggestions and let us completely shut off from the rationale and reasons that underlie the AI's recommendations, this will be a threat to autonomy.

6. Ways to handle smart machines

The above consideration may sound as if I am drawn into cultural pessimism – but this is not the case. Thinking about what it means if we are increasingly outsourcing cognitive tasks and moving towards technological dependence can inform us to answer the question of what it means to live a good life in the age of smart machines. Presupposing that the preliminary findings we have on homogenization and deskilling concern a careless use of GenAI, we can develop suggestions that could counteract the consequences of careless use of AI tools.

6.1 *Minimizing the pollution of our text corpora*

To counteract the pollution of our text corpora, we might establish diverse practices to certify verified machine-generated text. For example, academic journals could come to an agreement that they only publish papers in which all the references are checked. Eventually, we could develop institutions that offer the service to certify references.

It might also be an option to develop a cultural practice according to which one would emphasize and make it visible when a text is written by humans or at least verified by adequate expertise by introducing the distribution of watermarks, metadata identifications, cryptographic methods, or other techniques for proving provenance and authenticity of content.

Turning to the discussed measures of potential regulations, we see that, for example, the EU AI act recommends labeling machine-generated text to counteract the decrease of integrity and trust of our information ecosystem in the Recital 133 from the EU-AI act:

A variety of AI systems can generate large quantities of synthetic content that becomes increasingly hard for humans to distinguish from human-generated and authentic content. The wide availability and increasing capabilities of those systems have a significant impact on the integrity and trust in the information ecosystem, raising new risks of misinformation and manipulation at scale, fraud, impersonation and consumer deception. In light of those impacts, the fast technological pace and the need for new methods and techniques to trace origin of information, it is appropriate to require providers of those systems to embed technical solutions that enable marking in a machine readable format and detection that the output has been generated or manipulated by an AI system and not a human. Such techniques and methods should be sufficiently reliable, interoperable, effective and robust as far as this is technically feasible, taking into account available techniques or a combination of such techniques, such as watermarks, metadata identifications, cryptographic methods for proving provenance and authenticity of content, logging methods, fingerprints or other techniques, as may be appropriate. When implementing this obligation, providers should also take into account the specificities and the limitations of the different types of content and the relevant technological and market developments in the field, as reflected in the generally acknowledged state of the art. Such techniques and methods can be implemented at the level of the AI system or at the level of the AI model, including general-purpose AI models generating content, thereby facilitating fulfilment of this obligation by the downstream provider of the AI system. To remain proportionate, it is appropriate to envisage that this marking obligation should not cover AI systems performing primarily an assistive function for standard editing or AI systems not substantially altering the input data provided by the deployer or the semantics thereof. (European Commission 2024, EU-AI Act, Recital 133, available at <https://artificialintelligenceact.eu/recital/133>)

Furthermore, one finds in Article 50 ‘Transparency obligations for providers and deployers of certain AI systems’ the following recommendation:

7. The AI Office shall encourage and facilitate the drawing up of codes of practice at Union level to facilitate the effective implementation of the obligations regarding the detection and labelling of artificially generated or manipulated content. (Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 Laying down Harmonised Rules on Artificial Intelligence and Amending Regulations, 2024)

However, even if laws like this were implemented worldwide, which is not very probable, it would be nearly impossible to check who is following such laws. Therefore, it may be more effective to emphasize the value of verified texts more clearly, either by adding whether a machine-generated text has been verified by an expert or by making it explicit if a text has been written by a human author only, who takes responsibility for it.

6.2 *Avoiding homogenization of language and thinking*

It is difficult to imagine how one could avoid the homogenization effects as long as we observe an increasing use of GenAI in all aspects of life. However, one could hope that a cultural practice could be established that attributes special value to human-generated text. Furthermore, a fatigue effect could perhaps be developed, at least in view of the increasingly similar visual products that one encounters on social media, and thus may motivate us to search for human-made pictures.

Or perhaps it would be conceivable that, instead of the few huge models of the main global players, there would be more and more diverse LLMs, including smaller models that differ significantly in terms of their training data. Nicolas Bauer (2025), for example, suggested that it may be an option to train (fine-tune) a version of an LLM on one's original text, whereby the model could learn linguistic idiosyncrasies that writers like to use: "*Instead of relying on the same version of ChatGPT, which will exhibit streamlined and homogeneous behavior, creating one's own bot would benefit the variance of language produced in the future*" (Bauer 2025, p. 25).

When considering educational domains, it may be advisable to recommend the use of AI tools only after students have developed their own characteristic writing style. Furthermore, in the future, the assessment of exam performance could focus on creativity and originality, as well as the ability to explain content well orally.

6.3 *Minimizing Deskilling*

The main motivation for counteracting deskilling derives from the fact that human expertise is required to verify machine-generated results. In my opinion, this cannot be rendered obsolete by nearby assumed future advances in research as long as our systems are based on the current LLM's architecture. In addition, we cannot ignore that we naturally need human expertise in various domains to develop more advanced GenAI tools. These considerations may be taken as reasons to support the idea that cognitive offloading should only be practiced after the relevant cognitive abilities have been learned and should not be considered the standard procedure. Figuratively speaking, one could imagine that using cognitive abilities could be viewed similarly to how we keep our bodies functioning through exercise.

6.4 *Foster human influence in decision-making*

If we develop a more critical attitude towards careless GenAI use, outsourcing decision-making to AI systems may no longer be considered a desirable practice, and the discussion on rationales and reasons on which decisions are based could experience a vivid revival. This could be supported by further progress in research in the field of explainable AI, which aims to disclose the reasons behind decisions. However, here too, care must be taken to ensure that the explanations provided by the models really do refer to the criteria the AI tools developed through their pattern recognition.

7. Conclusion

In this paper, I have examined the potential negative consequences of careless use of GenAI. Due to limitations in reliability and the resulting erroneous outputs of AI tools, which are not verified and corrected when used carelessly, our environment in which we search for information can become permanently polluted. This could even undermine the strategies we have used to date in finding reliable information, thus escalating into an epistemological crisis. In addition, I reported from initial studies showing that the widespread use of AI tools leads to homogenization of language and thinking. Last but not least, I also presented the first initial empirical results showing that frequent use of AI tools is associated with deskilling with respect to cognitive abilities.

Although the use of AI tools can enhance individual performance, it can hinder the development of the cognitive skills necessary to verify AI tool outputs and foster critical thinking. This means that, especially in education, it seems important to place value on learning those cognitive abilities that can later be offloaded to AI tools. Dependence on AI is problematic because these systems are not always reliable, which makes it necessary to be able to verify the results.

Finally, there are indications that humans are increasingly having less influence on many decision-making processes, as over-reliance on AI tools leads to the practice that AI suggestions are simply followed without the reasons for these suggestions being accessible. Consequently, an increasing dependence on AI for decision-making could reduce human autonomy. And there is a risk that users will believe that AI recommendations reflect their own preferences without being actively involved in the decision-making process.

However, if we establish practices that counteract careless use, we may develop strategies for handling smart machines that at least minimize the negative consequences described above. For example, establishing practices for certifying machine-generated texts could help maintain the integrity of our information ecosystem. Similarly, promoting appreciation for human-generated texts could counteract the homogenization of language and thought. It is also worth considering whether the use of AI tools should only be recommended after users have developed their own writing style. Last but not least, we should negotiate jointly on which areas we want to insist on promoting human influence in decisions.

Statements and Declarations

The author has no relevant financial or non-financial interests to disclose. A pre-publication version of this manuscript can be found on my website (*reference deleted for anonymous review*). No LLM was used in writing this paper, apart from the cautious use of AI tools for the purpose of “AI-assisted copy editing.”

References

- Abd-alrazaq A, AlSaad R, Alhuwail D, Ahmed A, Healy PM, Latifi S, Aziz S, Damseh R, Alabed Alrazak S, Sheikh J (2023) Large Language Models in medical education: Opportunities, challenges, and future directions. *JMIR Medical Education*, 9, e48291. <https://doi.org/10.2196/48291>
- Agarwal A, Wong-Fannjiang C, Sussillo D, Hawley K, Firat O (2018) Hallucinations in neural machine translation. *ICLR*. <https://research.google/pubs/hallucinations-in-neural-machine-translation/>
- Agarwal D, Naaman M, Vashistha A (2025) AI suggestions homogenize writing toward Western styles and diminish cultural nuances. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, 1–21. <https://doi.org/10.1145/3706598.3713564>

- Anderson BR, Shah JH, Kreminski M (2024) Homogenization effects of large language models on human creative ideation. *Creativity and Cognition*, 413–425.
<https://doi.org/10.1145/3635636.3656204>
- Arnold KC, Chauncey K, Gajos KZ (2020) Predictive text encourages predictable writing. *Proceedings of the 25th International Conference on Intelligent User Interfaces*, 128–138.
<https://doi.org/10.1145/3377325.3377523>
- Banerjee S, Agarwal A, Singla S (2025) LLMs will always hallucinate, and we need to live with this. In: Arai K (ed) *Intelligent Systems and Applications*, Springer Nature Switzerland, Vol. 1554, pp 624–648. https://doi.org/10.1007/978-3-031-99965-9_39
- Bauer NF (2025) Does ChatGPT increase language homogenization? In: Vaih-Baur C, Mathauer V, von Gamm E-I, Pietzcker D (eds) *KI in Medien, Kommunikation und Marketing*, Springer Fachmedien Wiesbaden, pp 11–31. https://doi.org/10.1007/978-3-658-46344-1_2
- Bhat A, Agashe S, Oberoi P, Mohile N, Jangir R, Joshi A (2023) Interacting with next-phrase suggestions: How suggestion systems aid and influence the cognitive processes of writing. *Proceedings of the 28th International Conference on Intelligent User Interfaces*, 436–452.
<https://doi.org/10.1145/3581641.3584060>
- Bommasani R, Creel KA, Kumar A, Jurafsky D, Liang P (2022) Picking on the same person: Does algorithmic monoculture lead to outcome homogenization? *Proceedings of the 36th International Conference on Neural Information Processing Systems*, 3663–3678.
- Brown TB, Mann B, Ryder N, Subbiah M, Kaplan J, Dhariwal P, Neelakantan A, Shyam P, Sastry G, Askell A, Agarwal S, Herbert-Voss A, Krueger G, Henighan T, Child R, Ramesh A, Ziegler DM, Wu J, Winter C, Amodei D, et al (2020) *Language models are few-shot learners*.
<https://doi.org/10.48550/arXiv.2005.14165>
- Chakrabarty T, Laban P, Agarwal D, Muresan S, Wu C-S (2024) Art or artifice? Large Language Models and the false promise of creativity. *Proceedings of the CHI Conference on Human Factors in Computing Systems*, 1–34. <https://doi.org/10.1145/3613904.3642731>
- Charlotin D (n.d.) AI hallucination cases database. <https://www.damiencharlotin.com/hallucinations>. Accessed 29 November 2025.
- Chayka K (2025, June 25) A.I. is homogenizing our thoughts. *The New Yorker*.
<https://www.newyorker.com/culture/infinite-scroll/ai-is-homogenizing-our-thoughts>
- Chelli M, Descamps J, Lavoué V, Trojani C, Azar M, Deckert M, Raynier J-L, Clowez G, Boileau P, Ruetsch-Chelli C (2024) Hallucination rates and reference accuracy of ChatGPT and Bard for systematic reviews: Comparative Analysis. *Journal of Medical Internet Research*, 26, e53164.
<https://doi.org/10.2196/53164>
- Clark E, August T, Serrano S, Haduong N, Gururangan S, Smith NA (2021) *All that’s “human” is not gold: Evaluating human evaluation of generated text*.
<https://doi.org/10.48550/arXiv.2107.00061>
- Danaher J (2018) Toward an ethics of AI assistants: An initial framework. *Philosophy & Technology*, 31(4), 629–653. <https://doi.org/10.1007/s13347-018-0317-3>
- DeGeurin M (2024, March 19) AI-generated nonsense is leaking into scientific journals. *Popular Science*. <https://www.popsci.com/technology/ai-generated-text-scientific-journals/>
- Dergaa I, Chamari K, Zmijewski P, Ben Saad H (2023) From human writing to artificial intelligence generated text: Examining the prospects and potential threats of ChatGPT in academic writing. *Biology of Sport*, 40(2), 615–622. <https://doi.org/10.5114/biolsport.2023.125623>
- Doshi AR, Hauser OP (2024) Generative AI enhances individual creativity but reduces the collective diversity of novel content. *Science Advances*, 10(28), eadn5290.
<https://doi.org/10.1126/sciadv.adn5290>
- Dugan L, Ippolito D, Kirubarajan A, Callison-Burch C (2020) RoFT: A Tool for evaluating human detection of machine-generated text. *arXiv*. <https://doi.org/10.48550/arXiv.2010.03070>

- Duhaylungsod AV, Chavez JV (2023) ChatGPT and other AI users: Innovative and creative utilitarian value and mindset shift. *Journal of Namibian Studies : History Politics Culture*, 33. <https://doi.org/10.59670/jns.v33i.2791>
- Fauzi F, Tuhuteru L, Sampe F, Ausat AMA, Hatta HR (2023) Analysing the role of ChatGPT in improving student productivity in higher education. *Journal on Education*, 5(4), 14886–14891. <https://doi.org/10.31004/joe.v5i4.2563>
- Gao CA, Howard FM, Markov NS, Dyer EC, Ramesh S, Luo Y, Pearson AT (2023) Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers. *Npj Digital Medicine*, 6(1), 75. <https://doi.org/10.1038/s41746-023-00819-6>
- Gerlich M (2025) AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), 6. <https://doi.org/10.3390/soc15010006>
- Glynn A (n.d.) Suspected undeclared AI usage in the academic literature. *Academ-AI*. <https://www.academ-ai.info>. Accessed 29 November 2025.
- Grassini S (2023) Shaping the future of education: Exploring the potential and consequences of AI and ChatGPT in educational settings. *Education Sciences*, 13(7), 692. <https://doi.org/10.3390/educsci13070692>
- Hern A (2024, April 16) TechScape: How cheap, outsourced labour in Africa is shaping AI English. *The Guardian*. <https://www.theguardian.com/technology/2024/apr/16/techscape-ai-gadgest-humane-ai-pin-chatgpt>
- Hopster J (2021) What are socially disruptive technologies? *Technology in Society*, 67, 101750. <https://doi.org/10.1016/j.techsoc.2021.101750>
- Hosny A, Parmar C, Quackenbush J, Schwartz LH, Aerts HJWL (2018) Artificial intelligence in radiology. *Nature Reviews Cancer*, 18(8), 500–510. <https://doi.org/10.1038/s41568-018-0016-5>
- Jacobs P (2025) One-fifth of computer science papers may include AI content. *Science*, 389(6760), 558–559. <https://doi.org/10.1126/science.zxxd90o>
- Jakesch M, Bhat A, Buschek D, Zalmanson L, Naaman M (2023) Co-Writing with opinionated language models affects users' views. *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, 1–15. <https://doi.org/10.1145/3544548.3581196>
- Jumper J, Evans R, Pritzel A, Green T, Figurnov M, Ronneberger O, Tunyasuvunakool K, Bates R, Žídek A, Potapenko A, Bridgland A, Meyer C, Kohl SAA, Ballard AJ, Cowie A, Romera-Paredes B, Nikolov S, Jain R, Adler J, Hassabis D et al (2021) Highly accurate protein structure prediction with AlphaFold. *Nature*, 596, 583–589. <https://doi.org/10.1038/s41586-021-03819-2>
- Kalai AT, Nachum O, Vempala SS, Zhang E (2025) Why language models hallucinate. *arXiv*. <https://doi.org/10.48550/arXiv.2509.04664>
- Kim Y, Lee M, Kim D, Lee S-J (2023) Towards explainable AI writing assistants for non-native English speakers. *arXiv*. <https://doi.org/10.48550/ARXIV.2304.02625>
- Kleinberg J, Raghavan M (2021) Algorithmic monoculture and social welfare. *Proceedings of the National Academy of Sciences*, 118(22), e2018340118. <https://doi.org/10.1073/pnas.2018340118>
- Koos S, Wachsmann S (2023) Navigating the impact of ChatGPT/GPT4 on legal academic examinations: Challenges, opportunities and recommendations. *Media Iuris*, 6(2), 255–270. <https://doi.org/10.20473/mi.v6i2.45270>
- Kosmyna N, Hauptmann E, Yuan YT, Situ J, Liao X-H, Beresnitzky AV, Braunstein I, Maes P (2025) Your brain on ChatGPT: Accumulation of cognitive debt when using an AI assistant for essay writing task. *arXiv*. <https://doi.org/10.48550/arXiv.2506.08872>

- Lee P, Bubeck S, Petro J (2023) Benefits, limits, and risks of GPT-4 as an AI chatbot for medicine. *New England Journal of Medicine*, 388(13), 1233–1239. <https://doi.org/10.1056/NEJMSr2214184>
- Levent I, Shroff L (2023) The model is the message. *The Second Workshop on Intelligent and Interactive Writing Assistants*. https://cdn.glitch.global/d058c114-3406-43be-8a3c-d3afff35eda2/paper29_2023.pdf
- Liang W, Zhang Y, Wu Z, Lepp H, Ji W, Zhao X, Cao H, Liu S, He S, Huang Z, Yang D, Potts C, Manning CD, Zou J (2025) Quantifying large language model usage in scientific papers. *Nature Human Behaviour*. <https://doi.org/10.1038/s41562-025-02273-8>
- Macnamara BN, Berber I, Çavuşoğlu MC, Krupinski EA, Nallapareddy N, Nelson NE, Smith PJ, Wilson-Delfosse AL, Ray S (2024) Does using artificial intelligence assistance accelerate skill decay and hinder skill development without performers' awareness? *Cognitive Research: Principles and Implications*, 9(1), 46. <https://doi.org/10.1186/s41235-024-00572-8>
- Malik AR, Pratiwi Y, Andajani K, Numertayasa IW, Suharti S, Darwis A, Marzuki (2023) Exploring artificial intelligence in academic essay: Higher education student's perspective. *International Journal of Educational Research Open*, 5, 100296. <https://doi.org/10.1016/j.ijedro.2023.100296>
- Marr B (2024, August 19) Why AI models are collapsing and what it means for the future of technology. *Forbes*. <https://www.forbes.com/sites/bernardmarr/2024/08/19/why-ai-models-are-collapsing-and-what-it-means-for-the-future-of-technology/>
- Maruf R (2023, May 27) Lawyer apologizes for fake court citations from ChatGPT. *CNN Business*. <https://www.cnn.com/2023/05/27/business/chat-gpt-avianca-mata-lawyers>
- Marzuki, Widiati U, Rusdin D, Darwin, Indrawati I (2023) The impact of AI writing tools on the content and organization of students' writing: EFL teachers' perspective. *Cogent Education*, 10(2), 2236469. <https://doi.org/10.1080/2331186X.2023.2236469>
- McLuhan M (1964) *Understanding media: The extensions of man* (Repr). Mentor.
- Milmo D (2023, September 1) Mushroom pickers urged to avoid foraging books on Amazon that appear to be written by AI. *The Guardian*. <https://www.theguardian.com/technology/2023/sep/01/mushroom-pickers-urged-to-avoid-foraging-books-on-amazon-that-appear-to-be-written-by-ai>
- Mugaanyi J, Cai L, Cheng S, Lu C, Huang J (2024) Evaluation of Large Language Model performance and reliability for citations and references in scholarly writing: Cross-disciplinary study. *Journal of Medical Internet Research*, 26, e52935. <https://doi.org/10.2196/52935>
- Oblique Strategies. (2025) In: *Wikipedia*. https://en.wikipedia.org/w/index.php?title=Oblique_Strategies&oldid=1305924047#cite_note-15
- Open-AI. (2025). Eine Einführung in Deep Research. <https://openai.com/de-DE/index/introducing-deep-research>. Accessed 29 November 2025.
- Padmakumar V, He H (2024) Does writing with language models reduce content diversity? In: Kim B, Yue Y, Chaudhuri S, Fragkiadaki K, Khan M, Sun Y (eds), *International Conference on Representation Learning*, pp 642–669. https://proceedings.iclr.cc/paper_files/paper/2024/file/02dec8877fb7c6aa9a79f81661baca7c-Paper-Conference.pdf
- PLOS ONE Editors (2024) Retraction: A comparative analysis of blended learning and traditional instruction: Effects on academic motivation and learning outcomes. *PloS One*, 19(4), e0302484. <https://doi.org/10.1371/journal.pone.0302484>
- Pokkakillath S, Suleri J (2023) ChatGPT and its impact on education. *Research in Hospitality Management*, 13(1), 31–34. <https://doi.org/10.1080/22243534.2023.2239579>

- Porter B, Machery E (2024) AI-generated poetry is indistinguishable from human-written poetry and is rated more favorably. *Scientific Reports*, 14(1), 26133. <https://doi.org/10.1038/s41598-024-76900-1>
- Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence and amending regulations, Pub. L. No. (EU) 2024/1689 (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797, (EU) 2020/1828 (Artificial Intelligence Act) (2024). <http://data.europa.eu/eli/reg/2024/1689/oj/eng>
- Roemmele M (2021, July 8) Inspiration through observation: Demonstrating the influence of automatically generated text on creative writing. *Proceedings ICCV'21*. <https://doi.org/10.48550/arXiv.2107.04007>
- Santiago CS, Embang SI, Acanto RB, Ambojia KWP, Aperochio MDB, Balilo BB, Cahapin EL, Conlu MTN, Lausa SM, Laput EY, Malabag BA, Paderes JJ, Romasanta JKN (2023) Utilization of writing assistance tools in research in selected higher learning institutions in the Philippines: A text mining analysis. *International Journal of Learning, Teaching and Educational Research*, 22(11), 259–284. <https://doi.org/10.26803/ijlter.22.11.14>
- Sasaki Y (2017) Publishing nations: technology acquisition and language standardization for European ethnic groups. *The Journal of Economic History*, 77(4), 1007–1047. <https://doi.org/10.1017/S0022050717000821>
- Schwitzgebel E, Schwitzgebel D, Strasser A (2023) Creating a large language model of a philosopher. *Mind & Language*, 1–23. <https://doi.org/10.1111/mila.12466>
- Semrl N, Feigl S, Taumberger N, Bracic T, Fluhr H, Blockeel C, Kollmann M (2023) AI language models in human reproduction research: Exploring ChatGPT's potential to assist academic writing. *Human Reproduction*, 38(12), 2281–2288. <https://doi.org/10.1093/humrep/dead207>
- Shahid F, Dittgen M, Naaman M, Vashistha A (2025) *Examining Human-AI collaboration for co-writing constructive comments online*. <https://doi.org/10.1145/3757591>
- Shapira P (2024, March 31) Delving into “delve.” *Philip Shapira*. <https://pshapira.net/2024/03/31/delving-into-delve/>
- Strasser A (2025) Why a careless use of AI tools may contribute to an epistemological crisis. *Philosophy & Digitality*, 2(1), 78–96. <https://doi.org/10.18716/PD.V2I1.11662>
- Torrance EP (1966) *Torrance tests of creative thinking: Norms-technical manual*. Personal Press, Incorporated. <https://discover.library.unt.edu/catalog/b1469691>
- Tramèr F (2025, August 3) Trends in LLM-generated citations on arXiv. *SPY Lab*. <https://spylab.ai/blog/hallucinations/>
- Watts FM, Dood AJ, Shultz GV, Rodriguez J-MG (2023) Comparing student and generative Artificial Intelligence chatbot responses to organic chemistry writing-to-learn assignments. *Journal of Chemical Education*, 100(10), 3806–3817. <https://doi.org/10.1021/acs.jchemed.3c00664>
- Weber-Wulff D, Anohina-Naumeca A, Bjelobaba S, Foltýnek T, Guerrero-Dib J, Popoola O, Šigut P, Waddington L (2023) Testing of detection tools for AI-generated text. *International Journal for Educational Integrity*, 19(1), Article 1. <https://doi.org/10.1007/s40979-023-00146-z>
- Zhai C, Wibowo S, Li LD (2024) The effects of over-reliance on AI dialogue systems on students' cognitive abilities: A systematic review. *Smart Learning Environments*, 11(1), 28. <https://doi.org/10.1186/s40561-024-00316-7>
- Zhang S, Zhao X, Zhou T, Kim JH (2024) Do you have AI dependency? The roles of academic self-efficacy, academic stress, and performance expectations on problematic AI usage behavior. *International Journal of Educational Technology in Higher Education*, 21(1), 34. <https://doi.org/10.1186/s41239-024-00467-0>