

THE ETHICAL RISKS POSED BY NEW TECHNOLOGIES IN RESEARCH

Toni GIBEA ^{a*}, Radu USZKAI ^b, Mihail-Valentin CERNEA ^c

a, b, c Bucharest University of Economic Studies, Romania

ABSTRACT

While we have made progress in the field of Artificial Intelligence (AI) for decades, in the past couple of years we have seen that more and more applications of it have left the confines of laboratories and specialists and gained a place in our day to day lives, with countless implications being constantly explored, more recently, especially in the wake of the explosion of Large language models (LLM) like ChatGPT in November 2022. The purpose of our article is that of exploring what are the ethical risks posed by the deployment of such new technologies in education, with a specific focus on research. We will first begin with a brief overview of how AI is used in these fields. Secondly, we will analyze the normative framework for these new technologies highlighting why the dual-use dilemmas framework might be appropriate. Our paper will end with a series of recommendations regarding how the future of research will be shaped by these new technologies and their integration in our day-to-day academic life.

KEYWORDS: *Artificial Intelligence, ChatGPT, dual-use research, Large language models, research ethics.*

DOI: 10.24818/IMC/2023/05.07

1. INTRODUCTION

In the past couple of months, it has become quite a common trope in academic contexts to start a paper or a conference presentation by introducing the topic via a conversation with ChatGPT. Thanks to its impressive algorithm and the immense amount of data that it was trained on and has access to, ChatGPT is able to answer complex questions belonging to the field of research ethics. For instance, when we gave it the following prompt ("How will ChatGPT shape the future of academic research?"). The chatbot told us that it, alongside "similar language models, has the potential to significantly impact academic research in various ways." To be more specific, according to ChatGPT there are ten ways in which it can have such a potentially significant impact: by assisting researchers with *data analysis and interpretation* or in drafting *literature reviews*, helping them *brainstorm* new ideas, offering *writing assistance* (editing papers, correcting spelling errors, etc.), enhancing *interdisciplinary collaboration and communication* in a team of researchers (by drafting emails or helping with translation in the case of international teams), treating it as a virtual *Research Assistant*, for *educational support* (i.e. using it for providing understanding of complex concepts and technical problems) and *customized information retrieval*. In addition to this, ChatGPT claims that it can also help researchers understand the *ethical and social implications* of their research.

However, the chatbot is quick to draw attention to the fact that such tools "offer tremendous potential, ethical considerations, biases in training data, and the need for responsible and transparent use of AI

* Corresponding author. E-mail address: toni.gibea@man.ase.ro.

in research should be carefully addressed to ensure the positive impact on academia" (OpenAI, 2023). When prompted additionally with regards to what it thinks are the main ethical risks, the chatbot compiled a list of eight such main concerns: bias in training data, lack of explainability, unintended generation of misinformation, privacy, overreliance on its recommendations, unintended use and misinterpretation, emotional connections with the AI and security concerns associated with the generation of deceptive content or phishing attacks.

This sample conversation that we had with ChatGPT highlights multiple relevant aspects for our work. Firstly, from a strictly technical perspective, there are ways of citing such conversations in an academic paper, as many of citation style manuals like APA have adapted to this new reality (McAdoo, 2023). Secondly, chatbots like ChatGPT are able to provide some careful and nuanced answers to difficult, complex questions. Last but not least, however, chatbots like ChatGPT are by no means able to offer an exhaustive perspective on many issues, including the ethical risks of its deployment in research. One strong indication for why this is the case is the inclusion of security concerns like phishing instead of listing one issue that was on everyone's lips from the moment that ChatGPT became a part of everyday conversations: its use as a tool that enhances (or initiates a new form of) plagiarism in academia.

While there have been many studies dedicated to the impact that LLMs will have on the job market, from how it will change the world of computer coding (Tian et al., 2023) to marketing (Paul et al., 2023) or human resource management (Budhwar et al., 2023), a major focus was on how these new technologies will shape the field of education and research (Uszkai, 2024). The November 2022 launch of Open AI's ChatGPT 3.5 version (and their subsequent GPT-4) raised serious concerns in most educational institutions, with many rushing to ban it on their campus. In time, while some schools went back on their original decision and allowed it, many administrators and scholars have started to argue that the future of education should incorporate AI and bring about significant changes in school curricula (Hsu, 2023; Roose, 2023; Ceres & Hoover, 2023). While there is a clear upside of integrating AI into education, as chatbots can become useful companions to teachers and students in the learning process (Kasneci et al., 2023), many have argued that such new technologies would hamper the epistemic goals of education as they can be used by students as a cheap and effective way to do their homework with a minimal opportunity cost (Exance, 2023; Uszkai, 2024).

Unsurprisingly, LLMs also have the potential to disrupt how scientific research is conducted. In a recent paper published in *Nature*, Van Noorden and Perkel (2023) present the results of a survey they conducted on 1600 researchers who were polled on their opinions regarding the future of science in the wake of powerful AI tools. More than half of the respondents pointed out that AI will either be essential, or at least very important when thinking about the way in which the future of doing science will look like. What they highlighted in particular was the fact that AI will help them process data faster (66% of respondents), do computations that were previously unfeasible (58%) and it overall serve as a cost and time effective tool (55%). Furthermore, what the survey highlighted was that, the more researchers directly use AI in their work, the more they tend to believe that it will become essential in their field. The way in which they think about the barriers in using AI in research are also contingent upon their field of research and how they go about their work. For instance, the "researchers who directly studied AI were most concerned about a lack of computing resources, funding for their work and high-quality data to run AI on. Those who work in other fields but use AI in their research tended to be more worried by a lack of skilled scientists and training resources, and they also mentioned security and privacy considerations. Researchers who didn't use AI generally said that they didn't need it or find it useful, or that they lacked experience or time to investigate it." (Van Noorden & Perkel, 2023, p. 673).

The survey also pointed out that more than half of the respondents raised a wide array of concerns (both epistemic and moral) about the deployment of AI tools in research, like an increase in academic fraud (55%), bias (58%), replication problems (53%), more and potentially undetectable plagiarism

(68%), misinformation (68%) and the proliferation of inaccuracies and mistakes in academic papers (66%). Furthermore, 69% of the respondents express concern about how scientific research might change as a result of integrating AI tools as research might "lead to more reliance on pattern recognition without understanding" (Van Noorden & Perkel, 2023, p. 673). Last but not least, an additional worry regarding the future of academic publishing of scientific research in the age of AI was related to the ability of journal editors and reviewers to accurately evaluate papers that have used ChatGPT and other similar tools.

The worries that Van Noorden and Perkel (2023) identified in their survey are shared by other scholars as well. Following a paper published in *JAMA Ophthalmology* by Taloni et al. (2023), in which the Italian scholars created a fake clinical trial in support of an unverified hypothesis, Naddaf (2023) highlights how easy it is for researchers (especially in the natural sciences) to create believable data which, in a world with a dysfunctional peer-review system, becomes highly problematic. For Conroy (2023), we are already living in an "AI-assisted writing and reviewing" world filled with both possibilities and ethical risks. On the one hand, AI tools have the potential to transform all aspects of how we are doing science nowadays, from actual research to publishing and the communication of our results. It is especially a useful tool in a world with linguistic injustice (Uszkai, 2015) in which many non-native English speakers need to publish in English in order for their work to actually have an impact in their respective epistemic communities. However, their role in promoting more linguistic fairness is contingent upon market conditions and how effective free LLM tools will prove to be. The drawbacks of relying on LLMs are, however, quite obvious. Hallucinations and AI-generated fake data are currently posing serious issues to major publishers (Conroy, 2023; Else, 2023; Pournaras, 2023), with even the way in which research projects get funded being at risk by the ability of scholars to enlist ChatGPT to help write their funding applications (Parrilla, 2023). As a result, the whole research infrastructure is at risk and many reforms will need to be implemented in the near future, from devising a new ethical infrastructure to eliminating some redundancies.

2. THE NORMATIVE FRAMEWORK

As mentioned in the first part, LLMs can assist researchers in various tasks, from co-authoring papers (ChatGPT Generative Pre-trained Transformer & Zhavoronkov, 2022; O'Connor & ChatGPT, 2023) to guiding researchers on how to conduct their empirical studies (e.g. checking how to verify a hypothesis, proposing statistical tests, interpreting the statistical results, conducting qualitative analysis, and so forth).

For our analysis we considered both the current capabilities LLMs possess (ChatGPT, Bard or BLOOM) and potential more powerful versions of a LLM. For example, it is plausible that soon LLM would have access to up-to-date information, a necessary characteristic, especially for effectively assisting researchers during the literature review process. It is well-known that LLMs like ChatGPT lack this capacity and tend to confabulate, suggesting non-existing papers (Rahman et al., 2023). Also, GPT-4 made a huge step forward with its Advanced Data Analysis (ADA) module that allows users to upload and download files for analyzing data. In our paper we imagined an even more powerful version that conducts a complex set of statistical analyses and interpretations. Although we will discuss the ethical risks these probable versions of LLMs raise, we did not assume the existence of what is labeled by computer scientists as artificial general intelligence (AGI). Not because we are skeptical that it is less likely to happen, but because AGI would raise another more pressing normative debate regarding the normative status of such an artificial agent. Given the scope of our paper, it is enough to look at LLMs as simple research tools that raise complex ethical conundrums.

Parrilla (2023) advances a percentage without backing it up, that nearly 15% of researchers use ChatGPT for grant writing. He avoids qualifying this practice as cheating and prefers to favor a revisionist account of the grant application system itself, arguing that since certain 'dead documents'

can be effectively written by the AI, hence these documents should be eliminated altogether from the grant application process. We prefer to address the ethical queries directly although in the end we agree that the whole spectrum of how science is founded is not the only way we could advance our knowledge.

The general discussion about the normative frameworks of AI technologies is extensive, but Constantinescu et al. (2021) manage to summarize them and show that we can dissociate between at least two different normative approaches when assessing the ethics of responsible AI, a top-down approach specific for utilitarianism, deontology, and Principlism and a bottom-up one specific for virtue ethics. All these approaches are relevant to context, but we believe that the most effective way to identify the ethical risks is by looking at what kind of dual-use dilemmas LLMs generate. Dual-use dilemmas are known in biomedical sciences as situations when certain research outputs, which are not in itself bad, have a great potential for harming as well as for doing good (Miller & Selgelid, 2007). The same thing is valid for the case of LLMs; these digital tools are not bad in themselves, they can be used to conduct good research, but the same tools can be used to do a lot of bad research. We will present and discuss several dual-use dilemmas based on central concepts from three normative frameworks, namely deontology, utilitarianism and virtue ethics.

From a deontologist perspective, LLMs could enhance researchers' autonomy by assisting non-native English researchers with their writing. Not just that non-native English speakers must develop clear research ideas, but they also need to communicate them effectively in English. Some argue that this situation is also unjust, and native English speakers can be condoned as free riders (Parijs, 2011). The autonomy of a researcher is greatly enhanced in this respect because a LLM is quite effective at improving a written text. A similar kind of argument applies in the case of non-statisticians who need to confirm their assumptions empirically. LLMs can also effectively guide researchers when they are using a software that requires a certain level of coding skill for running a certain type of statistical test (R or Python). ChatGPT offers good answers without wasting time browsing through professional blogs or delaying the work until someone could help. The LLM could also easily harm the autonomy of the researchers. For example, LLMs might constrain the researchers to write their scientific papers in a certain style, template or even use a certain kind of wording. Filters deployed by the developers in LLM play a crucial role in this instance. We do not have access to these filters, nor is it easy to tune them. This seriously diminishes the autonomy of the researchers because they don't have access to the kind of assistance a LLM is capable of offering. Regarding the empirical work, there is a high risk that a LLM might guide researchers towards the mainstream empirical designs and analyses repeatedly.

From a utilitarian perspective, the concern for the greater good is paramount. We assumed a direct link between the rapid advancement of science and the greater good. We know how problematic this assumption is, especially if we ignore the problem of enhancing the moral capacities of humans (Persson & Savulescu, 2012). It is well beyond the scope of this paper to delve into these problems, although we agree that they are relevant. In this specific case, it is highly plausible that LLMs accelerate the empirical research process. Researchers would stop wasting time on writing reports or the descriptive parts of their papers because LLMs can successfully conduct such tasks without making mistakes. Instead, the researchers would have the freedom to focus on the important parts, like interpreting the results or how to advance their work at another level. Also, the ADA module of ChatGPT could be used to simulate data sets; this is useful when checking if every step of the process is robust before gathering the actual data. Better and more efficient data gathering should enhance our understanding in highly experimental fields like health sciences but also other interdisciplinary fields like political sciences, moral psychology, or behavioral economics. Since LLMs enable professionals from other fields to perform better empirical research, we can assume that these are also steps towards the greater good. Sadly, the same technological capabilities could also make us distrust empirical sciences altogether. As we already mentioned in the introduction, Taloni et al. (2023) and

her team used GPT-4 to generate a set of empirical data that confirms a given hypothesis. Data fabrication could reach new heights with ChatGPT. Researchers could be easily tempted to cheat and ask ChatGPT to tweak the data just a bit so that no one figures out what we have done, but enough to make our tests statistically significant. Researchers could also be way more effective at eliminating groups of responses that mess with our analysis although there is no good statistical reason for eliminating them. The ease of performing such accurate interventions could do a lot of harm.

An up-to-date LLM could also signal untested and plausible hypotheses from the recent literature. This would be especially useful for young researchers who did not have enough time to go through the literature extensively. Other researchers could also use such a LLM to generate agreeable hypotheses for raising the chances of publishing their work. We could also use LLMs to identify hypotheses which are agreeable by a group of researchers that we believe will evaluate our grant applications. These tools can guide researchers on how to become effective 'yes-men' if the main incentives are to ensure the funding of their teams or institutes. Even more damaging, if the main incentive is to bring as much money as possible to become as influential as possible in an institution, the goal of conducting good research is second at best in these dark scenarios. Most of the researchers might conclude that they are wasting their time thinking at how to do things differently and use a LLM as a helpful tool and become more obedient. Innovative new research designs are riskier in our publish or perish culture compared to mainstream science.

Honesty and integrity are essential character traits for researchers and they promote what many consider to be the true ethos of science (Evans et al., 2023). A catalyst for academic dishonesty is the need to be part of a research network. People tend to hide their true views regarding certain studies or methods when they believe that others might judge them. More locally, in an interdisciplinary team, a statistician tends to have great authority and most of the other members are reluctant to ask why a certain statistical test was used instead of another. The need for being liked is also present in multiple ways and at different stages, this fact makes us be less honest and give up our integrity just a bit at a time. LLMs could help us because there is no reason to care about what ChatGPT thinks about us or that it will judge us. A similar effect was confirmed in the case of patients' willingness to disclose sensitive information to an artificial agent (Lucas et al., 2014). We do not care if we ask a basic question like what 'p' is because it generates no reputational consequences. One might judge you quite harshly if you ask such a basic question after you've published several empirical studies or you've mentioned multiple empirical findings to support your ideas. Fortunately, with LLMs, you can practice your virtues as a researcher without being judged.

The other side of this scenario is that some people might be tempted to take credit for a LLM making some researchers even more dishonest. You know that there is no way ChatGPT will disclose the fact that it was the one that suggested how to modify your cool lines of code. Of course, even now, people might be tempted to pretend that the things they found on professional blogs belong to them. But in this very instance it is easy to catch someone because that piece of information is publicly available. Your interaction with a LLM is not publicly available. In case your interaction with LLM is productive in terms of 'original' research outcomes it is impossible to check to what extent it was the contribution of the human agent. Fewer people will have a true incentive to act as honest researchers and conduct their activity with integrity.

3. GUIDELINES AND RECOMMENDATIONS

Given the ethical and epistemic challenges outlined in the previous parts of our paper, there is an obvious need for the development of recommendations and guidelines for the use of Artificial Intelligence in scientific research. There is no shortage of documents that purport to provide governance for AI usage in a large variety of domains, created by a large variety of private and public organizations. A recent meta-analysis discovered 200 such documents, warning that, in actuality, this

is a small dataset, as many more documents are available in many languages around the world (Corrêa et al., 2023). Interestingly, most of these guidelines converge, according to the meta-analysis cited, on five aggregated principles: transparency/ explainability/ auditability, reliability/ safety/ security/ trustworthiness, justice/ equity/ fairness/ non-discrimination, privacy and accountability/ liability. While this research is not concerned with the conduct of science in particular, it's worth noting that some of the documents analyzed by Corrêa et al. come from research organizations and the principles are applicable in this context.

Literature into the right kind of guidelines for scientific research tends to emphasize the importance of accountability and transparency (Pournaras, 2023; Hine, 2021; Zielinski et al., 2023) and the use of tools such as impact statements, checklists and codes of ethics for the operationalization of such principles (Bernstein et al., 2021; Srikumar et al., 2022; Watkins, 2023).

Bockting et al. (2023) argue for the creation of living guidelines, i.e. guidelines that get revised as soon as new relevant data becomes available. The reasoning behind the necessity of living guidelines is the speed in which new developments seem to occur in the AI industry. Any body seeking to regulate the use of AI in science cannot just revise guidelines once every few years as the impact of new AI technology is immediate in our connected world.

Developed during two summits at the Institute for Advanced Study at the University of Amsterdam involving members of multinational scientific institutions and policy makers, the living guidelines center on three key principles: accountability, transparency and independent oversight. In this case, accountability refers to the need to keep humans in the loop specifically during the most crucial moments of the process of research – this keeps the researcher accountable for any errors that may happen. Transparency refers to the duty of researchers to disclose any use of generative AI and during which part of the development of the research it was used. The authors also assert that this principle should be extended to the companies that develop generative AI technologies in order for us to understand how they function and how they might affect the quality of research. Independent oversight is needed because of legitimate concerns regarding the ability of the AI industry to self-regulate – the incentives just do not seem to be there for self-regulation given the competitive nature and the high value of the industry. This is why an independent board needs to be established, set up in such a way that it can audit researchers and the AI industry without political and commercial interference (Bockting et al., 2023, p. 694).

These key principles are crystalized in 13 guidelines divided between the three main categories of stakeholders in research: (a) researchers, reviewers and editors of scientific journals, (b) LLM developers and companies and (c) research funding organizations. The authors then detail six steps to operationalize the living guidelines:

- 1 The establishment of a scientific body for auditing the use of AI technology in research.
- 2 The establishment of a second committee to “keep the living guidelines living” in close contact with the auditing body.
- 3 Obtaining the necessary funding to sustain the two committees and the guidelines in such a way as to maintain the independence and integrity of both bodies.
- 4 Seeking legal status for the guidelines so that they can be enforced both at the national level and the international level.
- 5 Finding ways to collaborate with the tech industry without compromising the independence of the two committees.
- 6 Address whatever possible problematic situations that have not been covered by the proposed guidelines.

The cornerstone of this most comprehensive proposal is the creation of the international oversight body that should audit and revise the living guidelines. We think that the authors are right in emphasizing the need for independent oversight. While many versions of guidelines and recommendations can be found in the literature, without some mechanism of enforcement, there is no

guarantee that researchers, editors or financing institutions will follow them. The global nature of the auditing body is necessary not only from the perspective of unifying the approach, given how much research is starting to be done by multinational interdisciplinary teams, but also from the perspective of the ease of communicating and correlating research findings by these teams.

There are two main obstacles that stand in the way of the proposal provided by Brocking et al. that need to be taken into consideration. First, the current geopolitical context and the dual use aspect of both using AI in scientific research and the research of AI poses significant problems. It is quite unclear whether the kind of international cooperation involved in building such an institution is possible at this time and whether the main state actors have any incentive to contribute constructively in this respect. Second, the financial cost of such an undertaking is not to be underestimated, given the organizational agility involved - the auditing body needs to react fast to keep researchers from the global community from misusing AI. The committee in charge of developing the guidelines needs to move even faster to provide the necessary regulatory framework for the auditing body. The high costs of these activities are bound to conflict with the objective of maintaining their independence in the long term.

REFERENCES

- Bernstein, M. S., Levi, M., Magnus, D., Rajala, B. A., Satz, D., & Waeiss, Q. (2021). Ethics and society review: Ethics reflection as a precondition to research funding. *Proceedings of the National Academy of Sciences*, 118, 52. <https://doi.org/10.1073/pnas.2117261118>.
- Bocking, C. L., van Dis, E. A., van Rooij, R., Zuidema, W., & Bollen, J. (2023). Living guidelines for generative AI—why scientists must oversee its use. *Nature*, 622, 7984, 693-696. <https://doi.org/10.1038/d41586-023-03266-1>.
- Budhwar, P. Chowdhury, S., Wood, G., Aguinis, H., Bamber, G.J., et al. (2023). Human resource management in the age of generative artificial intelligence: Perspectives and research directions on ChatGPT. *Human Resource Management Journal*, 33(3), 606-659. <https://doi.org/10.1111/1748-8583.12524>.
- Ceres, P. & Hoover, A. (2023, August 2023). Kids Are Going Back to School. So Is ChatGPT. *Wired*. Retrieved November 28, 2023, from <https://www.wired.com/story/chatgpt-schools-plagiarism-lesson-plans/>.
- ChatGPT Generative Pre-trained Transformer, & Zhavoronkov, A. (2022). Rapamycin in the context of Pascal's Wager: Generative pre-trained transformer perspective. *Oncoscience*, 9, 82–84. <https://doi.org/10.18632/oncoscience.571>
- Conroy, G. (2023). How Generative AI could disrupt scientific publishing. *Nature*, 622, 234-236. <https://doi.org/10.1038/d41586-023-03144-w>.
- Constantinescu, M., Voinea, C., Uszkai, R., & Vică, C. (2021). Understanding responsibility in Responsible AI. Dianoetic virtues and the hard problem of context. *Ethics and Information Technology*, 23(4), 803–814. <https://doi.org/10.1007/s10676-021-09616-9>.
- Corrêa, N.K., Galvão, C., Santos, J.W., Del Pino, C., Pinto, E.P., Barbosa, C., Massmann, D., Mambrini, R., Galvão, L., Terem, & E. de Oliveira, N. (2023). Worldwide AI ethics: A review of 200 guidelines and recommendations for AI governance. *Patterns*, 4(10). <https://doi.org/10.1016/j.patter.2023.100857>.
- Else, H. (2023). Abstracts Written by ChatGPT Fool Scientists. *Nature*, 613, 7944, 423–423.

- Evans, N., Schmolmueller, A., Stolper, M., Inguaggiato, G., Hooghiemstra, A., Tokalic, R., Pizzolato, D., Foeger, N., Marušić, A., van Hoof, M., Lanzerath, D., Molewijk, B., Dierickx, K., & Widdershoven on, G. (2023). VIRT2UE: A European train-the-trainer programme for teaching research integrity. *Research Ethics*, 0(0). <https://doi.org/10.1177/17470161231161267>.
- Extance, A. (2023). ChatGPT Enters the Classroom. *Nature*, 623, 474-477. <https://doi.org/10.1038/d41586-023-03507-3>.
- Hine, C. (2021). Evaluating the Prospects for University-based Ethical Governance in Artificial Intelligence and Data-driven Innovation. *Research Ethics* 17/4, 464–479. <https://doi.org/10.1177/17470161211022790>.
- Hsu, J. (2023). Should schools ban AI chatbots?. *New Scientist*, 257(3422), 15. [https://doi.org/10.1016/S0262-4079\(23\)00099-4](https://doi.org/10.1016/S0262-4079(23)00099-4).
- Kasneji, E., Sessler, K., Küchemann, S., Bannert, M., et al. (2023). ChatGPT for good? On opportunities and challenges of large language models for education. *Learning and Individual Differences*, 103, 102274. <https://doi.org/10.1016/j.lindif.2023.102274>.
- Lucas, G. M., Gratch, J., King, A., & Morency, L.-P. (2014). It's only a computer: Virtual humans increase willingness to disclose. *Computers in Human Behavior*, 37, 94–100. <https://doi.org/10.1016/j.chb.2014.04.043>.
- McAdoo, T. (2023, April 7). *How to cite ChatGPT*. Retrieved November 28, 2023, from <https://apastyle.apa.org/blog/how-to-cite-chatgpt>.
- Miller, S. & Selgelid, M. J. (2007). Ethical and Philosophical Consideration of the Dual-use Dilemma in the Biological Sciences. *Science and Engineering Ethics*, 13(4), 523–580. <https://doi.org/10.1007/s11948-007-9043-4>
- Naddaf, M. (2023). ChatGPT generates fake data set to support scientific hypothesis. *Nature*, 623, 895-896. <https://doi.org/10.1038/d41586-023-03635-w>.
- O'Connor, S. & ChatGPT. (2023). Open artificial intelligence platforms in nursing education: Tools for academic progress or abuse? *Nurse Education in Practice*, 66, 103537. <https://doi.org/10.1016/j.nepr.2022.103537>.
- OpenAI. (2023). *ChatGPT* (November 28 version) [Large language model]. Retrieved November 28, 2023, from <https://chat.openai.com/chat>.
- Paul, J., Ueno, A., & Dennis, C. (2023). ChatGPT and consumers: Benefits, Pitfalls and Future Research Agenda. *International Journal of Consumer Studies*, 47(4), 1213-1225. <https://doi.org/10.1111/ijcs.12928>.
- Parijs, P. V. (2011). *Linguistic Justice for Europe and for the World*. Oxford University Press.
- Parrilla, J. M. (2023). AI Takes on Grant Applications, *Nature*, 623, 443, <https://doi.org/10.1038/d41586-023-03238-5>.
- Persson, I. & Savulescu, J. (2012). *Unfit for the Future: The Need for Moral Enhancement*. Oxford University Press.
- Pournaras, E. (2023). Science in the era of ChatGPT, large language models and generative AI Challenges for research ethics and how to respond. In A. Sudmann, A. Echterhölter, M. Ramsauer, F. Retkowski, J. Schröter, A. Waibel (Eds.), *Beyond Quantity. Research with Subsymbolic AI* (pp. 275-291).
- Rahman, M. M., Terano, H. J., Rahman, M. N., Salamzadeh, A., & Rahaman, M. S. (2023). ChatGPT and Academic Research: A Review and Recommendations Based on Practical Examples. *Journal of Education, Management and Development Studies*, 3(1), 1–12. <https://doi.org/10.52631/jemds.v3i1.175>.

- Roose, K. (2023, January 12). Don't Ban ChatGPT in Schools. Teach With It. *The New York Times*. Retrieved November 28, 2023, from <https://www.nytimes.com/2023/01/12/technology/chatgpt-schools-teachers.html>.
- Srikumar, M., Finlay, R., Abuhamad, G., Ashurst, C., Campbell, R., Campbell-Ratcliffe, E., Hongo, H., Jordan, S.R., Lindley, J., Ovadya, A. & Pineau, J. (2022) Advancing ethics review practices in AI research. *Nature Machine Intelligence*, 4(12), pp.1061-1064. <https://doi.org/10.1038/s42256-023-00608-6>.
- Taloni, A., Scorcio, V., & Giannaccare, G. (2023). Large Language Model Advanced Data Analysis Abuse to Create a Fake Data Set in Medical Research. *JAMA Ophthalmology*. doi:10.1001/jamaophthalmol.2023.5162.
- Tian, H., Lu, W., Li, T.O., Tang, X., Cheung, S-C., Klein, J., & Bissyandé, T.F. (2023). Is ChatGPT the Ultimate Programming Assistant -- How far is it?. *arXiv: 2304.11938*. <https://doi.org/10.48550/arXiv.2304.11938>.
- Uszkai, R. (2015). Global Justice and Research Ethics: Linguistic Justice and Intellectual Property. *Public Reason*, 7(1-2), 13-28.
- Uszkai, R. (2024). Fostering intellectual and ethical virtues in the age of Artificial Intelligence. The need for educators-in-the-loop. In E. Guarcello and A. Longo (Eds.), *School Children and the Challenge of Managing AI Technologies. Fostering a Critical Relationship through Aesthetic Experiences*. Routledge (forthcoming).
- Van Noorden, R. & Perkel, J.M. (2023). AI and science: what 1,600 researchers think. *Nature*, 621, 672-675. <https://doi.org/10.1038/d41586-023-02980-0>.
- Watkins, R. (2023). Guidance for researchers and peer-reviewers on the ethical use of Large Language Models (LLMs) in scientific research workflows. *AI and Ethics*, 1-6. <https://doi.org/10.1007/s43681-023-00294-5>.
- Zielinski, C., Winker, M., Aggarwal, R., Ferris, L., Heinemann, M., Lapeña, J.F., Pai, S. & Citrome, L. (2023). Chatbots, ChatGPT, and Scholarly Manuscripts-WAME Recommendations on ChatGPT and Chatbots in Relation to Scholarly Publications. *Afro-Egyptian Journal of Infectious and Endemic Diseases*, 13(1), pp.75-79. <https://dx.doi.org/10.21608/aeji.2023.282936>.