



GenAI and Academic Writing

Petar Jandrić^{1,2}

Received: 24 March 2026 / Revised: 24 March 2026 / Accepted: 26 March 2026 /
Published online: 29 April 2026
© The Author(s) 2026

Abstract

This paper exposes some of the main problems of cowriting with General Artificial Intelligence (GenAI) based on author's experience of editing Springer Nature's *Postdigital Science and Education* journal, Postdigital Science and Education Book Series, the *Encyclopaedia of Postdigital Science and Education*, and Chapters in Education. The paper outlines the four main challenges in facing scholarly editors—the challenge of detection, the challenge of sourcing, the challenge of writing, and the challenge of attitude. It reveals the opacity and practical uselessness of publisher guidelines, followed by some examples of methodologically and ethically sound human–GenAI cowriting. While presented conclusions are developed within the context of one publisher and its several publishing outlets, they are to an extent generalizable to contemporary academic publishing at least in the humanities and social sciences. The paper ends with an announcement of a zero-tolerance policy towards any kind of fraudulent human–GenAI cowriting in *Postdigital Science and Education* journal, Postdigital Science and Education Book Series, the *Encyclopaedia of Postdigital Science and Education*, and Chapters in Education, followed by an invitation for a broad postdigital dialogue about the future of academic writing with GenAI.

Keywords Postdigital · Writing · Generative Artificial Intelligence · GenAI · Policy · Plagiarism · Aigiarism

✉ Petar Jandrić
pjandric@tvz.hr

¹ Sungkyunkwan University, Seoul, Korea, Republic of

² Tehničko Veleučilište u Zagrebu, Zagreb, Croatia

Introduction

These days, it is my impression that 9 out of 10 articles submitted to Postdigital Science and Education (PDSE)¹ and Chapters in Education² are cowritten with some sort of General Artificial Intelligence (GenAI). While many researchers (should) understand that their writing needs to be adequately theorized, methodologically elaborated, and ethical, I have noticed a steady increase in submitted articles and chapters cowritten with GenAI that do not meet basic academic requirements. Similar trends have been reported by other scholarly editors, often as anecdotal as my impressions, and for now still largely appearing in grey literature such as blogs and social media posts.³ As of recently, some authors have started to call this phenomenon AIgiarism—an extended definition of plagiarism ‘to include representing AI generated work as one’s own’ (Drisko 2024). While the word sounds great and an umbrella term may ease my prose, I believe that the problem runs much deeper than representation so I will remain with older terminology.

When GenAI-related problems started to appear, I was flabbergasted. Who would expect, for instance, that a researcher of academic ethics would submit a paper that is obviously co-written with a GenAI, without an acknowledgement? This is not a figure of speech; the example refers to the real-life paper (Walle et al. 2025), littered with hallucinated references, published in Springer Nature’s *Journal of Academic Ethics*, and retracted a few months ago (see Retraction Watch 2025). While it would be easy to mock a journal of academic ethics that publishes obviously unethical material, I refuse to throw the first stone. As Erja Moore said for Retraction Watch (2025), ‘on one hand this is hilarious that an ethics journal publishes this, but on the other hand it seems that this is a much bigger problem in publishing and we can’t really trust scientific articles any more’.

My day-to-day practice seems to confirm Moore’s suspicions. A significant portion of suspicious material that overflows the PDSE and Chapters in Education submission systems does not arrive from paper mills or inexperienced researchers; it now routinely arrives from academics with solid track records, people whose careers exhibit years of good work, deep knowledge of their disciplines, and moral integrity. Those who submit such papers inevitably receive my email version of a howler—a bright red letter that angrily screams the message to the recipient, popularized by Ron Weasley’s parents in *Harry Potter and the Chamber of Secrets* (Rowling 1998). I should definitely work on improving my anger management and communication skills. To my defence, however, these howlers are my only hope that my message will get heard. Despite howlers and other desperate measures, the problem intensifies by the day.

¹ The umbrella term Postdigital Science and Education (PDSE) refers to the publishing ecosystem consisting of the *Postdigital Science and Education* journal, the Postdigital Science and Education Book Series, <https://link.springer.com/series/16439>, and the *Encyclopedia of Postdigital Science and Education*, <https://link.springer.com/referencework/10.1007/978-3-031-35469-4>. Accessed 24 March 2026.

² See <https://link.springer.com/series/60490>. Accessed 24 March 2026.

³ See, for instance, Ben Williamson’s account of similar issues in *Learning, Media and Technology*, <https://codeactsineducation.wordpress.com/>, accessed 24 March 2026.

I don't want to become known as the howling editor. I also refuse to dwell on my editor's high horse, flogging otherwise good people for their alleged sins. After all, cowriting with GenAI is in its infancy, and its proper use is still far from crystalized. However, I need to do something about all those paper submissions, and I need to do it now. I begin this paper with a brief outline of what I see as some of the main theoretical problems of cowriting with GenAI. I then expose the four main challenges in my day-to-day work as editor—the challenge of detection, the challenge of sourcing, the challenge of writing, and the challenge of attitude. I proceed with displaying the uselessness of publisher guidelines, followed by some examples of methodologically and ethically sound human-GenAI cowriting. I end the paper with an announcement of a zero-tolerance policy towards any kind of fraudulent human-GenAI cowriting in *Postdigital Science and Education* journal, Postdigital Science and Education Book Series, the *Encyclopaedia of Postdigital Science and Education*, and Chapters in Education.

What Seems to Be the Problem, Editor?

As Ingrid Forsler and I recently wrote elsewhere, 'the way we express knowledge fundamentally shapes the content and quality of that knowledge' (Jandrić and Forsler 2026). What we call scholarly knowledge is usually known by the concept of *episteme*, or true justified belief. Another concept from ancient Greece, *doxa*, describes common sense, belief, or wisdom, that does not necessarily need to be true. Truth does not depend on what its author is made of; both humans and GenAI produce a lot of *doxa* and some *episteme*. However, that does not mean that human and GenAI authors are equal—GenAI is structurally unable to distinguish between *doxa* and *episteme*, while (at least some) humans are (at least theoretically) able to do so (Traxler and Jandrić 2026). Since the beginning of Western civilisation, this ability has been the key to knowledge development.

The problem at hand, therefore, is not who produced knowledge, but how knowledge is assessed and verified. These have traditionally been the main duties of scholarly editors. While we are well versed in the assessment and verification of human writings, and while it seems reasonably easy to handle texts completely written by GenAI, the first problem in my work as editor are the many outcomes of human-GenAI cowriting and the associated interplay between *episteme* and *doxa*. How should such knowledge be assessed and verified?

The second problem is aesthetics. I edit because I love good writing—and I have yet to read a GenAI-written text that will sweep me off my feet. Structurally, all GenAI writing is one or other outcome of algorithmic processing of previously written words and a statistical calculation of a plausible sentence for a given context. Together with problems such as recursive repetitions of the same thought in different words, and beautifully sounding paragraphs that say next to nothing, I find GenAI writing just aesthetically unpleasing—sterile and boring, like those great-looking industrially grown tomatoes with no taste or smell (see Selwyn 2025; Molloy 2026). Aesthetics, in this context, is not just a matter of taste, as it directly links

to knowledge (Jandrić and Forsler 2026). But how can I formulate these aesthetic preferences, and how much weight should they have in my decision-making?

The third problem is that the political economy of scholarly work is closely related to publication. Publishing provides authors with various benefits from publicity through employment to cash payments. The incentive to publish more (cited) articles, that often prevails the incentive to publish better articles, pushes academic authors towards various shortcuts including cowriting with GenAI and towards unethical practices such as not acknowledging that usage. How can I even approach the text, when I do not know who wrote it, and when I can reasonably suspect that the author will try and hide cowriting with GenAI?

These problems, alongside many unmentioned problems such as the general impact of GenAI on knowledge development (Xiao et al. 2025; Traxler and Jandrić 2026), the poisonous role of fraudulent knowledge in postdigital colonial relationships (Anamika et al. 2026), conflict and war (Green et al. 2026), the environmental impacts of GenAI (Jandrić et al. 2026), and so on, act in a dynamic interplay. Collectively, we currently just do not understand many aspects of that interplay. We need to continue exploring pros and cons of cowriting with GenAI, anchored in older research in related topics such as the postdigital backlash (Forsler et al. 2025) and general studies of technology acceptance (Granić and Marangunić 2019), hoping to develop a (more) comprehensive solution sometime in the future.

Insights and policies of the future are shaped and configured in the research practice of today (Crain et al. 2026). It is hugely important to shape today's research practice for the future we would like to live in! This shaping needs to be open for everyone to analyse and critique, and this paper is just one small step towards these goals. As I wrote years ago in a slightly different context, '[w]hile the presented life cycle of a double-blind peer-reviewed scholarly article is based on the example of *Postdigital Science and Education*, these practices are fairly standard for journals across the humanities and social sciences and may be of interest to scholars in diverse fields and disciplines' (Jandrić 2020: 37). Conclusions and policies presented in this article are also deeply interlinked with the context of their making, yet they are to an extent generalizable as well.

Practical Challenges

In my experience, using GenAI in academic writing yields 4 main classes of practical challenges: the challenge of detection, the challenge of sourcing, the challenge of writing, and the challenge of attitude.

The Challenge of Detection

Only in the past year or so, I have experienced attempts at hiding GenAI usage by all sorts of writers. Scholars from early researchers to accomplished academics, some of them close friends, have submitted articles cowritten with GenAI without any acknowledgement. For more times than I want to even remember, I uttered in disbelief:

Et tu, Brute? Sadly, the first lesson from my practice sounds like a line from *Breaking Bad* (Gilligan 2008–2013): trust no one, including your closest friends and family.

This brings forward the problem of detection. I have developed a decent editors' hunch about unacknowledged GenAI usage (see Veletsianos et al. 2025)—but that hunch is far from perfect, often kicks in late, and, in some cases, it just fails. Machines are not of much help either, as available GenAI plagiarism detectors are far from accurate. Ridiculous results offered by some GenAI plagiarism detectors such as the claim that the Declaration of Independence is GenAI-generated now seem to be addressed (Rao 2025), yet those detectors are still very far from an ability to accurately assess complex academic language. So far, the winning formula consists of one or another combination of automated plagiarism detection with human insight. While it seems fairly accurate so far, this formula suffers from two main problems.

First, fraudulent use of GenAI cannot be 'proven' in the same way that we can prove traditional plagiarism. While some scholars consciously misrepresent their use of GenAI, others might even not realize themselves if something is cowritten with GenAI, or at least not to what extent, as many universities now have GenAI assistants integrated in word processing software and the border between cowriting and word processing is blurring. This sometimes results in long, uncomfortable, and unproductive discussions, and in some cases, there is also a tiny glimmer of doubt: what if some of these people are indeed right?

Second, accurate detection takes a lot of time. Combined with the vast increase in the number of submissions, decision-making based on extensive human labour is just not sustainable—not to mention underlying injustice in that dishonest submissions tend to get much more (undeserved!) editors' time and attention than honest ones. Third, when I do detect unacknowledged use of GenAI, I can never be sure that I found out everything. In their responses, authors often make claims such as 'ideas are completely mine; I used GenAI only for referencing... or grammar... or style... or whatever', and I can never be sure whether they purposefully misrepresent or are truly unaware of the nature of their actions. In practical terms, however, an author's understanding of their actions makes no difference to the reader.

The Challenge of Sourcing

GenAI often 'proves' its claims using made-up or hallucinated sources (a popular word that should be avoided for its anthropomorphism). One blatant example is the aforementioned paper in *Journal of Academic Ethics* where the whistleblower analysed 'all of the paper's 29 references and found at least 19 of them appear to be fabricated' (Retraction Watch 2025). GenAI is getting better, so these days most papers come with a lesser percent of fabricated references, and the fabricated references get less obvious; leading to the need for manual cross-checking of each and every reference in each and every submitted paper. However tedious, this detection is only the first part of the problem—the real question, of course, is what should be done when an otherwise decent-looking paper has some fabricated references?

My policy is simple: verifiable arguments cannot be based on unverifiable sources, so such papers are immediately rejected. Surprisingly, authors often object,

saying things such as ‘the argument is mine, I only asked GenAI to help me with the boring work of sorting references’. Well, your bad. Garbage in–garbage out. Sloppy referencing implies sloppy thinking; take care of the pennies and the pounds will take care of themselves.

The Challenge of Writing

At the level of sentence, GenAI writing is almost impossible to detect. Looking at bigger chunks of text, however, GenAI tends to produce texts that sound convincing but convey little to no meaning; circular arguments that do not offer any new information; repetitions (in many versions including completely different paragraphs that convey the same meaning); and many other writing *faux pas* (see Abdul-Jabbar and Bhatt 2025). The depth and extent of these issues depend on the topic, the GenAI system used, and the quality of prompting. Combined with human intervention, these problems can be flattened to the point where some articles may look fairly convincing. Some of these papers may even bring a new quality and pass as original research!

In an early example, PDSE published the paper ‘Theodor W. Adorno, Artificial Intelligence, and Democracy in the Postdigital Era’ (Park 2024). I edited the text myself, and through several feedback iterations, I decided that the paper is worth publishing. A year later, I received a whistleblower email saying:

I believe an article published in Postdigital Science and Education contains plagiarism ... This author already has one retracted article (Park 2025). In that case the evidence was based on fabricated references. That does not appear to be the issue here. Instead, the text is rife with superficial paraphrasings that is common to the way genAI ‘paraphrases’ text. Even if not GenAI, the level of superficial paraphrasing would constitute mosaic plagiarism.⁴

I checked the article thoroughly, and I realized that the whistleblower is correct. The article has been cowritten between a GenAI and the human author, yet many typical GenAI problems such as superficial paraphrasing and circular argument have remained.

What should I do with this article? Paraphrasing is standard academic practice that predates GenAI; if all references are in place (and they are), it has never been sanctionable. All sources exist and are used correctly. While it is obvious that the author wrote with GenAI, this cannot be proven empirically. Publishing this article was clearly my mistake, that can be justified only by the fact that this was one of the first GenAI articles submitted to PDSE, and I just had too little experience. Having said that, was publishing this paper really a mistake? After all, the paper has passed human peer review has already been read and cited by other authors; perhaps, this time, the human-GenAI collaboration has managed to produce some clean *episteme*.

Author’s unacknowledged use of GenAI is deeply problematic. But I cannot prove that use, and I cannot sanction the author based on hunch, so I just have no ground for retracting this article. With my current experience, I can reasonably safely vouch that such a thing will not happen again. But technology develops, new

⁴ Personal email correspondence, lightly edited for clarity.

uses of GenAI appear by the day, and I really cannot say that some other author, using some other strategy, will not trick me again. The ivory tower of academic publishing leaks like a sieve – and no-one can say what kind of damage leaked water will make in the long run.

The Challenge of Attitude

Different writing practices bring about different ways of thinking and different kinds of knowledge. I will refrain from delving into deep epistemological questions (see Abdul-Jabbar and Bhatt 2025 for a brilliant take on these) and will only cover, very superficially, some changes in authors' attitude about academic writing. Judging from my correspondence with approximately 1000 authors per year, an increasing number of authors either plead ignorance or genuinely believe that their practice is unproblematic.

In a paradigmatic example, the author submitted an article that had some great ideas yet was clearly cowritten with GenAI; the use of GenAI was not acknowledged. I responded with a long email—explaining that I like the ideas and that I can clearly see that he invested a lot of effort in thinking, but the paper suffers from these and that GenAI writing problems, and, of course, that the use of GenAI should be acknowledged. Since I felt that only about half of the text was GenAI generated, while the other half was written by the author, I took a couple of hours to write extensive feedback, and I asked the author to address this feedback—by himself, or with an appropriately acknowledged use of GenAI. The author responded with an equally long email, thanking me for my work, addressing my feedback point by point, and promising that they'll rewrite problematic parts themselves.

The second version of the article soon came back; text edited, GenAI unacknowledged. The article still looked weird, so I spent another couple of hours to prepare another round of feedback. Then the third version of the article arrived, and I started to work on the text again... After about an hour of editing, the article just did not feel right, so I checked it for GenAI plagiarism. What I found out, was that all three versions of the article, including two point-to-point responses to my feedback, were written by a GenAI! I felt really bad. This person literally wrote, 'I wrote this myself', and then he submitted a paper that was written by a GenAI! I had no reason to distrust him, so I wasted 5 or 6 hours of labour fighting windmills. When I emailed a howler to the author, saying that his actions are unethical and personally demeaning, he just wrote back, 'Thanks a lot for your effort', followed with a smiley face emoji. Recalling this exchange months later, I cannot help but think of several nouns and adjectives for this person, none of which are suitable for an academic journal.

After this experience, I would not trust this author to make me a cappuccino – and don't get me even started on what I think about his moral integrity. I now have a blacklist of known PDSE and Chapters in Education cheaters and liars. Sadly, my blacklist grows rapidly; sadder, it now contains some people I used to trust and even consider friends. What could be the reason for such behaviour? Do those people genuinely believe that unacknowledged cowriting with GenAI is ethical? Do they perhaps use my feedback as their Turing testbed, probing the

limits of what can go undetected? Do they not see that their dishonesty transforms my editing into policing? And do they not realize how miserable I must feel when lied to my face?

This issue is charged emotionally and appears beyond academic editing—I have seen senior professors with tears in their eyes when recalling how they spent hours giving feedback on something that later turned out to be GenAI-generated. Yet so many people seem to feel that such actions are permissible; perhaps there is a terrain of different moral understanding that I do not see. Whatever the future verdict of history, something needs to be done now. Unsurprisingly, publishers have already responded with guidelines and regulations.

Publisher Guidance

This paper is based on my experience of editing PDSE and Chapters in Education, so I will examine one guidance document by their publisher, Springer Nature. While I seriously doubt that other mainstream publishers and documents will be radically different, I do allow that a different object of analysis may lead to different conclusions. In its ‘Guidance for researchers on use of AI in publishing’, Springer Nature (2026) (Table 1) recommends:

Table 1 Guidance for researchers on use of AI in publishing (Springer Nature 2026)

Guidance for authors
• Be accountable. You—not AI—are responsible for the accuracy, originality, and integrity of your manuscript. • Do not list AI tools as authors. LLMs such as ChatGPT cannot take accountability and do not meet authorship criteria. • Declare AI use transparently. If an AI tool helped generate text, data analysis, or content, acknowledge it in the Introduction, Preface, or Acknowledgements. • AI-assisted copy editing need not be declared when used only to improve readability, grammar, or formatting—but authors remain accountable for the final version. • Do not use AI to fabricate, plagiarise, or create false content. Verify and reference any AI-assisted output. • Avoid generative AI images or figures. These are not permitted for publication unless they meet specific exceptions (legally sourced, directly about AI, or based on verifiable scientific data) and are clearly labelled as ‘AI-generated.’ • Craft prompts responsibly. Do not include personal, sensitive, confidential, or copyrighted material. Avoid biased or harmful prompts and test outputs for appropriateness. • Follow Springer Nature’s Authors’ Code of Conduct and ensure you uphold your research integrity.
Guidance for editors and peer reviewers
• Take authorship and accountability for all editorial decisions and review content—these must reflect your expert judgment. • Do not upload manuscripts into generative AI tools. Manuscripts contain confidential and sensitive information. • If you have used an AI tool to assist your evaluation, declare its use transparently in your review report. • Maintain confidentiality and integrity throughout the peer-review process. • Keep a human in the loop. Editorial and peer-review assessments must be made and verified by humans. (Springer Nature 2026)

Some author guidelines, such as ‘You—not AI—are responsible for the accuracy, originality, and integrity of your manuscript’ and ‘Do not list AI tools as authors’, seem like no-brainers. Others are much more problematic; some to the point of uselessness. For instance, the guideline that ‘AI-assisted copy editing need not be declared when used only to improve readability, grammar, or formatting—but authors remain accountable for the final version’, is just unenforceable in practice, as it is impossible to discern between the use of GenAI to improve readability, grammar, or formatting, and the use of GenAI to develop argument.

When an author ticks a box saying that they used GenAI only for readability, grammar, or formatting, the publisher and the editor are absolved from responsibility. While virtue signalling and legal protection may mean a lot in the corporate world of academic publishing, people do misrepresent their use of GenAI, and the system allows them to do so without consequence. Guidelines such as ‘Craft prompts responsibly’ are similarly useless, as GenAI responses to prompts are not reproducible, responsibility means different things to different people, and one or two sentences explaining what Springer Nature sees as responsibility are at best unclear. Therefore, I cannot help but see these performative box-ticking practices as cynical attempts at moral and legal whitewashing with little or no epistemic consequence.

Guidance for editors and peer reviewers is just as problematic. In my favourite example, the second guideline, ‘Do not upload manuscripts into generative AI tools’, and the third guideline, ‘If you have used an AI tool to assist your evaluation, declare its use transparently’, directly collide. First, does Springer really expect its editors to fight the machine guns of GenAI-produced content with bows and arrows of sole human judgement? Second, if I am really not supposed to use GenAI tools such as plagiarism detectors in my work – why are you asking me, literally in the next sentence, to declare their usage for the purpose? Written by people whose bread and butter is academic publishing (if the authors of the guidance are even human!), it is hard to believe that this obfuscation is accidental. So what am I exactly supposed to do this afternoon, when I will be deciding on the fate of articles submitted to PDSE and Chapters in Education since yesterday?

Without further ado, this short analysis confirms my practice-informed suspicions. ‘Guidance for researchers on use of AI in publishing’ (Springer Nature 2026) offers moral and legal protection of publisher’s business. However, it can be interpreted in too many contradicting ways and bears little if any epistemic consequence. For those of us who are here for the love of *episteme*, such guidance is largely useless.

In People We Trust

The easiest solution to these problems is to ban all GenAI writing. In the context of postdigital research, however, this would be just hypocritical. As we explore entanglements of human-nonhuman agencies in theory, we just need to walk our own talk and try them out in practice. In my inaugural editorial for *Postdigital Science and Education*, published now more than 8 years ago, I wrote: ‘I would gladly accept an article written by a cat, or by an Artificial Intelligence, if it holds merit.’ (Jandrić

2019: 2). I still hold on to these words, but now I feel the full weight of the phrase ‘if it holds merit’. So what does it mean to hold merit?

Speaking of research design, there are many interesting examples where collaboration between GenAI and human authors has produced new knowledge. In one of the earliest examples, Michael A. Peters (2023) wrote a Foreword for *Post-digital Research: Genealogies, Challenges, and Future Perspectives* by interviewing ChatGPT about the book’s topic. While this approach is the lowest-hanging fruit of GenAI experimentation, it was also amongst the first. Peters wrote the piece less than 2 months after ChatGPT appeared in public, and his approach has been debated widely – making a timely if ephemeral contribution to knowledge development.

One article at a time, authors have expanded such experimentation in various directions. In ‘The “Tragic Dilemma” of GenAI in Education: An Exploratory Human-AI Conversation’, Agnieszka Palalas and colleagues (2025) have pushed Peters’ dialogue with GenAI horizontally (by adding more interlocutors) and vertically (stretching their dialogue to 3 months). Others have turned their attention to GenAI image generators. One of my favourite articles in this style, Carlos Escaño’s ‘*Die Mensch-Maschine: Visual Dialogue Between a Human and an Artificial Intelligence*’ (2024), presents a visual inquiry into in-built GenAI ideology; similarly, Jasper Roe (2025) has collaborated with GenAI to explore ‘Generative AI as Cultural Artifact: Applying Anthropological Methods to AI Literacy’.

In another popular genre of scholarship, authors explore the many ways to use GenAI in teaching and learning. ‘Schooling the (Artificially) Intelligent Writer: GenAI Technologies as Thinking Infrastructures in School Education’ by Toon Tierens and colleagues (2025) is an especially powerful example, not the least because it has provoked vivid debate. Niñoval Pacaol’s (2025) ‘Doubting AI and Infrastructure Literacy: Response to “Schooling the (Artificially) Intelligent Writer: GenAI Technologies as Thinking Infrastructures in School Education” (Tierens et al. 2025)’ critiques and expands Tierens et al.’s (2025) argument, and Tierens (2026) response to Pacaol’s (2025) critique. The dialogue continues.

PDSE and Chapters in Education have a strong focus on (educational) futures (Macgilchrist et al. 2024; Crain et al. 2026; Jandrić 2026), and this scholarship tends to experiment with GenAI in more speculative ways. A nice example of this genre is Sara Pastore’s “‘The Classroom Robot”, by ChatGPT: Unpacking the Imaginary of the AIED Classroom Through an AI-Generated Story’ (2025), where a fictional GenAI-generated text is used for academic purposes. Similarly, in a wonderful piece of human-written academic fiction, Mark Curcher describes a hypothetical future in which students write their essays with GenAI and teachers mark the essays with GenAI, with little or no human thinking involved. The protagonist bitterly concludes:

‘Lots of learning going on in this process, machine learning that is,’ he thought sardonically, ‘hypnotising chickens!’

There are no losers in this. The students get what they want, the staff made a living, the University meets all its metrics, and the Ministry of Education, Enterprise and Innovation keeps up in the international league tables for

the number of graduates. Everyone is a success, their potential unlimited. A whole lot of qualifications and not a whole lot of education, but everyone in the game is a winner! Outside the sun set in the purple sky, turning it the colour of a dead TV channel. (Curcher 2023: 559).

Perhaps, as Stefan Hrastinski and I wrote in our editorial for *PDSE* 5(3) (Hrastinski and Jandrić 2023), we do need fictional stories like this to expose paradoxical future consequences of our present actions!

Finally, there are papers that do not speak about GenAI, but are written with the help of GenAI. One recent example includes a paper on academic peer review that I coauthored with George Veletsianos and his students, ‘Postdigital Peer Review: A Vulnerable Space Between Freedom and Responsibility’ (Veletsianos et al. 2025). While I never use GenAI in my own writing, students felt that they needed some help with the language, so they used GenAI to refine their parts of the text. Fine with me! We added a brief acknowledgement at the end: ‘The paper was reviewed, edited, and refined with the assistance of Gemini Pro 2.5 as of September 2025, complementing the human editorial process to address grammar, flow, and style.’ That’s all, folks! – the paper is ready to go.

This brief and incomplete overview of recent postdigital research clearly shows that papers cowritten with GenAI can yield solid academic knowledge. GenAI can offer useful suggestions, yet for as long as it is structurally unable to distinguish between *doxa* and *episteme*, judgement of GenAI suggestions remains with humans. The key issue here is trust – and an important aspect of this trust is that editors and readers just need to know who wrote which part of the text (who created an image, and so on) and why. However, this trust is now deeply eroded, as the simple act of asking, ‘Has this been written by a GenAI?’ is grounded in distrust.

As I wrote ages ago in the context of post-truth, we could try and reinvent Freirean critical pedagogy of trust (Jandrić 2018a, b). This reinvention takes time, and social problems such as trust are at least partially unsusceptible to educationalisation (see Peters et al. 2019). Alongside pedagogy, we therefore need some sort of regulation to keep rotten apples out of scholarly discourse. Given the uselessness of publisher guidance, the best thing I can do is sweep my own doorstep and develop PDSE-specific and Chapters-specific policy.

Conclusion

A well-put assemblage between a human and GenAI, and an adequate and ethical approach to their entanglement, should work like a bicycle – a technology that allows us to travel faster, and further, while maintaining health benefits and environmental sustainability (see Illich 1974). Obviously, bicycles are not for all people, cannot be driven in all terrains, and allow much lesser reach than a car or an aeroplane; to get the best of bicycles, we need to understand their advantages and limitations. If used right, GenAI also can help us achieve things that would be

difficult, if not impossible, to achieve otherwise. To get the best of GenAI, however, we need to understand what it can be used for, and the limits of its usage.

As GenAI articles proliferate, editing becomes less and less focused on concepts and ideas and more and more focused on detection and management of fraudulent behaviour. I am sick and tired of this, as are many other editors. But we cannot abandon ship, as we are the last keepers of humankind's sanity. If we open the door to automated knowledge-makers that cannot distinguish between *episteme* and *doxa*, we will enter a world that most of us would just not like to inhabit. GenAI seems here to stay, so we remain in the trenches.

In a recent editorial, I wrote: 'While we strive very hard to develop better scholarship, we strive even harder to develop better scholars.' (Jandrić 2026) This paper, for all it is worth, is an effort to that end. I use this brutally honest account of my editing experience to develop insight that will hopefully be useful within and beyond my immediate context. I try and develop a critical pedagogy of trust between authors and editors. I want to raise critical consciousness about cowriting with GenAI, and I especially want to reach those good people who may have gone astray for various reasons such as lack of knowledge or time. Arguably most importantly, and this was my original motivation for writing this article, I desperately need to set out some actionable conclusions.

As I write these words, it has been more than three years since ChatGPT opened the floodgates of GenAI cowriting. Given the amount of time and effort required for detection and management of fraudulent GenAI practice, I just need to announce a zero-tolerance policy. Any attempt at improper and especially unacknowledged use of GenAI in submissions to *Postdigital Science and Education* journal, *Postdigital Science and Education* Book Series, the *Encyclopaedia of Postdigital Science and Education*, and *Chapters in Education*, will be immediately sanctioned: submitted work will be rejected, and the author will be black-listed. Excuses such as 'I thought that ChatGPT would sort my references like Zotero, had no idea that it can make up references' or 'I used ChatGPT only to improve my grammar, so I didn't feel the need to acknowledge that usage', just don't fit the bill anymore. While I don't believe that once a cheater will always be a cheater, I just refuse to play cop games.

Ancient Roman Law gave us an important principle that persists in worldwide legal systems to this day: ignorance of the law does not excuse from responsibility (*ignorantia juris non excusat*). Reckless driving is an offence, so is reckless usage of GenAI – and from now on, in the PDSE and *Chapters in Education* publishing universe, it will be unanimously treated as such. To anyone who wants to produce decent scholarship, the zero-tolerance policy towards improper and unacknowledged use of GenAI is not a threat but an important benefit. This policy is a filter that allows us to clearly distinguish between *episteme* and *doxa*, and a mechanism that makes PDSE and *Chapters in Education* a verifiable and respectable publishing ecosystem, ensuring that its authors' work counts.

Technologies advance, circumstances change, and exact solutions are impossible to prescribe. Therefore, this policy is also an invitation for a respectful and productive postdigital dialogue (Jandrić 2019) about the future of academic publishing. PDSE and *Chapters in Education* will continue to produce highest quality research

that advances our knowledge of human-GenAI cowriting and collaboration. In the meantime, if you are unsure of any aspect of GenAI use in your scholarship, ask your editor! We will discuss the matter, seek further help if needed, and come up with an ethical solution that will make your research shine.

Acknowledgements Ideally, policy of GenAI use should result from wide community consensus. This paper is a mere patch to the burning problem of rapid increase in unacknowledged GenAI submissions, and presented policy should hopefully serve the community until a more permanent solution is developed. Many thanks to James Lamb (co-editor of the PDSE book series), Ingrid Forsler (editor of the Postdigital Dialogues section of *PDSE*), and Mia Čarapina (co-editor of Chapters in Education) for their invaluable comments and contributions to this paper. Input and approval from fellow editors is the policy's reality check coming straight from practice, and an important factor of harmonization between different publishing outlets. Furthermore, many thanks to Sarah Hayes and Jeremy Knox; co-founders of *PDSE*, who have been developing postdigital scholarship and community for now more than a decade, and whose input and approval is the policy's reality check that links our past and present. A publishing ecosystem that has such a team does not need to worry about its future!

Author Contributions This paper is proudly and completely written by a human being. P.J. wrote the whole paper on his own.

Funding Open Access funding enabled and organized by SungKyunKwan University

Data Availability No datasets were generated or analysed during the current study.

Declarations

Competing interests The authors declare no competing interests. The author is the Editor-in-Chief of the journal *Postdigital Science and Education*. He was not involved in the review process.

Open Access This article is licensed under a Creative Commons Attribution-NonCommercial-NoDerivatives 4.0 International License, which permits any non-commercial use, sharing, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons licence, and indicate if you modified the licensed material. You do not have permission under this licence to share adapted material derived from this article or parts of it. The images or other third party material in this article are included in the article's Creative Commons licence, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons licence and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder. To view a copy of this licence, visit <http://creativecommons.org/licenses/by-nc-nd/4.0/>.

References

- Abdul-Jabbar, W. K., & Bhatt, I. (2025). Post-authorship and AI-Assisted Writing: Singularity, Communication, and Pedagogy in Higher Education. *Postdigital Science and Education*, 7(4), 1136–1149. <https://doi.org/10.1007/s42438-025-00586-5>.
- Anamika, Conradi, T., & Jandrić, P. (2026). Media, Knowledge, and Learning: Decolonizing Educational Technology?. *Journal of Educational Media, Memory, and Society*.
- Crain, D. E., Jandrić, P., Peters, M. A., Besley, T., Davis, H., Green, B. J., Kamenarac, O., & Suoranta, J. (2026). Preconfiguring Postdigital Futures: Educational Development for Sustainability and Justice. *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-026-00627-7>.
- Curcher, M. (2023). The Pseudo Uni. *Postdigital Science and Education*, 5(3), 558–560. <https://doi.org/10.1007/s42438-022-00384-3>.

- Drisko, J. W. (2024). AIgiarism: Computer Generated Text, Plagiarism, and How to Address it in Teaching. *Journal of Teaching in Social Work*, 45(1), 1–15. <https://doi.org/10.1080/08841233.2024.2433795>.
- Escaño, C. (2024). *Die Mensch-Maschine*: Visual Dialogue Between a Human and an Artificial Intelligence. *Postdigital Science and Education*, 6(4), 1055–1081. <https://doi.org/10.1007/s42438-024-00518-9>.
- Forsler, I., Hayes, S., Jandrić, P., Macgilchrist, F., Selwyn, N., & Hansén, S. (2025). Hijacking the Digital Backlash in Education. *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-025-00601-9>.
- Gilligan, V. (2008–2013). *Breaking Bad* [TV Series]. High Bridge Entertainment, Gran Via Productions, and Sony Pictures Television.
- Granić, A., & Marangunic, N. (2019). Technology acceptance model in educational context: A systematic literature review. *British Journal of Educational Technology*, 50(5), 2572–2593. <https://doi.org/10.1111/bjet.12864>.
- Green, B. J., Kamenarac, O., & Jandrić, P. (2026). Postdigital War and Peace. *Postdigital Science and Education*.
- Hrastinski, S., & Jandrić, P. (2023). Imagining Education Futures: Researchers as Fiction Authors. *Postdigital Science and Education*, 5(3), 509–515. <https://doi.org/10.1007/s42438-023-00403-x>.
- Illich, I. (1974). *Energy and equity*. London: Marion Boyars.
- Jandrić, P. (2018a). Post-truth and critical pedagogy of trust. In M. A. Peters, S. Rider, M. Hyvönen & Tina Besley (Eds.), *Post-Truth, Fake News: Viral Modernity & Higher Education* (pp. 101–111). Singapore: Springer. https://doi.org/10.1007/978-981-10-8013-5_8.
- Jandrić, P. (2018b). Welcome to Postdigital Science and Education!. *Postdigital Science and Education*, 1(1), 1–3. <https://doi.org/10.1007/s42438-018-0013-8>.
- Jandrić, P. (2019). Welcome to postdigital science and education!. *Postdigital Science and Education*, 1(1), 1–3. <https://doi.org/10.1007/s42438-018-0013-8>.
- Jandrić, P. (2020). A Peer-Reviewed Scholarly Article. *Postdigital Science and Education*, 3(1), 36–47. <https://doi.org/10.1007/s42438-020-00202-8>.
- Jandrić, P. (2026). Postdigital Research: Current Issues and Future Directions. *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-026-00637-5>.
- Jandrić, P., & Forsler, I. (2026). Knowledge and Writing: Tinkering with The Scholarly Article. *Postdigital Science and Education*.
- Jandrić, P., Knox, J., Rapanta, C., Hayes, S., Tolbert, S., Kostakis, V., Misiasek, G. W., & Lee, K. (2026). EdTech and the Environment: A Research Program. *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-026-00635-7>.
- Macgilchrist, F., Allert, H., Cerratto Pargman, T., & Jarke, J. (2024). Designing Postdigital Futures: Which Designs? Whose Futures?. *Postdigital Science Education*, 6(1), 13–24. <https://doi.org/10.1007/s42438-022-00389-y>.
- Molloy, K. (2026). Review of Neil Selwyn (2025). *Digital Degrowth: Radically Rethinking Our Digital Futures*. *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-025-00621-5>.
- Pacaol, N. (2025). Doubting AI and Infrastructure Literacy: Response to ‘Schooling the (Artificially) Intelligent Writer: GenAI Technologies as Thinking Infrastructures in School Education’ (Tierens et al. 2025). *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-026-00626-8>.
- Palalas, A., Pegrum, M., Gunawardena, C. N., Robertson, R. S., & Ventura, J. (2025). The ‘Tragic Dilemma’ of GenAI in Education: An Exploratory Human-AI Conversation. *Postdigital Science and Education*, 7(4), 1451–1477. <https://doi.org/10.1007/s42438-025-00600-w>.
- Park, S. (2024). Theodor W. Adorno, Artificial Intelligence, and Democracy in the Postdigital Era. *Postdigital Science and Education*, 6(4), 1287–1303. <https://doi.org/10.1007/s42438-023-00424-6>.
- Pastore, S. (2025). ‘The Classroom Robot’, by ChatGPT: Unpacking the Imaginary of the AIEd Classroom Through an AI-Generated Story. *Postdigital Science and Education*, 7(3), 699–718. <https://doi.org/10.1007/s42438-025-00564-x>.
- Peters, M. A. (2023). In Search of The Postdigital: A Conversation with ChatGPT. In P. Jandrić, A. MacKenzie, & J. Knox (Eds.), *Postdigital Research: Genealogies, Challenges, and Future Perspectives* (pp. ix–xiv). Cham: Springer.
- Peters, M. A., Jandrić, P., & Hayes, S. (2019). The curious promise of educationalising technological unemployment: What can places of learning really do about the future of work? *Educational Philosophy and Theory*, 51(3), 242–254. <https://doi.org/10.1080/00131857.2018.1439376>.

- Rao, T. (2025). The AI Said the U.S Declaration of Independence Was AI-Generated. Medium, 3 June. <https://medium.com/illumination/the-ai-said-the-u-s-declaration-of-independence-was-ai-generated-85b9b6b853df>. Accessed 13 March 2026.
- Retraction Watch. (2025). The case of the fake references in an ethics journal. <https://retractionwatch.com/2025/12/02/fake-references-chatgpt-journal-academic-ethics-springer-nature-whistleblowing/>. Accessed 13 March 2026.
- Roe, J. (2025). Generative AI as Cultural Artifact: Applying Anthropological Methods to AI Literacy. *Postdigital Science and Education*, 7(4), 1107–1124. <https://doi.org/10.1007/s42438-025-00547-y>.
- Rowling, J. K. (1998). *Harry Potter and the Chamber of Secrets*. London: Bloomsbury.
- Selwyn, N. (2025). *Digital Degrowth: Radically Rethinking Our Digital Futures*. Cambridge, UK: Polity Press.
- Springer Nature. (2026). Guidance for researchers on use of AI in publishing. <https://group.springernature.com/gp/group/ai-ai-guidance-for-our-researchers-and-communities>. Accessed 24 March 2026.
- Tierens, T., Decuyper, M., Hartong, S., & Alirezabeigi, S. (2025). Schooling the (Artificially) Intelligent Writer: GenAI Technologies as Thinking Infrastructures in School Education. *Postdigital Science and Education*, 7(4), 1250–1269. <https://doi.org/10.1007/s42438-025-00585-6>.
- Tierens, T. (2026). Literate Attachments: A Response to ‘Doubting AI and Infrastructure Literacy’ (Pac-aol 2026). *Postdigital Science and Education*.
- Traxler, J., & Jandrić, P. (2026). GenAI and Why It’s Impossible to Learn or Understand Language. In A. J. Means, & K. J. Saltman, (Eds.), (2026). *Critical Studies of AI in Education*. New York: Cambridge University Press.
- Veletsianos, G., Arabadzy, G., Athilat, V., Bartucz, J., Dogra, S., Fredrickson, N. R., Goeke, M. Jeon, S. T., Urena, M., Jandrić, P., Green, B. J., & Jackson, L. (2025). Postdigital Peer Review: A Vulnerable Space Between Freedom and Responsibility. *Postdigital Science and Education*. <https://doi.org/10.1007/s42438-025-00604-6>.
- Walle, Y. M., Gedefaw, S. T., & Bezabih, H. A. (2025). RETRACTED ARTICLE: Whistleblowing in Ethiopian Public Educational Institutions: Perspectives of Individuals with Disabilities. *Journal of Academic Ethics*, 23, 1829–1846. <https://doi.org/10.1007/s10805-025-09630-2>.
- Xiao, J., Bozkurt, A., Nichols, M., Pazurek, A., Stracke, C. M., Bai, J. Y. H., Farrow, R., Mulligan, D., Nerantzi, C., Sharma, R. C., Singh, L., Frumin, I., Swindell, A., Honeychurch, S., Bond, M., Dron, J., Moore, S., Leng, J., Slagter van Tryon, P. J., Garcia, M., Terentev, E., Tlili, A., Chiu, T. K. F., Hodges, C. B., Jandrić, P., Sidorkin, A., Crompton, H., Hrastinski, S., Koutropoulos, A., Cukurova, M., Shea, P., Watson, S., Zhang, K., Lee, K., Costello, E., Sharples, M., Vorochkov, A., Alexander, B., Bali, M., Moore, R. L., Zawacki-Richter, O., Asino, T. T., Huijser, H., Zheng, C., Sani-Bozkurt, S., Duart, J. M., & Themeli, C. (2025). Venturing into the Unknown: Critical Insights into Grey Areas and Pioneering Future Directions in Educational Generative AI Research. *TechTrends*, 69(3), 582–597. <https://doi.org/10.1007/s11528-025-01060-6>.

Publisher’s Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.