

Peer Review

Review of: "The Role of Artificial Intelligence in the Lifecycle of Scientific Manuscripts: Authoring, Reviewing, and Editorial Selection"

Nadine Schlicker¹

1. Phillips-Universität Marburg, Germany

1. Summary: The paper "The Role of Artificial Intelligence in the Lifecycle of Scientific Manuscripts: Authoring, Reviewing, and Editorial Selection" provides an interesting starting point to consider and discuss the current problems of peer review and how they are amplified by AI-(supported) peer review. The solutions suggested in the provided framework are overall valuable. I have some minor recommendations that may help to further improve the manuscript.
2. Introduction: In the introduction, it is claimed that "50% used AI tools when peer reviewing manuscripts". At this point, it may be valuable to clarify which parts of peer review LLMs are used for. Especially considering the later part of the manuscript, it should be clarified what exactly the problem of using LLMs in peer review is. So, for instance, is it a problem if reviewers try to polish their language to make their recommendations clearer? Or if reviewers ask an LLM whether their comments are helpful to the authors? It may help to specify here exactly which parts of the reviewing process the 50% refer to and to contextualise this before speaking of an "undisciplined adoption", as some reviewers may indeed spend more time reviewing when only editing their language with an LLM. And in this case, I think there would also be no "paradox"; AI could also "alleviate burdens like linguistic inequality" in this regard.
3. Additional References: When reading the manuscript, this paper came to mind: https://www.nature.com/articles/s42256-025-01159-8?error=cookies_not_supported&code=2f1082e4-2800-4251-bc52-aa1d1fbb17e9. It addresses the attribution problem, i.e., that authors are not at all credited for their original work and hence do not receive intellectual credit when LLMs summarise,

rephrase, and blend ideas. Probably, the author may find value in these thoughts and wants to integrate them into their manuscript.

4. Figure 1: The AI use for Figure 1 is disclosed. But the figure needs some clarification. The arrow of Stage 1 leads to Stage 3. Is this intentional? Why do the arrows start at the AI part and not at the human part in Stages 1 and 2? How does this reflect human involvement; is the AI triggering the next step? Only in Stage 3 does the arrow lead from AI to Human – does this suggest a sequence, while the other stages are run in parallel? What do the exclamation marks refer to? I would suggest elaborating on the “visual language” used in the figure in the caption and properly controlling the AI-generated figure to match your suggested framework to avoid confusion.

5. Terminology: It seems as if the terms AI and LLM are used interchangeably. If possible, the author may want to consider sticking to one term or clarifying the usage of both terms. Especially in Table 2, this double usage confused me.

6. Conclusion: Maybe soften this statement: “The only ethically and scientifically defensible path is Scenario 3” and open the solution space for other scenarios that were not outlined in the manuscript.

Declarations

Potential competing interests: No potential competing interests to declare.