



## Editorial

# The indispensable human author: Integrating large language models in anesthesiology research with accountability

Lalit Gupta<sup>1\*</sup>, Devang Bharti<sup>1</sup>

<sup>1</sup>Dept. of Anaesthesia and Intensive Care, Maulana Azad Medical College and Associated Lok Nayak Hospital, New Delhi, India

Received: 30-09-2025; Accepted: 11-10-2025; Available Online: 31-10-2025

This is an Open Access (OA) journal, and articles are distributed under the terms of the [Creative Commons Attribution-NonCommercial-ShareAlike 4.0 License](https://creativecommons.org/licenses/by-nc-sa/4.0/), which allows others to remix, tweak, and build upon the work non-commercially, as long as appropriate credit is given and the new creations are licensed under the identical terms.

For reprints contact: [reprint@ipinnovative.com](mailto:reprint@ipinnovative.com)

The increasing use of Large Language Models (LLMs) like ChatGPT into academia is transforming scholarly communication. In anesthesiology and pain medicine, where clinicians balance immense clinical workload with publication ambitions, these tools offer a lucrative opportunity: to streamline the tedious process of manuscript preparation.<sup>1</sup> However, their use is a double-edged sword, and demands a proactive, ethical, and focused response from the scientific community. The critical question is no longer if we should use them, but how to use them responsibly, without compromising the scientific integrity, originality, and human expertise that form the bedrock of our scientific publication process.

## 1. The Potential: A Powerful Force Multiplier

There are multiple potential advantages of the use of LLMs for scientific publications. LLMs can serve as a powerful force multiplier, ensuring enhanced productivity across different domains:

1. Idea generation and hypothesis formation: They can accelerate the initial phase of a project by scanning voluminous literature and identifying gaps in existing knowledge.
2. Overcoming writer's block: For many, the biggest hurdle is the blank page. LLMs can help by suggesting an outline or by simplifying complicated methodologies, making the writing process easier.

3. Enhancing language clarity: These tools can help improve grammar and clarity, especially for researchers who are not publishing in their first language. This ensures that their work is judged not on the quality of written language, but on the scientific merit.
4. Technical support: LLMs can assist with technical tasks such as formatting references to specific journal guidelines or suggesting appropriate methods for statistical analysis, thereby saving valuable time for busy clinicians.

## 2. The Peril: Risks and Ethical Dilemma

These benefits are shadowed by serious risks that must be addressed.<sup>2</sup>

1. Factual inaccuracy and hallucination: LLMs generate plausible text, not factual truth. They can invent references, fabricate data, and present incorrect information with exaggerated confidence, threatening the foundation of scientific trust.
2. Amplification of bias: These models can perpetuate and amplify biases present in their training data, potentially skewing the scientific narrative in subtle ways.
3. Breach of confidentiality: Submitting patient details or unpublished results into a public LLM constitutes a major breach of privacy and trust.

\*Corresponding author: Lalit Gupta  
Email: [lalit.doc@gmail.com](mailto:lalit.doc@gmail.com)

4. Undermining critical thinking: Over-reliance on these models may erode essential analytical skills, while excessive use in drafting can create issues like accountability and plagiarism.

### 3. Guidelines for Ethical Use: Transparency and Accountability

Given these risks, clear guidelines are essential. The following principles and boundaries must be established.

#### 3.1. Authorship and disclosure

LLMs must never be considered authors because they cannot possess the accountability or intellectual contribution that authorship requires.<sup>3</sup> An LLM is a tool, just like a statistical software. However, its use must be transparently disclosed. While the specific location of disclosure (e.g., methods section for substantive contributions, acknowledgements for language editing) can be debated, the principle cannot. Additionally, this disclosure must be explicitly mentioned in the cover letter to the editor. Remember, ultimately it is the human authors who retain full responsibility for every word written in the manuscript.

#### 3.2. Appropriate and prohibited uses

The appropriate use of LLMs varies across the manuscript lifecycle (**Figure 1**). They are well-suited for preparatory and supportive tasks: brainstorming, drafting outlines, and polishing language. However, their use must be strictly prohibited in core scientific functions. They must not be used to draft the results section, as the presentation of original data must be precise. Their role in interpreting results and drawing conclusions should be highly restricted, as these tasks require intellectual expertise. They should never be used to generate citations or for peer review, which requires confidential expert assessment.

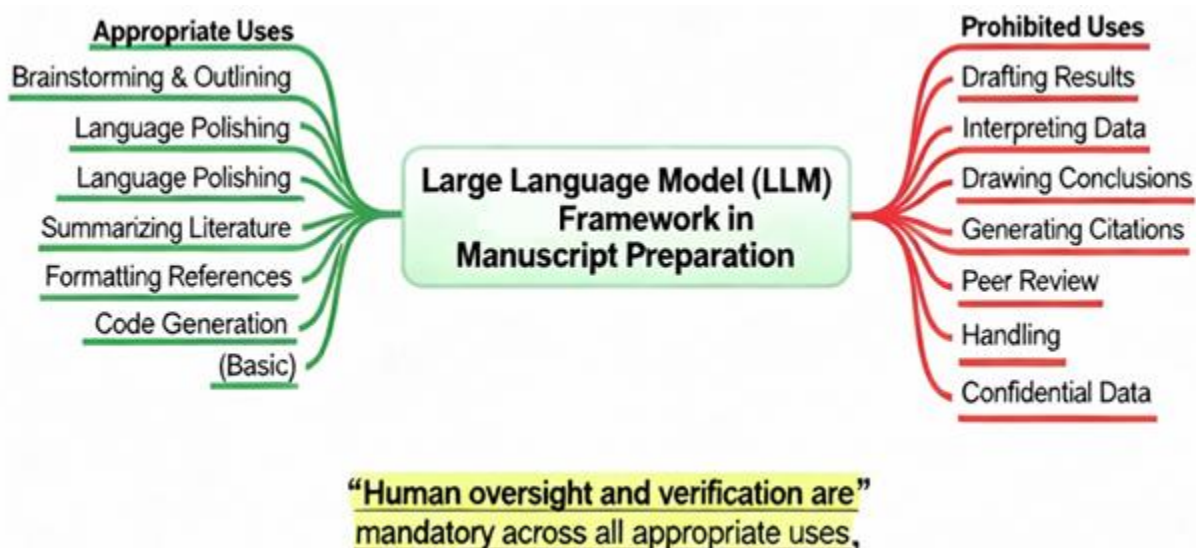
### 4. The Essential Principles and the Role of Journals

This leads to the essential ethical principles that must guide every researcher: Transparency, accountability, verification, integrity, and confidentiality. Using an LLM without disclosure breaches trust, while using its output without rigorous, expert verification is academic malpractice.<sup>4</sup> The human expert must remain the final arbiter of all content.

Journals have a critical role in providing clear guardrails. They must establish standardized policies that prohibit AI authorship, mandate disclosure, and include strong clauses on author accountability.<sup>5</sup> This ensures a level playing field and protects the integrity of the publication process. Editors may use LLMs for administrative tasks like technical checks but must never delegate editorial decisions to them.

### 5. Broader Implications and the Path Forward

Beyond immediate risks, we must consider emerging issues. The "AI-washing" of low-quality submissions threatens to overwhelm peer review.<sup>6</sup> The long-term erosion of writing skills and critical thinking poses a risk to the holistic development of budding scientists. Legal questions surrounding the copyright of AI-generated content, along with the growing issue of data contamination, where models are trained on AI-created text, are emerging as substantial challenges.<sup>7</sup> Current legal consensus holds that copyright requires human authorship; thus, the raw outputs of AI systems cannot be copyrighted, although derivative works involving human creativity may be eligible.<sup>8</sup> Multiple ongoing lawsuits focus on both, the use of copyrighted material for training AI models and on the potential for infringement when AI-generated works closely resemble existing protected works.



**Figure 1:** A framework for LLM use in manuscript preparation

## 6. Future Directions and Collaborative Efforts

As LLMs continue to evolve, ongoing research is essential to better understand their capabilities, limitations, and ethical implications in academic publication. Collaborative efforts between anesthesiologists, AI developers, ethicists, and journal editors will be crucial to develop standardized guidelines, best practices, and verification tools that preserve scientific integrity while maximizing the benefits of these technologies. Additionally, periodic updates to policies will be necessary to keep pace with rapid advancements in AI, ensuring that the responsible use of LLMs remains aligned with emerging challenges and opportunities.<sup>9</sup>

In conclusion, LLMs present a dual reality for academic anesthesiology. They are powerful tools that can enhance productivity. However, they are also potential vectors for error, bias, and ethical breach. Their value is entirely dependent on the expertise, judgment, and integrity of the human user. We must embrace their potential while instituting robust safeguards centered on transparency, human oversight, and unwavering accountability. The goal should be to use these tools to augment our capabilities, ensuring that the human intellect remains the definitive author of all scientific progress.

## 7. Conflict of Interest

None.

## References

1. van Dis EAM, Bollen J, Zuidema W, van Rooij R, Bockting CL. ChatGPT: five priorities for research. *Nature*. 2023;614(7947):224–6. <https://doi.org/10.1038/d41586-023-00288-7>.
2. Alkaissi H, McFarlane SI. Artificial Hallucinations in ChatGPT: Implications in Scientific Writing. *Cureus*. 2023;15(2):e35179. <https://doi.org/10.7759/cureus.35179>.
3. International Committee of Medical Journal Editors (ICMJE). Defining the role of authors and contributors [Internet]. ICMJE; [cited 2025 Sep 25]. Available from: <https://www.icmje.org/recommendations/browse/roles-and-responsibilities/defining-the-role-of-authors-and-contributors.html>
4. Hosseini M, Horbach SPJM. Fighting reviewer fatigue or amplifying bias? Considerations and recommendations for use of ChatGPT and other large language models in scholarly peer review. *Res Integr Peer Rev*. 2023;8(1):4. <https://doi.org/10.1186/s41073-023-00133-5>.
5. Thorp HH. ChatGPT is fun, but not an author. *Science*. 2023;379(6630):313. <https://doi.org/10.1126/science.adg7879>.
6. Stokel-Walker C. ChatGPT listed as author on research papers: many scientists disapprove. *Nature*. 2023;613(7945):620–1. <https://doi.org/10.1038/d41586-023-00107-z>.
7. Gao CA, Howard FM, Markov NS, Dyer EC, Ramesh S, Luo Y, et al. Comparing scientific abstracts generated by ChatGPT to real abstracts with detectors and blinded human reviewers. *NPJ Digit Med*. 2023;6(1):75. <https://doi.org/10.1038/s41746-023-00819-6>.
8. Elmahjub E. The algorithmic muse and the public domain: why copyright's legal philosophy precludes protection for generative AI outputs. *Sci Technol Hum Values*. 2025;50(2):453–73.
9. models in anaesthesiology: Current insights and future directions—A narrative review. *Indian J Clin Anaesth*. 2025;12(2):190–7. <https://doi.org/10.18231/j.ijca.2025.033>.

**Cite this article:** Gupta L, Bharti D. The indispensable human author: Integrating large language models in anesthesiology research with accountability. *Indian J Clin Anaesth*. 2025;12(4):564–566.