

The Ten Commandments for Responsible Use of Artificial Intelligence

Why Education—Not Fear—Determines Whether AI Helps or Harms

Letter to the Editor

¹Lev Neymotin, PhD
N5Digital, Inc.

This letter is written in response to the editorial “*ChatGPT Warns: ‘Don’t Trust ChatGPT’*” by Petrikovsky et al., published in *Neonatal Intensive Care*, Vol. 39, No. 1 (Winter 2026), pages 9–10. The authors appropriately raise concerns regarding uncritical reliance on artificial intelligence in medicine, particularly in high-risk fields such as obstetrics and gynecology. However, the central implication of the article—that the primary danger lies in the AI tool itself—deserves reconsideration.

Artificial intelligence, like any sophisticated instrument, is neither inherently safe nor inherently dangerous. Its impact is determined entirely by the education, intent, and discipline of the user. History offers endless analogies. A sharp knife can nourish or injure. A drill can build or destroy. A saw can create or amputate. As tools become more powerful and precise, the margin for misuse grows, and with it the requirement for training. We do not blame the tool for improper outcomes; we train the operator. AI belongs squarely in this category of advanced instruments.

Attempts to evaluate the safety of AI by analyzing its responses to poorly formulated, underspecified, or ambiguous questions are conceptually immature. A vague question necessarily produces a vague or misleading answer. This is not a failure of intelligence but a failure of communication. AI systems respond strictly to the information, context, and constraints provided. When these are absent, the system must fill gaps with generalities or assumptions. Blaming AI for such outcomes is equivalent to blaming a calculator for an incorrect result when the wrong equation was entered.

The real risk identified in the referenced article is not AI itself, but an uneducated public—and, increasingly, untrained professionals—using a powerful tool without understanding its operating principles or limitations. No advanced tool can be made intrinsically safe without simultaneously rendering it ineffective. Safety and effectiveness arise only through education. What is required, therefore, is not fear-based discouragement or categorical distrust, but a modest, practical framework that teaches users how to engage with AI responsibly.

To that end, we propose a simple educational construct, deliberately framed as “Ten Commandments for Responsible AI Use.” These principles emphasize that AI is a tool rather than an authority; that poorly formulated questions produce unreliable answers; that relevant context and clearly defined goals are essential; that AI must never replace professional judgment in high-risk decisions; and that increasing stakes demand increasing

¹ Lev Neymotin, PhD, former researcher, Brookhaven National Laboratory

levels of human oversight. Such guidance is easily taught, easily remembered, and sufficient to prevent the very misapplications highlighted in the editorial.

Artificial intelligence can meaningfully enhance education, efficiency, and clinical reasoning when used correctly. It becomes dangerous only when users abandon precision, accountability, and professional responsibility. The solution, therefore, lies not in distrusting the tool, but in educating the user.

To translate this educational imperative into practice, we propose a concise and lay-accessible framework, “The Ten Commandments for Responsible AI Use,” provided as an Appendix for readers who wish to apply these principles in clinical and public-facing contexts.

Reference

Petrikovsky BM, Dynkin BI, Klune T, Alyeshmerni D. *ChatGPT Warns: “Don’t Trust ChatGPT.”* Neonatal Intensive Care. 2026;39(1):9–10.

Appendix

Educational Framework for Safe and Effective AI Use

1. Artificial intelligence is a tool, not an authority. It informs but does not decide.
2. Poorly formulated questions produce unreliable answers.
3. Context is essential; relevant background information must always be provided.
4. The objective of the inquiry should be clearly stated.
5. Artificial intelligence cannot infer personal circumstances unless they are explicitly described.
6. AI must never be used as the sole basis for medical, legal, financial, or safety-critical decisions.
7. Important information generated by AI should be independently verified.
8. Confidence and fluency of language do not guarantee accuracy.
9. AI should enhance human reasoning, not replace human judgment or responsibility.
10. As the potential consequences increase, mandatory human oversight becomes essential.

