

Oksana Dubrova

*Berdyansk State Pedagogical University
(Berdyansk, Ukraine)*

HALLUCINATED GRAMMAR: SYSTEMATIC MORPHOSYNTACTICAL ERRORS IN LARGE LANGUAGE MODELS AND THEIR PEDAGOGICAL IMPLICATIONS

The growing interest in the use of Large Language Models (LLMs) across various educational fields – from essay writing to automated feedback – has raised fundamentally new questions for educators regarding the grammatical reliability of texts generated by artificial intelligence. Systematic morphosyntactic deviations have been studied far less than lexical and factual “hallucinations” of LLMs (E.Bender et al., 2021).

Examining and analyzing hallucinated grammatical constructions in contemporary English is relevant because of several factors, including students' frequent use of various AI tools to complete grammar exercises, teachers' lack of awareness of the particular nature of AI-generated errors, and the possibility that students will internalize incorrect grammatical patterns as a result of being exposed to grammatically created AI-generated texts regularly.

This topic is becoming increasingly popular in the analysis and research conducted by both international and Ukrainian scholars. For example, Emily M. Bender with other researches in the article “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?”, systematically described for the first time the mechanism of generating “form without meaning” in LLMs, laying the theoretical groundwork for studying grammatical hallucinations. The authors argue that “models operate on statistical patterns rather than grammatical rules, which inevitably generate plausible but structurally incorrect constructions” (E.Bender et al., 2021).

In the article “Dissociating language and thought in large language models” K. Mahowald and co-authors empirically confirmed that LLMs demonstrate formal linguistic competence (grammatical correctness of form), yet systematically fail in the

functional aspects of language – specifically, in the use of tense and aspect forms under conditions of pragmatic ambiguity (K. Mahowald, 2024).

In their study titled “A comprehensive survey of grammatical error correction,” Y. Wang and a team of researchers applied corpus analysis to the outputs of GPT models and identified “recurring morphosyntactic blind spots”: specifically, difficulties with the article system in abstract nouns and errors in the use of the perfect tense versus the simple past tense in narrative contexts (Y. Wang, 2021).

In turn, O. Chugai and K. Havrylenko in their paper “ChatGPT: Attitudes and Experiences of Technical University Students in Ukraine,” explore the cognitive mechanisms underlying Ukrainian students’ acquisition of English grammar by analyzing interference errors that are typologically related to AI hallucinations, making their theoretical framework suitable for describing machine errors (O. Chugai and K. Havrylenko, 2024). O. Pavlenko & A. Syzenko in their research, raise the issue of students’ critical attitude toward AI-generated texts and propose a “grammatical audit” methodology as a tool for media literacy in language education (O. Pavlenko & A. Syzenko, 2024).

Based on an analysis of text fragments generated by GPT-4 and Claude 3 by students while performing practical tasks, we suggest that four main types of grammatical hallucinations can be identified.

1. Article Misuse (this system is one of the most systematically “hallucinated” aspects of English grammar in LLMs. Models demonstrate a consistent tendency to use the definite or indefinite article with uncountable abstract nouns, as well as to ignore generic semantics when choosing between a/an/the/-:

- “*The research provides an evidence...*” – „... provides evidence” (zero article with uncountable nouns);

- “*A knowledge of the grammar is a foundation of academic writing*” – „knowledge” is an uncountable noun (does not take ‘a’); “grammar,” in the sense of a system, is used

without an article; “the foundation” is justified if a specific, unique foundation is implied; however, “a knowledge” is a gross article error.

2. Perfect–Preterite Confusion (confusion between the Present Perfect and the Past Simple is the second most common type of hallucination. LLMs regularly use the Perfect in contexts where a time marker (a specific date, the adverb “ago,” “in + year,” etc.) requires the Simple Past, and conversely, they avoid the Perfect in resultative contexts without a time marker:

- “*Scientists have discovered this effect in 1998...*” – “Scientists discovered this effect in 1998” (Past Simple with a specific time marker); “

- “*Darwin has published ‘On the Origin of Species’ in 1859 and has changed the course of biology*” – the specific year (in 1859) is a deictic marker of a completed past action – the Simple Past is mandatory; the Present Perfect here is a typical L1 interference and an AI hallucination.

3. Modal Stacking / Epistemic Overcautioning (LLMs have a pronounced tendency to accumulate modal operators and hedging constructions within a single sentence. This phenomenon is linked to training based on the RLHF (Reinforcement Learning from Human Feedback) principle, which encourages “cautious” responses, but leads to pragmatically redundant and grammatically anomalous constructions:

- “*The results would seem to suggest that it might possibly indicate...*” – excessive layering of hedges without pragmatic necessity;

- “*It could perhaps be argued that there may potentially exist certain possible factors which might conceivably influence the outcome*” – six hedging elements (could, perhaps, may potentially, certain possible, might conceivably) render the sentence pragmatically meaningless; this is a classic case of “modal hallucination” resulting from excessive caution.

4. Syntactic Overembedding:

- “The data which was collected by researchers who worked at the lab that had been established in...” – five levels of relative clause nesting; the optimal solution is nominalization and compression in a postpositive construction;

- “This is a problem that has been identified by scholars who have been studying the phenomenon that emerges in contexts where the conditions necessary for its development are present” – accumulation of that-clauses (four levels) without semantic gain; restructuring via a participle clause and compression eliminates excessive nesting.

Thus, current grammatical hallucinations in LLMs are becoming a logical result of the architectural constraints of transformer models, which function on probabilistic patterns without access to grammatical rules as such, rather than a random artifact. For the educational community, this means shifting from naively trusting AI-generated texts to fostering a robust critical stance among students toward machine-generated content. Longitudinal investigations into the effects of AI-generated texts on language learners’ acquisition of grammatical competence, and the creation of reliable instruments to evaluate the grammatical reliability of LLMs, should be the main areas of future research.

REFERENCES

1. Павленко, О., & Сизенко, А. (2024). Using ChatGPT as a learning tool: A study of Ukrainian students' perceptions. *Arab World English Journal (AWEJ)*, Special Issue on ChatGPT, 252–264. <https://dx.doi.org/10.24093/awej/ChatGPT.17>
2. Чугай, О., & Гавриленко, К. (2024). ChatGPT: ставлення та досвід студентів технічного університету в Україні. *Information Technologies and Learning Tools*, 101(3), 15–27. <https://doi.org/10.33407/itlt.v101i3.5559>
3. Bender, E. M., Gebru, T., McMillan-Major, A., & Shmitchell, S. (2021). On the dangers of stochastic parrots: Can language models be too big? In *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency (FAccT '21)* (pp. 610–623). Association for Computing Machinery. <https://dl.acm.org/doi/10.1145/3442188.3445922>
4. Mahowald, K., Ivanova, A. A., Blank, I. A., Kanwisher, N., Tenenbaum, J. B., & Fedorenko, E. (2024). Dissociating language and thought in large language models. *Trends in Cognitive Sciences*, 28(6), 529–541. <https://doi.org/10.1016/j.tics.2024.01.011>
5. Wang, Y., Wang, Y., Dang, K., Liu, J., & Liu, Z. (2021). A comprehensive survey of grammatical error correction. *ACM Transactions on Intelligent Systems and Technology*, 12(5), 1–51. <https://doi.org/10.1145/3474840>