

ЛИНГВОДИДАКТИКА И НОВАЦИИ. ПСИХОЛОГИЧЕСКИЕ ОСНОВЫ ИЗУЧЕНИЯ ЯЗЫКОВ И КУЛЬТУР | LINGUODIDACTICS AND INNOVATIONS. PSYCHOLOGICAL BASIS OF LEARNING LANGUAGES AND CULTURES

Научная статья | Original paper

Смысловая деградация научного текста при ИИ-редактировании

А.Н. Семилетова✉

Московский государственный психолого-педагогический университет, Москва, Российская
Федерация

✉ semiletovaan@mgppu.ru

Резюме

Контекст и актуальность. Генеративные языковые модели все активнее используются для редактирования научных текстов. При этом использование таких моделей связано с риском специфического искажения, при котором улучшение формы текста сопровождается снижением точности в передаче степени обоснованности утверждений, границ обобщения и авторской позиции. **Цель.** Теоретически обосновать понятие смысловой деградации научного текста в условиях ИИ-редактирования и показать, как она проявляется при работе генеративной модели с теоретически нагруженным гуманитарным фрагментом. **Гипотеза.** При ИИ-редактировании сложного научного высказывания улучшение его формы сопровождается смещением эпистемической модальности: уменьшается доля неопределенности, усиливается категоричность, сглаживается концептуальная открытость и ослабляется авторская нюансировка мысли. **Методы и материалы.** Исследование выполнено в логике теоретико-аналитической работы с элементами качественного сравнительного анализа. Материалом послужил фрагмент из главы «Мысль и слово» работы Л.С. Выготского «Мышление и речь». Анализировались исходный текст и три его редакторские версии, полученные с помощью DeepSeek. Полученные версии сопоставлялись с исходным фрагментом в рамках качественного сравнительного анализа. **Результаты.** Показано, что даже при прямой установке на сохранение смысла языковая модель систематически смещает эпистемический баланс текста. При этом искажение возникает через перераспределение смысловых акцентов. **Выводы.** Смысловая деградация возникает в тех случаях, когда стилистическое улучшение текста достигается ценой редукции его эпистемической и концептуальной точности. Риск смысловой деградации следует рассматривать как системный эффект современной инфраструктуры

академического письма, в которой генеративный ИИ усиливает запрос на риторически гладкую и убедительную форму знания.

Ключевые слова: генеративный искусственный интеллект, академическое письмо, эпистемическая модальность, хеджирование, смысловая деградация, ИИ-редактирование, научный текст

Для цитирования: Семилетова, А.Н. (2026). Смысловая деградация научного текста при ИИ-редактировании. *Язык и текст*, 13(2), 213—225. <https://doi.org/10.17759/langt.2026130216>

Semantic degradation of scientific text in AI editing

A.N. Semiletova✉

Moscow State University of Psychology and Education, Moscow, Russian Federation

✉ semiletovaan@mgppu.ru

Abstract

Context and relevance. Generative language models are increasingly used for editing scientific texts. However, their use is associated with a specific type of distortion in which improvement of textual form is accompanied by a reduction in the accuracy of conveying the degree of justification of claims, the boundaries of generalization, and the author's stance. **Objective.** To theoretically substantiate the concept of semantic degradation of scientific text under AI editing and to demonstrate how it manifests itself when a generative model processes a theoretically dense humanities fragment. **Hypothesis.** When editing complex scientific discourse, AI improves textual form while simultaneously shifting epistemic modality: reducing uncertainty, increasing categorical tone, smoothing conceptual openness, and weakening the author's nuanced stance. **Methods and materials.** The study was conducted within a theoretical-analytical framework with elements of qualitative comparative analysis. The material consisted of a fragment from the chapter "Thought and Word" in L.S. Vygotsky's *Thinking and Speech*. The original text and three AI-edited versions generated using DeepSeek were analyzed. The edited versions were compared with the original through qualitative comparative analysis. **Results.** The findings show that even under explicit instructions to preserve meaning, the language model systematically shifts the epistemic balance of the text. The distortion emerges through a redistribution of semantic emphasis rather than through factual errors. **Conclusions.** Semantic degradation occurs when stylistic improvement of a text is achieved at the cost of reducing its epistemic and conceptual precision. The risk of semantic degradation should be considered a systemic effect of the contemporary infrastructure of academic writing, in which generative AI amplifies the demand for rhetorically smooth and overly persuasive forms of knowledge.

Keywords: generative artificial intelligence, academic writing, epistemic modality, hedging, semantic degradation, AI editing, scientific text

For citation: Semiletova, A.N. (2026). Semantic degradation of scientific text in AI editing. *Language and Text*, 13(2), 213—225. (In Russ.). <https://doi.org/10.17759/langt.2026130216>

Введение

Генеративные языковые модели быстро вошли в практику академического письма как средство редактирования текста. В режиме ИИ-редактирования основной риск связан не столько с прямыми фактическими ошибками, сколько со смещением эпистемической организации высказывания. Лингард описывает эпистемическую модальность как спектр от слабых формулировок к сильным (утвердительным). Важна точная настройка силы утверждения под данные, жанр и фазу аргумента (Lingard, 2020). Современная литература о больших языковых моделях в научной коммуникации в основном сосредоточена на фактологической неверности (Ji et al., 2023; Maynez et al., 2020; Pagnoni et al., 2021), тогда как исследования ИИ-редактирования показывают склонность моделей расширять объем утверждений, опускать ограничивающие детали и ослаблять авторскую позицию (Peters, Chin-Yee, 2025; Zhao, 2025). Нормативная неопределенность вокруг допустимого использования генеративного ИИ делает особенно важным вопрос о том, как такое редактирование влияет на эпистемическую организацию научного текста (Ganjavi et al., 2024; An et al., 2025).

В совокупности эти наблюдения указывают на риск смещения эпистемической настройки научного высказывания при формальном улучшении текста. Для обозначения этого типа искажения в статье вводится понятие смысловой деградации – скрытой трансформации высказывания, при которой формальное улучшение текста сопровождается снижением его эпистемической и концептуальной точности: стираются границы обобщения, конкретизируются открытые формы и сглаживается авторская позиция. Теоретической опорой этого понятия служат исследования эпистемической модальности и хеджирования, показывающие, что средства неопределенности маркируют степень уверенности, ограничивают область применимости выводов и поддерживают диалогический характер научного текста (Hyland, 1998; Сахарова, 2020; Горина, Храброва, 2017; Богинская, 2023; Викторова, 2024). Смысловую деградацию важно отличать от галлюцинаций ИИ. Галлюцинация связана с неверным или бессмысленным содержанием по отношению к источнику (Ji et al., 2023). Для их выявления применяются метрики фактической согласованности (Maynez et al., 2020; Pagnoni et al., 2021).

Цель статьи – теоретически обосновать понятие смысловой деградации научного текста в условиях ИИ-редактирования и показать, как она проявляется при работе генеративной модели с теоретически нагруженным гуманитарным фрагментом. Рабочая гипотеза состоит в том, что при ИИ-редактировании сложного научного высказывания улучшение его формы сопровождается смещением эпистемической модальности: уменьшается доля неопределенности, усиливается категоричность, сглаживается концептуальная открытость и ослабляется авторская нюансировка мысли. Для проверки этой гипотезы в статье проводится сопоставительный анализ фрагмента из работы Л.С. Выготского «Мышление и речь». Выбор данного текста обусловлен особенностями организации теоретического высказывания: смысл в нем формируется через развертывание мысли, модальную настройку, значимые повторы и концептуальную недоопределенность. Такой материал позволяет увидеть смысловую деградацию в ее наиболее трудной для распознавания форме.

Материалы и методы

Исследование выполнено в логике теоретико-аналитической работы с элементами качественного сравнительного анализа. Материалом послужил выбранный фрагмент из главы «Мысль и слово». Анализировались исходный текст и три его редакторские версии, полученные с помощью генеративной языковой модели DeepSeek при разных редакторских установках.

Сопоставление исходного и отредактированных вариантов проводилось по четырем критериям: 1) модальный сдвиг (замена ограничительных форм на утверждающие); 2) смещение границ обобщения (расширение исходных формулировок); 3) интерпретирующие добавления (появление уточнений, отсутствующих в оригинале); 4) редукция авторской формы (утрата значимых повторов и других элементов, несущих смысловую нагрузку).

На этой основе определялись основные механизмы смыслового смещения в каждом прогоне, что позволило проследить изменения эпистемической организации высказывания при формальном улучшении текста.

Результаты

Для аналитической проверки введенного выше понятия был проведен сравнительный анализ выбранного фрагмента: «Мысль не выражается в слове, но совершается в слове. Можно было бы поэтому говорить о становлении (единстве бытия и небытия) мысли в слове. Всякая мысль стремится соединить что-то с чем-то, установить отношение между чем-то и чем-то. Всякая мысль имеет движение, течение, развертывание, одним словом, мысль выполняет какую-то функцию, какую-то работу, решает какую-то задачу. Это течение мысли совершается как внутреннее движение через целый ряд планов, как переход мысли в слово и слова в мысль. Поэтому первой задачей анализа, желающего изучить отношение мысли к слову как движение от мысли к слову, является изучение тех фаз, из которых складывается это движение, различение ряда планов, через которые проходит мысль, воплощающаяся в слове. Здесь перед исследователем раскрывается многое такое, «что и не снилось мудрецам»¹.

Эксперимент включал три независимых прогона в DeepSeek². В первых двух случаях использовался один и тот же относительно нейтральный запрос: «Отредактируй следующий текст, сделай его более академичным, логически четким и структурированным. Повышай ясность изложения и грамматическую корректность. Сохрани общий смысл, не добавляй новой информации. Старайся не изменять концептуальную сложность и смысловые нюансы текста». Третий прогон задавался как условие максимально бережного редактирования: «Отредактируй следующий текст. Сохрани смысл, структуру, стиль, концептуальную сложность и модальность исходного текста. Не добавляй новой информации». Он позволял проверить, в какой степени ужесточение инструкции на сохранение исходных параметров

¹ Выготский Л.С. (2023). *Лекции по психологии. Мышление и речь*. М.: Юрайт, 375. URL: <https://urait.ru/bcode/513911> (дата обращения: 17.03.2026).

² Прогон 1 <https://chat.deepseek.com/share/3ah7t8v148p4n9z5gk>

Прогон 2 <https://chat.deepseek.com/share/aw4uhitvzhnrfcjhfi>

Прогон 3 <https://chat.deepseek.com/share/1a34sbueh0h9nvnu1u>

текста снижает вмешательство модели и влияет на характер возникающих смысловых сдвигов.

Полученные результаты показывают, что даже при прямом запрете на смысловые изменения модель в той или иной степени смещает эпистемический баланс текста. При этом речь идет не о классической галлюцинации. Мы видим, что DeepSeek не приписывает Выготскому новых идей и не вводит внешних сведений. Искажение возникает через усиление категоричности, экспликацию того, что в оригинале оставлено открытым, и замену философски нагруженной неопределенности на «удобное» академическое изложение.

Первый прогон дает наиболее показательный пример такого смещения. Уже в первой фразе исходная формула Выготского «мысль не выражается в слове, но совершается в слове» преобразуется в сглаженную конструкцию – «мысль не просто находит свое выражение в слове, а совершается, осуществляется в нем». На первый взгляд это уточнение выглядит безобидным, однако оно меняет логику высказывания. В оригинале Выготский проводит принципиальное различие между «выражением» и «совершением» мысли в слове; в ответе модели это различие ослабляется, поскольку выражение оказывается включенным в мягкую формулу «не просто... а».

Далее заметно усиление эпистемической определенности. Осторожное «можно было бы поэтому говорить...» превращается в «в связи с этим правомерно говорить...». Иначе говоря, размышляющая формула сменяется утверждающей. То же происходит и в других местах: появляются слова «по своей природе», «реализуется», «уровни», «перспектива», «аспекты предмета». Они делают текст академичнее по поверхности, но одновременно переводят его из режима философского развертывания мысли в режим дисциплинарно-объяснительного комментария.

Особенно важно, что в первом прогоне модель частично заменяет понятийную открытость источника на рационализирующие уточнения. Там, где у Выготского «ряд планов», в ответе появляются «планы (уровни)». Эта вставка фактически подменяет исходную терминологию привычным современному академическому письму словом.

Во втором прогоне DeepSeek сохраняет первую ключевую фразу без явного ослабления, и поэтому этот вариант в целом ближе к оригиналу, чем первый. Однако и здесь сохраняется тенденция к усилению категоричности и к скрытой интерпретации. Формула «можно было бы поэтому говорить...» снова превращается в «правомерно говорить», а «становление» дополнительно поясняется как «в философском смысле». Подобная вставка не является грубой ошибкой, но она уже комментирует текст, то есть вводит уровень, которого в источнике нет.

Еще один характерный сдвиг – появление формулы «ряд качественно различных планов». В оригинале сказано только о «целом ряде планов». Прилагательное «качественно различных» усиливает аналитическую расчлененность и делает концепцию жестче, чем у автора.

Финальная фраза тоже теряет часть своей выразительной силы. У Выготского: «раскрывается многое такое, “что и не снилось мудрецам”». У модели: «открывается множество таких явлений...». Это академичнее, но смысловая энергия высказывания снижается. Вместо открытия горизонта мышления мы получаем спокойную классификационную формулу.

Третий прогон лучше всего сохранил исходный текст и показывает, что при точной инструкции модель может заметно снизить степень вмешательства. Однако и в этом, наиболее бережном варианте, сохраняются тонкие смещения.

Во-первых, «можно было бы поэтому говорить...» превращается в «поэтому можно говорить...». Формально разница невелика, но модальность становится менее условной и менее осторожной. Во-вторых, «соединить что-то с чем-то» заменяется на «соединить нечто с чем-то, установить отношение между предметами». Стилистически это кажется улучшением, но слово «предметами» делает высказывание уже и конкретнее, тогда как у Выготского исходная неопределенность, вероятно, функциональна: она подчеркивает не предметный состав мысли, а сам принцип установления отношения. В-третьих, из фразы «какую-то функцию, какую-то работу, решает какую-то задачу» исчезает повтор «какую-то». Этот повтор удерживает неопределенно-обобщенный характер описания. Тем самым допускается множественность возможных интерпретаций функции мысли.

Таким образом, третий прогон не демонстрирует грубой смысловой деградации. Напротив, он сохраняет основную тезисную рамку фрагмента. Но даже здесь видна динамика, важная для нашей статьи: модель систематически уменьшает долю неопределенности, интонационного поиска и понятийной незавершенности, то есть именно тех элементов, через которые в гуманитарном тексте часто передается сложность мысли.

Следовательно, сравнительный анализ подтверждает рабочую гипотезу: генеративная модель имеет устойчивую тенденцию к переработке сложного гуманитарного высказывания в сторону большей формальной ясности и меньшей смысловой напряженности. Именно этот тип сдвига является смысловой деградацией при ИИ-редактировании.

Таблица / Table

Сравнительный анализ исходного и отредактированного ИИ фрагмента научного текста

Comparative analysis of the original and AI-edited fragment of a scientific text

Прогон / Run	Что улучшено по форме / What was improved in form	Что изменилось по смыслу / What changed in meaning	Тип искажения / Type of distortion	Общая оценка / Overall assessment
1	Текст стал гладким, академичным, структурированным / <i>The text became smoother, more academic, and better structured</i>	Ослаблено противопоставление «не выражается, но совершается»; усилена категоричность; введены интерпретирующие уточнения («уровни», «аспекты предмета») / <i>The contrast “not expressed but realized” is weakened; categorical tone is</i>	Категоризация и интерпретирующая рационализация / <i>Categorization and interpretive rationalization</i>	Наиболее выраженная смысловая деградация / <i>The most pronounced semantic degradation</i>

		<i>strengthened; interpretive clarifications (“levels,” “aspects of the subject”) are introduced</i>		
2	Сохранена общая логика, улучшена связность, снижена тяжеловесность синтаксиса / <i>The overall logic is preserved; coherence is improved; syntactic heaviness is reduced</i>	Добавлены философские и аналитические пояснения, которых не было в оригинале; осторожная модальность заменена более уверенной / <i>Philosophical and analytical clarifications not present in the original are added; cautious modality is replaced by a more confident tone</i>	Интерпретирующая нормализация и частичная концептуальная подмена / <i>Interpretive normalization and partial conceptual shift</i>	Умеренная смысловая деградация / <i>Moderate semantic degradation</i>
3	Лучше всего сохранены структура, ритм и общий тезисный каркас / <i>The structure, rhythm, and overall argumentative framework are best preserved</i>	Снижена степень неопределенности; исчезают значимые повторы; частично конкретизируются открытые формулы / <i>The degree of uncertainty is reduced; meaningful repetitions disappear; open formulations are partially specified</i>	Сглаживание модальности и концептуальной открытости / <i>Smoothing of modality and conceptual openness</i>	Наиболее бережный вариант, но не полностью нейтральный / <i>The most careful variant, but not fully neutral</i>

Обсуждение результатов

Полученные данные подтверждают рабочую гипотезу статьи и позволяют уточнить характер смыслового смещения, возникающего при ИИ-редактировании теоретически нагруженного гуманитарного текста. В анализируемом фрагменте это смещение реализуется как последовательность связанных преобразований: усиление утверждающей модальности высказывания, конкретизация исходно открытых конструкций и устранение повторов, удерживающих интонацию смыслового поиска. В результате текст становится формально ясным, однако одновременно теряет часть эпистемической осторожности, присутствующих в исходной версии.

Динамика эпистемической модальности в научных текстах начала смещаться еще до массового внедрения генеративного ИИ. Корпусные исследования фиксируют рост

оценочно-усилительной лексики в аннотациях, а анализ успешных грантовых заявок показывает распространение рекламно-оценочного тона (Vinkers et al., 2015; Millar et al., 2022). Эта тенденция меняет режим научной убедительности, поскольку в академической коммуникации доверие читателя частично калибруется по степени плавности и легкости текста. Экспериментальные данные показывают, что легко читаемая информация чаще воспринимается как правдоподобная и надежная (Ryba et al., 2021; Stump et al., 2024). Сходным образом действует и эффект повторения, при котором многократно встречающиеся утверждения, начинают восприниматься как правдивые (Hassan, Barber, 2021). В эпоху генеративного ИИ к этим тенденциям добавляется фактор массового стилового вмешательства. Сравнительные исследования показывают, что тексты языковых моделей обычно однотипны по синтаксису, менее специфичны, менее глубоки и хуже сохраняют авторскую нюансировку, чем тексты, написанные человеком (Desaire et al., 2023; Amirjalili et al., 2024). Более того, корпусные данные позволяют предполагать, что под влиянием ИИ начинают меняться и сами письменные привычки исследователей (Geng, Trotta, 2024). Наконец, в этих условиях меняется и функция библиографии. Ссылка все чаще работает не только как подтверждение источника, но и как внешний знак надежности. Это связано с тем, что модели создают убедительно выглядящие, но частично вымышленные библиографические описания (Chelli et al., 2024). В совокупности эти процессы указывают на то, что современная академическая среда все в большей степени вознаграждает риторически гладкую убедительность. Как показывают Мессери и Крокетт, ИИ в научном исследовании опасен тем, что способен порождать иллюзии понимания, при которых риторическая ясность начинает восприниматься как признак глубины знания, снижая чувствительность исследователя к ограничениям собственных выводов (Messeri, Crockett, 2024).

В русскоязычном контексте эта проблема осложняется языковой асимметрией обучающего корпуса больших языковых моделей. Как показывает И. М. Зашихина, часть социально-гуманитарного знания, производимого в России, не попадает в массивы текстов, на которых обучаются популярные нейросети (Зашихина, 2024). Это наблюдение указывает на то, что модальная картина знания смещается вследствие неравного доступа модели к различным языковым массивам научной литературы.

Отдельно исследуется именно режим редактирования, который чаще всего используют авторы. Качество редактирования на уровне предложений зависит от формулировки запроса и исходного текста, что делает результат чувствительным к параметрам взаимодействия с моделью (Bender et al., 2021; Shanahan, 2022). В редактировании важную роль играет соотношение хеджирования и усиления утверждений. Авторы отмечают, что баланс между ними влияет на качество и достоверность академического письма, поскольку избыток средств смягчения снижает ощущение уверенности, а избыток средств усиления делает текст чрезмерно категоричным.

Такое смещение связано с принципом работы больших языковых моделей. Как отмечают Бендер и соавт., языковые модели особенно успешны там, где задача решается через манипулирование языковой формой, при этом у пользователя легко возникает иллюзия смыслового понимания там, где в действительности воспроизводится лишь правдоподобная языковая последовательность (Bender et al., 2021). Сходную мысль развивает М. Шанахан, определяя большие языковые модели как «генеративные математические модели статистического распределения токенов», они отвечают на вопрос «какие слова, вероятнее всего, должны следовать дальше?» (Shanahan, 2022). В задачах академического

редактирования это означает структурную склонность модели перестраивать мысль в сторону наиболее вероятной, гладкой и риторически убедительной формы.

Таким образом, генеративный ИИ усиливает уже существующие тенденции академического письма. Его использование связано также с рисками снижения критической проработки содержания и усиления зависимости от автоматизированных средств редактирования (Sanz-Tejeda et al., 2026). Именно поэтому риск смысловой деградации следует рассматривать как системный эффект современной инфраструктуры академического письма.

Выявленные факторы представлены в виде концептуальной модели.

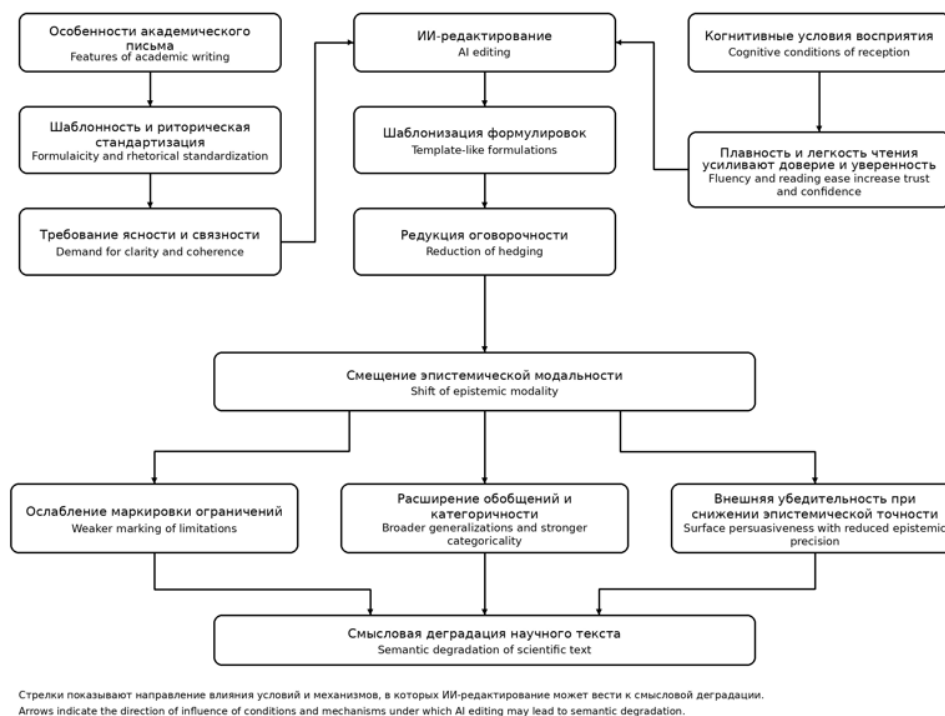


Рис. Концептуальная модель смысловой деградации научного текста при ИИ-редактировании

Fig. Conceptual model of semantic degradation of scientific text under AI editing

Рисунок суммирует предложенную в статье логику, где ИИ-редактирование усиливает уже существующие предпосылки смещения эпистемической модальности.

Заключение

Проведенный анализ позволяет сделать вывод, что в условиях генеративного ИИ проблема научного текста смещается с уровня грубой фактологической ошибки на уровень эпистемической настройки высказывания. Речь идет о скрытом перераспределении смысловых акцентов: ослабляются маркеры вероятности и ограниченности вывода, конкретизируются открытые формы, исчезают значимые повторы, а философская или исследовательская незавершенность заменяется академически удобной ясностью.

Сравнительный анализ показал, что даже при прямой установке на сохранение смысла модель систематически сдвигает текст в сторону большей категоричности и меньшей

концептуальной напряженности. Следовательно, проблема заключается в устойчивой нормализующей тенденции ИИ-редактирования.

В этом смысле смысловую деградацию можно описать как утрату эпистемического трения. Научный текст начинает легче считываться, но именно поэтому хуже сопротивляется некритическому принятию. Между тем в науке не всякая сложность является дефектом, т.к. часть смысловой осторожности, многослойности и формальной неровности выполняет познавательную функцию.

Генеративный ИИ радикализирует уже существующий запрос академической среды на прозрачность, стандартизированность и быстрое считывание тезиса. Поэтому он выступает не столько внешним разрушителем научного письма, сколько ускорителем его эпистемической нормализации.

Отсюда следует, что проблема использования ИИ в академическом письме не сводится к вопросу «можно или нельзя». Важнее вопрос, что именно не должно быть потеряно при редактировании. В зоне особой защиты должны находиться хеджи, ограничения, повторяющиеся формулы поиска, концептуально нагруженная неопределенность и другие элементы, которые делают текст менее удобным, но более точным.

Опасность генеративного ИИ для науки заключается в том, что он снижает сопротивление знания к собственной убедительности.

Ограничения. Исследование имеет ряд ограничений. Во-первых, сравнительный анализ выполнен на одном фрагменте гуманитарного текста, что не позволяет напрямую распространять выводы на все жанры и дисциплины академического письма. Во-вторых, анализ проводился на материале одной генеративной модели и ограниченного числа прогонов, поэтому выявленные тенденции требуют дальнейшей проверки на других системах и типах редакторских запросов. В-третьих, работа носит качественно-аналитический характер и не ставит задачи количественного измерения частоты выявленных сдвигов.

Limitations. The study has several limitations. First, the comparative analysis is based on a single fragment of a humanities text, which limits the generalizability of the findings across genres and disciplines of academic writing. Second, the analysis was based on a single generative model and a limited number of runs; therefore, the identified tendencies require further validation using other systems and types of editing prompts. Third, the study adopts a qualitative-analytical approach and does not aim to provide a quantitative assessment of the frequency of the observed shifts.

Список источников / References

1. Богинская, О.А. (2023). Лексико-синтаксические актуализаторы хеджирования в русском языке: опыт анализа отзывов о диссертации. *Русистика*, 21(1), 18–32.
<https://doi.org/10.22363/2618-8163-2023-21-1-18-32>
Boginskaya, O.A. (2023). Lexico-syntactic actualizers of hedging in the Russian language: An analysis of dissertation reviews. *Rusistika*, 21(1), 18–32. (In Russ.).
<https://doi.org/10.22363/2618-8163-2023-21-1-18-32>
2. Викторова, Е.Ю. (2024). Хеджинг-стратегия в оценочном научном дискурсе: междисциплинарное исследование. *Вестник Волгоградского государственного университета. Серия 2, Языкознание*, 23(1), 5—17.

- <https://doi.org/10.15688/jvolsu2.2024.1.1>
Viktorova, E.Yu. (2024). Transdisciplinary study of hedging strategy in evaluative academic discourse. *Science Journal of Volgograd State University. Linguistics*, 23(1), 5—17. (In Russ.). <https://doi.org/10.15688/jvolsu2.2024.1.1>
3. Горина, О.Г., Храброва, В.Е. (2017). Лингвистический хеджинг как коммуникативная структура (в русле корпусных исследований). *Вестник Новосибирского государственного университета. Серия: Лингвистика и межкультурная коммуникация*, 15(3), 44–53. <https://doi.org/10.25205/1818-7935-2017-15-3-44-53>
Gorina, O.G., Khrabrova, V.E. (2017). Linguistic hedging as a communicative structure (within corpus-based research). *Vestnik Novosibirskogo gosudarstvennogo universiteta. Seriya: Lingvistika i mezhkul'turnaya kommunikatsiya*, 15(3), 44–53. (In Russ.). <https://doi.org/10.25205/1818-7935-2017-15-3-44-53>
 4. Зашихина, И.М. (2024). Научные публикации и большие языковые модели: поймет ли нейросеть российскую науку? *Научный редактор и издатель*, 9(1 Suppl. 2), 2S31–2S46. <https://doi.org/10.24069/SEP-24-11>
Zashikhina, I.M. (2024). Scientific publications and large language models: Will a neural network understand Russian science? *Nauchnyi redaktor i izdatel'*, 9(1 Suppl. 2), 2S31–2S46. (In Russ.). <https://doi.org/10.24069/SEP-24-11>
 5. Сахарова, А.В. (2020). Языковые средства выражения объективной эпистемической модальности в научном дискурсе. *Научный диалог*, (4), 151–163. <https://doi.org/10.24224/2227-1295-2020-4-151-163>
Sakharova, A.V. (2020). Linguistic means of expressing objective epistemic modality in scientific discourse. *Nauchnyi dialog*, (4), 151–163. (In Russ.). <https://doi.org/10.24224/2227-1295-2020-4-151-163>
 6. Amirjalili, F., Neysani, M., Nikbakht, A. (2024). Exploring the boundaries of authorship: A comparative analysis of AI-generated text and human academic writing in English literature. *Frontiers in Education*, 9, 1347421. <https://doi.org/10.3389/feduc.2024.1347421>
 7. An, Y., Yu, J.H., James, S. (2025). Investigating the higher education institutions' guidelines and policies regarding the use of generative AI in teaching, learning, research, and administration. *International Journal of Educational Technology in Higher Education*, 22, Article 10. <https://doi.org/10.1186/s41239-025-00507-3>
 8. 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). New York: ACM. <https://doi.org/10.1145/3442188.3445922>
 9. Chelli, M., Shabani, M., Hsu, C., Bender, J. (2024). Accuracy of references generated by ChatGPT and Bard for systematic reviews. *Journal of Medical Internet Research*, 26, Article e53164. <https://doi.org/10.2196/53164>
 10. Desaire, H., Chua, A.E., Isom, M., Jarosova, R., Hua, D. (2023). Distinguishing academic science writing from humans or ChatGPT with over 99% accuracy using off-the-shelf machine learning tools. *Cell Reports Physical Science*, 4(6), Article 101426. <https://doi.org/10.1016/j.xcrp.2023.101426>

11. Ganjavi, C., Eppler, M.B., Pekcan, A., Biedermann, B., Abreu, A., Collins, G.S., Gill, I.S., Cacciamani, G.E. (2024). Publishers' and journals' instructions to authors on use of generative artificial intelligence in academic and scientific publishing: Bibliometric analysis. *BMJ*, 384, Article e077192. <https://doi.org/10.1136/bmj-2023-077192>
12. Geng, M., Trotta, R. (2024). Is ChatGPT transforming academics' writing style? In *Proceedings of the 41st International Conference on Machine Learning (ICML), Workshop "The Next Generation of AI Safety"*. OpenReview. <https://openreview.net/pdf/723496dea6538f9253334c3e923b5b533bd667c6.pdf> (viewed: 27.03.2026)
13. Hassan, A., Barber, S.J. (2021). The effects of repetition frequency on the illusory truth effect. *Cognitive Research: Principles and Implications*, 6, Article 38. <https://doi.org/10.1186/s41235-021-00301-5>
14. Hyland, K. (1998). Hedging in scientific research articles. *Amsterdam: John Benjamins*
15. Ji, Z., Lee, N., Frieske, R., Yu, T., Su, D., Xu, Y., Ishii, E., Bang, Y., Madotto, A., Fung, P. (2023). Survey of hallucination in natural language generation. *ACM Computing Surveys*, 55(12), Article 248. <https://doi.org/10.1145/3571730>
16. Lingard, L. (2020). The academic hedge: Modal tuning in research writing. *Perspectives on Medical Education*, 9(1), 5–8. <https://doi.org/10.1007/s40037-019-00559-y>
17. Maynez, J., Narayan, S., Bohnet, B., McDonald, R. (2020). On faithfulness and factuality in abstractive summarization. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics (ACL 2020)* (pp. 1906–1919). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2020.acl-main.173>
18. Messeri, L., Crockett, M.J. (2024). Artificial intelligence and illusions of understanding in scientific research. *Nature*, 627, 49–58. <https://doi.org/10.1038/s41586-024-07146-0>
19. Millar, N., Batalo, B., Budgell, B. (2022). Trends in the use of promotional language (hype) in abstracts of successful National Institutes of Health grant applications, 1985–2020. *JAMA Network Open*, 5(8), Article e2228676. <https://doi.org/10.1001/jamanetworkopen.2022.28676>
20. Pagnoni, A., Balachandran, V., Tsvetkov, Y. (2021). Understanding factuality in abstractive summarization with FRANK: A benchmark for factuality metrics. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (NAACL-HLT 2021)* (pp. 4812–4829). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2021.naacl-main.383>
21. Peters, U., Chin-Yee, B. (2025). Generalization bias in large language model summarization of scientific research. *Royal Society Open Science*, 12, 241776. <https://doi.org/10.1098/rsos.241776>
22. Ryba, R., Doubleday, Z.A., Dry, M.J., Semmler, C., Connell, S.D. (2021). Better writing in scientific publications builds reader confidence and understanding. *Frontiers in Psychology*, 12, Article 714321. <https://doi.org/10.3389/fpsyg.2021.714321>
23. Sanz-Tejeda, A., Domínguez-Oller, J.C., Baldaquí-Escandell, J.M., Gómez-Díaz, R., García-Rodríguez, A. (2026). The impact of generative AI on academic reading and writing: A

synthesis of recent evidence (2023–2025). *Frontiers in Education*, 10, 1711718.
<https://doi.org/10.3389/educ.2025.1711718>

24. Shanahan, M. (2022). Talking about large language models. *arXiv preprint arXiv:2212.03551*.
<https://doi.org/10.48550/arXiv.2212.03551>
25. Stump, A., Voss, A., Rummel, J. (2024). The illusory certainty: Information repetition and impressions of truth enhance subjective confidence in validity judgments independently of the factual truth. *Psychological Research*, 88, 1288–1297. <https://doi.org/10.1007/s00426-024-01956-7>
26. Vinkers, C.H., Tijdink, J.K., Otte, W.M. (2015). Use of positive and negative words in scientific PubMed abstracts between 1974 and 2014: Retrospective analysis. *BMJ*, 351, h6467.
<https://doi.org/10.1136/bmj.h6467>
27. Zhao, W. (2025). Reconstructing stance in EFL doctoral thesis writing through generative artificial intelligence. *Humanities and Social Sciences Communications*, 12, 1963.
<https://doi.org/10.1057/s41599-025-06249-x>

Информация об авторе

Анна Николаевна Семилетова, кандидат педагогических наук, доцент кафедры педагогической психологии имени профессора В.А. Гуружапова, факультет «Психология образования», Московский государственный психолого-педагогический университет, Москва, Российская Федерация, ORCID: <https://orcid.org/0009-0007-8555-3155>, e-mail: semiletovaan@mgppu.ru

Information about the author

Anna N. Semiletova, Candidate of Science (Pedagogy), Associate Professor, Department of Pedagogical Psychology named after Prof. V.A. Guruzhapov, Faculty of Educational Psychology, Moscow State University of Psychology and Education, Moscow, Russian Federation, ORCID: <https://orcid.org/0009-0007-8555-3155>, e-mail: semiletovaan@mgppu.ru

Конфликт интересов

Автор заявляет об отсутствии конфликта интересов.

Conflict of interest

The author declares no conflict of interest.

Поступила в редакцию 31.03.2026
Поступила после рецензирования 28.05.2026
Принята к публикации 10.06.2026
Опубликована 30.06.2026

Received 2026.03.31
Revised 2026.05.28
Accepted 2026.06.10
Published 2026.06.30