

Политика журнала в отношении искусственного интеллекта¹: на какие вопросы нужно ответить?

Материалы к докладу конференции [НИМУ-2024](#).

Зельдина Марина Михайловна, руководитель проекта Elpub.Образование
zeldina@neicon.ru

[Инструменты с ИИ в издательском деле](#)

[Инструменты с LLM](#)

[Ограничения чат-ботов с LLM](#)

[Что ученые думают об использовании ИИ?](#)

[Как ученые используют чат-боты с LLM?](#)

[Ключевые рекомендации по использованию ИИ при подготовке и публикации статей и политики научных журналов](#)

1 Программа с искусственным интеллектом – это техническая или программная система, способная решать задачи, традиционно считающиеся творческими, принадлежащие конкретной предметной области, знания о которой хранятся в памяти такой системы. Различные системы искусственного интеллекта различаются по уровню автономности и адаптивности после развертывания. Генеративный искусственный интеллект – это подмножество алгоритмов машинного обучения, способных генерировать текст, изображения или другие медиаданные в ответ на подсказки. Генеративный ИИ использует генеративные модели, такие как большие языковые модели, для статистической выборки данных на основе набора обучающих данных, который использовался для их создания. Больше нюансов и различий описаны здесь:

Понятный гайд по ИИ: сравниваем традиционный и генеративный искусственный интеллект

<https://habr.com/ru/companies/itglobalcom/articles/752150/>

Важное уточнение: мы говорим о применении инструментов, взглядах представителей научно-издательского сообщества, но не говорим о технических деталях разработки программ с использованием ИИ, а также об их внедрении в существующее ПО.

ChatGPT стремительно становится именем нарицательным, как «ксерокс» для копировальной машины. ChatGPT в этом материале мы называем все программы, которые используют большие языковые модели, в том числе аналогичные программы, разработанные Яндексом и Сбером. Эти же программы могут быть названы чат-ботами с ИИ, чат-ботами с LLM, программами с генеративным искусственным интеллектом.

[Для чего необходимо разрабатывать политику научного журнала в отношении ИИ?](#)

[Политика журнала в отношении искусственного интеллекта: на какие вопросы нужно ответить?](#)

- [1. Допускает ли редакция использование инструментов с ИИ при проведении исследования и подготовке статьи? Если допускает, то какие это инструменты, в каком объеме и для решения каких задач их можно использовать. Допускается ли использование ИИ при работе с изображениями?](#)
- [2. Может ли ИИ быть автором или соавтором статьи? Может ли ИИ быть указан в перечне лиц, внесших вклад в проведение исследования и подготовку статьи?](#)
- [3. Как автор должен раскрыть информацию об использовании ИИ в статье? Нужно ли приводить информацию о том, что ИИ не был использован? Если да, то как об этом нужно написать? В какой части статьи должна быть приведена эта информация?](#)
- [4. Есть ли исключения для применения существующей политики?](#)
- [5. Может ли рецензент использовать программы с ИИ при подготовке рецензии? Если да, то в каком объеме и для решения каких задач? Какие изменения коснутся работы рецензента после начала действия существующей политики?](#)
- [6. Может ли редактор использовать программы с ИИ при работе со статьей? Если да, то в каком объеме и для решения каких задач?](#)
- [7. В каких случаях автор должен проявить особенную осторожность при использовании ИИ при подготовке статьи?](#)
- [8. При каких условиях автору следует отказаться от использования ИИ при подготовке статьи?](#)
- [9. Каким образом информация о появлении политики в отношении ИИ будет доведена до сведения всех заинтересованных участников редакционно-издательского процесса?](#)

[Шаблон политики в отношении искусственного интеллекта и выдержки из ключевых рекомендаций](#)

[Резюме](#)

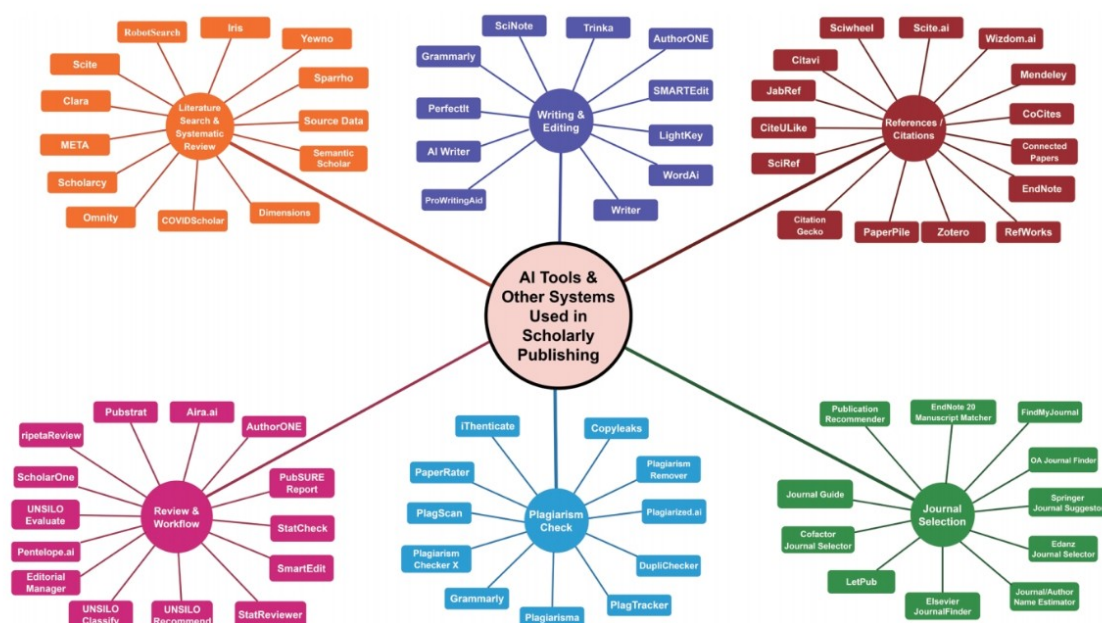
[Список литературы](#)

[Дополнительные материалы](#)

[НЭИКОН](#) связан с темой искусственного интеллекта с 2019 г. - в это время началась работа над проектом [Нейроассистент научного издательства](#), который включает в себя несколько сервисов для авторов и издателей. Нейроассистент подбирает журналы для публикации, ищет потенциальные пропуски источников, рекомендует литературу, подбирает ключевые слова и может оценить наличие ключевых разделов/элементов статьи в файле рукописи. Нейроассистент обучен на данных [платформы Elpub](#).

Инструменты с ИИ в издательском деле

Мы начнем обсуждение искусственного интеллекта в издательском деле, используя иллюстрацию к статье об инструментах ИИ, опубликованную в журнале Science Editing [1].



На этой иллюстрации приведены типы программ с ИИ, которые помогают автоматизировать рутинные процессы как со стороны авторов, так и со стороны редакторов журналов: обзор литературы, поиск информации, анализ данных, работа с текстом рукописи, управление библиографией и цитированием, выбор подходящих для публикации журналов, предотвращение плагиата, рецензирование, оценка качества рукописи, управление редакционными процессами.

К сожалению, большинство этих программ недоступны в России для использования: некоторые являются частью глобальных издательских систем либо требуют оплаты, почти все ориентированы на работу прежде всего с английским языком.

В другой статье [2] авторы рассказывают об инструментах с ИИ, отталкиваясь от задач редакции и используя схему стандартного редакционно-издательского процесса: от выбора журнала до рецензирования, принятия решения о публикации и так далее.

Инструменты с LLM

Развитие ПО, в основу которого положены большие языковые модели (large language models, LLM), в частности, появление осенью 2022 г. и широкого распространение чат-бота с искусственным интеллектом ChatGPT и подобных ему программ, изменили вектор возможных обсуждений программ с ИИ в издательском деле чуть более, чем полностью: на первый план вышло обсуждение этических вопросов, связанных с использованием LLM всех участников редакционно-издательского процесса: авторов, редакторов, рецензентов. А также возникла необходимость очень быстрого обобщения

существующего опыта и разработки рекомендаций для всех участвующих в процессе издания научного журнала сторон.

Почему возникла необходимость действовать быстро? Потому что все чат-боты с LLM широко распространены и просты в использовании, дают ответы на вопросы на любом языке, их стиль письма очень похож, а во многих случаях совершенно неотличим от текста, написанного человеком.

Наконец, с помощью чат-бота с LLM действительно можно решить задачу написания текста, который выглядит, как научный, и потратить на это минимальное количество времени и сил.

Но несмотря на кажущуюся привлекательность, использование чат-ботов с LLM несет в себе серьезные риски.

Люди и машины не всегда могут распознать тексты, сгенерированные ИИ. В исследовании Gao et al., 2022 (цитируется по [3]) рецензенты академических рефератов смогли идентифицировать только 68% контента, созданного ChatGPT, при этом уровень ложноположительных результатов составил 14%.

Даже специализированное ПО для определения сгенерированного текста ошибается: уровень обнаружения составляет всего 54,8% в примерах работ, представленных на платформе Turnitin (Perkins et al., 2023, цитируется по [3]), а неточности в ряде инструментов для обнаружения сгенерированного текста были выявлены и обсуждены Weber-Wulff et al. (2023) (цитируется по [3]).

Обманчивая легкость и доступность для использования этого инструмента вызывает множество опасений в научном сообществе. Ученые проводят исследования, в которых пытаются сформулировать плюсы и минусы использования чат-ботов с LLM применительно к проведению научных исследований, академическому письму, а также определить границы, которые, с одной стороны, позволят программам с LLM помогать ученому в работе (но не заменять его и не снимать с него ответственность), но при этом не будут ставить под сомнение ценность, целостность, воспроизводимость полученного нового знания.

Ограничения чат-ботов с LLM

В статье, которая посвящена SWOT-анализу ChatGPT [4], называют следующие минусы программ с LLM, которые могут серьезно повлиять на результаты проводимых исследований:

1. Ограниченные возможности – неспособность предоставить надежные факты и источники, ограничения в понимании сложных научных концепций, ограниченный объем знаний, отсутствие ответственности.
2. Отсутствие гарантированно правильных ответов и действий.
3. Необходимость дополнительной проверки.
4. «Галлюцинации» - создание несуществующих источников информации.
5. Плагиат – не всегда точное указание источников, пропуск источников.

6. Угроза конфиденциальности (личная информация будет доступна модели ИИ для обучения (запрашивая у ChatGPT проверку и отправляя ей текст статьи или фрагменты, вы не имеете возможности контролировать, как, кому и в каком виде эта информация будет передана).
7. Не существует программ, которые могли бы однозначно определять «руку» ChatGPT в научных статьях; программы для обнаружения сгенерированного текста тоже иногда ошибаются.
8. ChatGPT по умолчанию использует устаревшую информацию, не может генерировать новые идеи и не поможет получить новые результаты (эта характеристика почти перестала быть актуальной - новые версии ChatGPT могут обновлять данные «на лету», но того же нельзя сказать о множестве программ-посредников. Кроме того, пользователь далеко не всегда обращает внимание на версию программы, которую использует, и вряд ли вникает так глубоко в характеристики ее данных).
9. Простота использования и связное изложение результатов запроса может негативно сказаться на способности молодых людей обучаться мыслить и излагать свои мысли.

Ключевые нюансы, о которых всегда необходимо помнить при работе с любым чат-ботом с LLM:

1. ChatGPT никогда не может быть назван автором или соавтором статьи и любого материала, созданного с помощью этой программы. Ответственность за использование всех материалов, полученных в результате взаимодействия с чат-ботом, несет только человек.
2. ChatGPT не умеет говорить «нет». Если программа не знает точного ответа на вопрос, она придумает ответ, выглядящий правдоподобно, но не имеющий никакого отношения к реальности. Это и есть так называемые «галлюцинации».
3. ChatGPT ошибается. Он может признаться в ошибке, но только в случае если ему будет задан конкретный вопрос о ней. Можно ли определить ошибку там, где не ожидаешь ее найти?
4. ChatGPT не компенсирует недостаток опыта человека, который ее использует. Чем меньше опыта у автора в исследуемой теме - тем больше вероятность, что он не сможет найти неточности в предложенном чат-ботом ответе [5].
5. Использование ChatGPT всегда создает риск угрозы конфиденциальности и нарушения прав. Условия общедоступных инструментов с генеративным ИИ часто разрешают повторное использование входных данных в обучении, и любое обучение данные могут случайно или намеренно отображаться в качестве выходных данных из инструмента с генеративным ИИ без соответствующих сообщений о лицензировании или условий распространения [6].
6. Документация работы с ChatGPT при подготовке статьи и планировании исследования - не излишняя нагрузка, а необходимость: многие журналы требуют предоставить не только описание той работы, которая была проделана с помощью ChatGPT, но и приложить запросы к программе и ответы на них. В разные моменты времени ChatGPT может отвечать по-разному на один и тот же вопрос, поэтому необходимо сохранить важную для дальнейшей работы информацию [7].

Что ученые думают об использовании ИИ?

Ученые слишком доверяют ИИ — так [считают](#) антрополог Лиза Миссери и когнитивист Молли Крокетт [8]. Они изучили 100 научных статей, книг, препринтов исследований и материалов конференций. И сделали такой вывод: многие считают, будто ИИ повысит производительность и объективность, преодолев недостатки человека. В реальности же он может, наоборот, усугубить когнитивные искажения, присущие людям.

В основном ученые воспринимают ИИ как технологию для обработки потока знаний, превышающего когнитивные возможности человека, генерации гипотез на их основе, сбора, подготовки и анализа данных, рецензирования научных статей и заявок на гранты. Это может привести к множеству потенциальных проблем, в частности к таким иллюзиям понимания:

- иллюзия глубины понимания — ученый считает, что понимает явление, смоделированное ИИ, лучше, чем на самом деле;
- иллюзия широты исследования — ученый думает, что изучает все проверяемые гипотезы, а на самом деле лишь те, что доступны для проверки инструментами ИИ;
- иллюзия объективности — ученый уверен, что ИИ лишен предвзятости или представляет все возможные точки зрения, в то время как у технологии есть только те данные, на которых она обучалась.

Источник информации: тг-канал: [Science&Health Writing](#), автор Екатерина Кушнир
Разрешение на использование фрагмента поста получено от автора канала

Интересны результаты опроса издательства IOP Press в отношении возможностей использования программ с генеративным ИИ в процессе рецензирования. Около 35% респондентов считают, что инструменты с генеративным ИИ окажут негативное влияние на процесс экспертной оценки, 36% относятся нейтрально либо считают, что генеративный ИИ не окажет никакого влияния на процесс рецензирования. 29% респондентов считают, что такие инструменты окажут положительное влияние на процесс рецензирования [9].

Как ученые используют чат-боты с LLM?

Используя данные Dimensions и опробованный ранее словарь лексики, характерной для ответов чатбота, Эндрю Грей оценил распространенность применения AI в академических текстах 2023 года в 1% (~60000 публикаций) [10].

Результаты исследования подтвердили, что чат-боты с искусственным интеллектом широко

используются в научных исследованиях, хотя многие исследователи представляют свои публикации как собственные работы, не отмечая использования чат-ботов с ИИ [11].

Из более чем 1600 ученых, ответивших на опрос Nature 2023 года, почти 30% заявили, что использовали генеративный ИИ для написания статей, а около 15% заявили, что использовали его для собственных обзоров литературы и для написания заявок на гранты [12].

Тексты, сгенерированные ИИ, используются в различных разделах статьи (наиболее значимые: формулирование гипотез, обобщение выводов). В странах/регионах, в которых английский не является официальным языком, влияние ChatGPT на академические тексты сильнее [13].

Ключевые рекомендации по использованию ИИ при подготовке и публикации статей и политики научных журналов

Необходимость создания рекомендаций по использованию ИИ при работе с научными публикациями отмечают многие ученые и организации.

Ключевые рекомендации на сегодняшний день разработаны COPE, WAME, STM, уточнения в отношении использования инструментов с ИИ включены в рекомендации ICMJE.

[Artificial intelligence \(AI\) in decision making](#), COPE

[COPE Position Statements: Authorship and AI Tools](#)

[WAME RecommendationsonChatbots and Generative ArtificialIntelligencein Relation to Scholarly Publications](#)

[Рекомендации ICMJE.](#)

[Generative AI in Scholarly Communications: Ethical and Practical Guidelines for the Use of Generative AI in the Publication Process](#), STM

Рекомендации STM включают примеры издательских политик в отношении ИИ: AAS, AIP Publishing, Association of Computing Machinery, Cambridge University Press, Elsevier, Emerald, JAMA, MDPI, Springer Nature, Taylor&Francis, Wiley.

В исследовании [14] указаны некоторые расхождения и разные акценты в документах, но в целом подход к использованию ИИ в различных издательствах и организациях схож (сейчас речь о зарубежных рекомендациях, в России подобных официальных рекомендаций нет).

Документы активно используются мировым сообществом. Политики, созданные на основе этих рекомендаций, используются авторитетными международными издательствами.

Например, рекомендации ICMJE используются такими журналами, как NEJM, рекомендации WAME были адаптированы такими журналами, как BMJ, политика Lancet аналогична политике Elsevier [14].

Для более углубленного анализа текстов политик можно использовать данные Таблицы 1 препринта A BIBLIOMETRIC ANALYSIS OF PUBLISHER AND JOURNAL INSTRUCTIONS TO AUTHORS ON GENERATIVE-AI IN ACADEMIC AND SCIENTIFIC PUBLISHING Table 1. Top-100 Publisher Authors Guidelines on Generative Artificial Intelligence Выдержки из политик 100 издателей [15].

Тем не менее, ученые отмечают и пробелы, недостаточно четкие указания и недостаточно четкую проработку деталей в описании некоторых нюансов. Кроме того, разнообразие руководящих принципов может сбивать исследователей с толку.

Также можем отметить, что опубликовано крайне мало информации о реальных кейсах по использованию ИИ при подготовке научных статей (если исключить все упоминания недобросовестного использования: генерация текста вместо автора научной статьи, а не вместе с ним, сокрытие факта использования ИИ).

В ближайшие месяцы 4000 исследователей из различных дисциплин и стран будут обсуждать руководящие принципы, которые могут быть широко приняты в академических изданиях. Последние полтора года эти издания борются с чат-ботами и другими проблемами, связанными с искусственным интеллектом. Группа, стоящая за этим проектом, стремится заменить разрозненный набор текущих рекомендаций единым набором стандартов, представляющим собой консенсус исследовательского сообщества.

Эта инициатива, известная как CANGARU, представляет собой партнерство между исследователями и издателями, включая Elsevier, Springer Nature, Wiley, а также представителей журналов eLife, Cell и The BMJ и отраслевого органа – Комитета по публикационной этике. Группа надеется выпустить окончательный набор рекомендаций к августу, который будет обновляться ежегодно из-за быстрого развития этой технологии, — говорит Джованни Каччамани, уролог из Университета Южной Калифорнии, возглавляющий CANGARU. В рекомендациях будет представлен список способов, которыми авторам не следует использовать большие языковые модели (LLM), на которых работают чат-боты, а также инструкции по раскрытию информации о других видах использования [16].

В качестве иллюстрации приведем пример сравнения ключевых политик с описанием выявленных различий: <https://osf.io/m7sqd>

**Этот фрагмент текста был подготовлен с помощью ChatGPT4o*

Для чего необходимо разрабатывать политику научного журнала в отношении ИИ?

Политика в отношении искусственного интеллекта решает несколько задач как для самого научного журнала, так и для научного сообщества, поскольку не только описывает принятые журналом (и научным сообществом, как мы увидим ниже) нормы использования инструментов с ИИ, но и является в некотором смысле каналом для их распространения. Чем больше ученых будут знать о возможностях и ограничениях инструментов с ИИ, тем больше вероятность того, что по крайней мере непреднамеренных ошибок, связанных с незнанием тонкостей работы с ИИ, удастся избежать или хотя бы снизить их количество. Кроме того, имея четкую политику в отношении ИИ, редакция журнала защищает себя от возможных проблем, связанных с некорректным использованием инструментов с ИИ со стороны авторов.

На сегодняшний день существуют два возможных сценария при разработке политики в отношении ИИ:

1. Полный запрет на использование инструментов с ИИ на всех этапах работы со статьей. Этот запрет будет касаться как авторов, так и рецензентов и редакторов. Если говорить о тех журналах, которые имеют принятые политики в отношении ИИ, этот сценарий встречается крайне редко.
2. Частичный запрет на использование инструментов с ИИ: как правило, содержит ограничение на возможность использования ИИ для работы с изображениями, а также касается редакторов и рецензентов. Никакие конфиденциальные материалы не должны быть переданы чат-боту с генеративным ИИ.

Отсутствие любой политики в отношении ИИ одним из вариантов сценария считать вряд ли можно: отсутствие любой информации о нормах использования ИИ будет трактоваться, скорее, как полная вседозволенность и отсутствие ограничений.

Политика журнала в отношении искусственного интеллекта: на какие вопросы нужно ответить?

1. Допускает ли редакция использование инструментов с ИИ при проведении исследования и подготовке статьи? Если допускает, то какие это инструменты, в каком объеме и для решения каких задач их можно использовать. Допускается ли использование ИИ при работе с изображениями?

Большинство издательств и редакций разрешают использовать генеративный ИИ при подготовке статьи в целом [17], однако почти все налагают запрет на генерацию и редактирование изображений. Исключения допускаются только в случае если получение изображений либо их корректировка с помощью ИИ является частью плана

исследования. В таком случае авторы должны прозрачно описать, какие изменения были внесены. Также может потребоваться предоставление двух версий изображений: оригинальной и измененной с помощью ИИ.

Обязательным условием является полная прозрачность при раскрытии информации об использовании ИИ. Автор может сделать это в разделе “Методы”, “Благодарности” либо описать, какая работа была проделана, во введении.

2. Может ли ИИ быть автором или соавтором статьи? Может ли ИИ быть указан в перечне лиц, внесших вклад в проведение исследования и подготовку статьи?

Никакая программа с ИИ и ни при каких условиях не может быть указана в качестве автора или соавтора статьи. Программа с ИИ не может быть указана в перечне лиц, внесших вклад в проведение исследования и подготовку статьи.

3. Как автор должен раскрыть информацию об использовании ИИ в статье? Нужно ли приводить информацию о том, что ИИ не был использован? Если да, то как об этом нужно написать? В какой части статьи должна быть приведена эта информация?

Автор может сделать это в разделе “Методы” (если ИИ использовался для сбора данных, их анализа, создания изображений), “Благодарности” (если ИИ использовался для работы с языком рукописи) либо описать, какая работа была проделана, во введении. К статье должен быть приложен полный текст с запросами к чат-боту и ответами на них. В статье обязательно должна быть приведена ссылка на это приложение. Результат работы с чат-ботом публикуется вместе со статьей.

Описание работы, проведенной с помощью ИИ, должно включать:

- Название, версию и разработчика используемых инструментов искусственного интеллекта (например, ChatGPT, версия от 25 сентября, на основе GPT-4, разработанная OpenAI).
- Указание на разделы и объем вмешательства инструмента с ИИ (например, «В разделе «Обсуждение» примерно 20 % текста изначально было составлено ИИ»)
- Описание типа и цели сгенерированного контента, который был включен в статью (например, «Текст, сгенерированный ИИ, предназначен для предоставления структурированного резюме, а также основных выводов. Этот сгенерированный контент был позже отредактирован и уточнен авторами, чтобы обеспечить согласованность, точность и актуальность»).
- Описание подсказок/промтов, которые давались программе, вместе с датой/временем (например, ссылку или снимок экрана чата).

4. Есть ли исключения для применения существующей политики?

Как правило, при обсуждении исключений называют: программы для обнаружения некорректных заимствований, ПО для работы со списком литературы (reference manager), инструменты для проверки правописания. Также некоторые журналы указывают, что политика не препятствует использованию инструментов с ИИ для помощи при планировании исследований или разработке методов исследования.

5. Может ли рецензент использовать программы с ИИ при подготовке рецензии? Если да, то в каком объеме и для решения каких задач? Какие изменения коснутся работы рецензента после начала действия существующей политики?

Несмотря на то, что в некоторых статьях можно увидеть предложения, связанные с использованием LLM в рецензировании [2, 4, 9, 12, 18, 19], все без исключения рекомендации и политики издательств четко прописывают: LLM не может использоваться при подготовке рецензии. Фрагменты текста либо рецензии могут нарушить конфиденциальность полученной от автора информации.

Рецензентам следует ознакомиться с политикой издателя в отношении искусственного интеллекта. Когда рецензенты подозревают нарушение этого правила политики, они должны сообщить об этом редактору, занимающемуся рукописью как часть процесса рецензирования. Если есть опасения, что статья либо ее фрагменты были созданы с помощью ИИ, это может быть отмечено в обзоре как фактор, влияющий на его точность и/или пригодность к публикации.

6. Может ли редактор использовать программы с ИИ при работе со статьей? Если да, то в каком объеме и для решения каких задач?

В большинстве случаев редактору не разрешено использовать LLM при работе со статьей. Причина - та же, что и для рецензентов: высокий риск нарушения конфиденциальности при загрузке в чат-бот как фрагмента статьи, так и фрагмента рецензии.

7. В каких случаях автор должен проявить особенную осторожность при использовании ИИ при подготовке статьи?

Автору не следует полностью полагаться на искусственный интеллект. Он должен помнить о том, что передача любых данных чат-боту несет риски, связанные с нарушением конфиденциальности полученных данных - как чужих, так и его собственных.

8. При каких условиях автору следует отказаться от использования ИИ при подготовке статьи?

Если автор не уверен в правильности ответов ИИ, но также и в том случае если автор уверен в рекомендациях ИИ: если источники реальные, точны и актуальны, возможно, лучше прочитать эти оригинальные источники, чтобы самостоятельно проанализировать и понять их, а также выбрать статьи, которые наиболее лучшим образом отвечают цели исследования. Такой вариант лучше чем использование интерпретации чат-бота. По умолчанию предполагается, что авторы принимают все меры для минимизации рисков, связанных с использованием чат-бота с генеративным ИИ: от анонимизации своих данных на получение разрешение на использование данных для передачи, а также проверки полученных результатов. Если эти меры принять невозможно, авторам может быть рекомендовано рассмотреть альтернативные способы использования чат-ботов с ИИ либо отказаться от их использования.

9. Каким образом информация о появлении политики в отношении ИИ будет доведена до сведения всех заинтересованных участников редакционно-издательского процесса?

Редакционная статья, пресс-релиз, новость на сайте, рассылка.

Пример редакционной статьи, в которой презентуется политика журнала в отношении ИИ: [Best Practices for Using AI When Writing Scientific Manuscripts](#)

Шаблон политики в отношении искусственного интеллекта и выдержки из ключевых рекомендаций

Этот [шаблон](#) был разработан автором статьи [Towards an AI policy framework in scholarly publishing](#)

Резюме

В большинстве случаев темы, которые поднимаются в связи с использованием генеративного ИИ, содержат обсуждение прежде всего обсуждения негативных последствий, связанных с его использованием. Конечно, необходим контроль, прозрачность раскрытия информации об использовании и оправданная настороженность, связанная с любым вмешательством не-человека в публикационный процесс.

Однако нельзя не заметить, что конкретных примеров позитивного использования генеративного ИИ описано очень мало. В первые месяцы после появления ChatGPT появилось достаточно много статей, написанных искусственным интеллектом про искусственный интеллект, но подобные примеры сложно назвать руководством к действию. Очень не хватает реальных примеров того, как с помощью генеративного ИИ применительно к области научных публикаций действительно можно упростить себе жизнь (а не талантливо и незаметно обойти принятые научным сообществом нормы публикационной этики).

Знаете ли вы о таких кейсах? Поделитесь с нами! Мы обобщим результаты и опубликуем их в открытом доступе.

Адрес для связи: zeldina@neicon.ru

Список литературы

1. [Artificial intelligence-assisted tools for redefining the communication landscape of the scholarly world](#)
2. [Artificial intelligence to support publishing and peer review: A summary and review](#)
3. [Academic publisher guidelines on AI usage: A ChatGPT supported thematic analysis](#)
4. [A SWOT \(Strengths, Weaknesses, Opportunities, and Threats\) Analysis of ChatGPT in the Medical Literature: Concise Review](#)
5. [Potential of ChatGPT in facilitating research in radiation oncology?](#)
6. [Generative AI in scholarly communications](#)
7. [How to cite ChatGPT](#)
8. [Artificial intelligence and illusions of understanding in scientific research](#)
9. [State of peer-review 2024](#)
10. [ChatGPT “contamination”: estimating the prevalence of LLMs in the scholarly literature](#)
11. [Ethical concerns for using artificial intelligence chatbots in research and publication: Evidences from Saudi Arabia](#)
12. [Is ChatGPT corrupting peer review? Telltale words hint at AI use](#)
13. [Have AI-Generated Texts from LLM Infiltrated the Realm of Scientific Writing? A Large-Scale Analysis of Preprint Platforms](#)
14. [Towards an AI policy framework in scholarly publishing](#)
15. [A BIBLIOMETRIC ANALYSIS OF PUBLISHER AND JOURNAL INSTRUCTIONS TO AUTHORS ON GENERATIVE-AI IN ACADEMIC AND SCIENTIFIC PUBLISHING](#)
16. [Should researchers use AI to write papers? Group aims for community-driven standards](#)
17. [Responsible integration of AI in academic research: detection, attribution, and documentation](#)
18. [The dangers of using large language models for peer review](#)

19. [Fighting reviewer fatigue or amplifying bias? Considerations and recommendations for use of ChatGPT and other large language models in scholarly peer review](#)

Дополнительные материалы

[Конференция «Обнаружение заимствований»](#)

[Круглый стол «Искусственный интеллект: автор, соавтор, партнер, инструмент»](#)

[Искусственный интеллект в высшем образовании: практические кейсы](#)

[Искусственный интеллект в издательском деле. Вебинар Elpub](#)

[Круглый стол «Может ли GPT "прирастить" научное знание?»](#)

[Научная этика в эпоху ИИ](#)

[Elpub и Антиплагиат: вместе для качества научных](#)

[публикаций](#)

[Рекомендации Антиплагиат по работе со сгенерированными текстами](#)

[Перспективы применения ChatGPT для высшего образования: обзор международных исследований](#)

[«Мы оба с ним как будто из металла, но только он — действительно металл», или Как перестать беспокоиться и начать использовать генеративные модели ИИ](#)