

Использование искусственного интеллекта для генерации текста и его распознавание

Полина Дмитриевна Асманова

Студент

Нижегородский государственный педагогический университет им. Козьмы Минина

Нижний Новгород, Россия

asmanovapolina14@gmail.com

ORCID 0000-0000-0000-0000

Александр Викторович Поначугин

Кандидат экономических наук, доцент

Нижегородский государственный педагогический университет им. Козьмы Минина

Нижний Новгород, Россия

Ponachygin_AV@mininuniver.ru

ORCID 0000-0001-5518-5565

Иван Александрович Галкин

Студент

Нижегородский государственный педагогический университет им. Козьмы Минина

Нижний Новгород, Россия

ivan.galkin13@gmail.com

ORCID 0000-0000-0000-0000

Поступила в редакцию 05.03.2025

Принята 25.04.2025

Опубликована 30.05.2025

УДК 004.9(075.8)

DOI 10.25726/y2405-6382-9286-u

EDN VPOMAS

BAK 5.8.1. Общая педагогика, история педагогики и образования (педагогические науки)

OECD 05.03.HA. EDUCATION & EDUCATIONAL RESEARCH

Аннотация

В статье рассматриваются современные подходы к генерации текстов с использованием искусственного интеллекта (ИИ) и методам их распознавания. Во введении (Introduction) обозначена актуальность проблемы: рост применения ИИ для создания контента в образовании, медиа и бизнесе требует разработки инструментов для идентификации машинных текстов, что связано с вопросами академической честности, авторского права и борьбы с дезинформацией. В разделе методов (Methods) проанализированы технологии генерации текста на основе нейросетей, включая обучение моделей на больших корпусах данных, а также методы распознавания – статистический анализ (равномерное распределение слов, шаблонные конструкции), алгоритмы машинного обучения и специализированные программы. Результаты (Results) демонстрируют, что тексты ИИ часто характеризуются предсказуемостью синтаксиса и отсутствием эмоциональной вариативности, однако усложнение моделей (например, внедрение ошибок для имитации человеческого стиля) затрудняет их детекцию. Обсуждение (Discussion) подчеркивает этические и юридические вызовы: риски мошенничества, необходимость законодательного регулирования и защиты персональных данных. Отдельное внимание уделено перспективам гибридных решений, сочетающих автоматическую генерацию с человеческой редактурой, а также разработке мультимодальных систем, интегрирующих текст, аудио и визуальные

данные. Статья акцентирует, что, несмотря на прогресс ИИ, ключевая роль в контроле качества и творческом осмыслении сохраняется за человеком.

Ключевые слова

искусственный интеллект, генерация текста, распознавание текста, машинное обучение, этика ИИ.

Введение

В последние годы искусственный интеллект (ИИ) стал очень важной частью нашей жизни, особенно широко он применяется в банковском секторе экономики, в точных науках, лингвистике, биологии, медицине, психологии, машиностроении и в других областях. Он проник уже в самые разные сферы человеческой деятельности – от создания контента до автоматизации задач. На основе информационной помощи машинного обучения вычислительные системы имитируют человеческий интеллект. Искусственный интеллект может обнаружить хакерские атаки и взломы и автоматически принять необходимые меры для устранения таких угроз.

Благодаря стремительному развитию искусственный интеллект подходы к производству и потреблению информации претерпели значительную трансформацию, так как ИИ многократно упрощает работу с большими массивами информации.

Одной из интересных и актуальных тем является распознавание текстов, созданных данной технологией. Генерацию текстовой информации с помощью ИИ освоили уже даже школьники, в связи с этим встает ряд вопросов: как можно отличить статью, написанную человеком, от сгенерированного алгоритмом текста? Сложно ли это сделать, и какие технологии и принципы для этого используют? Цель данного исследования – рассмотреть, как проходит распознавание ИИ, каковы методы, использующиеся для этого, каковы вызовы и перспективы развития искусственного интеллекта.

Материалы и методы исследования

Первая информация о зарождении искусственного интеллекта появилась очень давно – в середине XX века. Одной из главных предпосылок создания искусственных нейронных сетей является идея нейрофизиолога У. Маккалока и У. Питтсома в статье 1943 года «Логическое исчисление идей, относящихся к нервной активности», опубликованной в «Бюллетене математической биофизики». В 1957 году Ф. Розенблатт создал перцептронный алгоритм распознавания образов, который впоследствии был реализован на компьютере. В XXI веке ценность нейросетей возросла благодаря использованию распределенных вычислений и графических процессоров (ПикOVER, 2023).

В 1950 году уже А. Тьюринг опубликовал статью «Вычислительные машины и разум» (Computing Machinery and Intelligence) в журнале «Mind», логической основой которой стала теория о том, что машины, как и люди, могут использовать общедоступную информацию, а также разум для решения проблем. Сейчас каждый студент знакомится с этим понятием, изучая дисциплину «Теория алгоритмов».

Искусственный интеллект развивается достаточно быстро и сейчас используется почти во всех сферах человеческой деятельности: в науке, медицине, торговле, образовании, промышленности и т.д.

Результаты и обсуждение

Так зачем же нам нужно распознавание текстов ИИ? Сейчас это становится все более популярным. Благодаря достижениям в обработке естественного языка различные системы могут создавать тексты, статьи, эссе и даже книги. В связи с этим возникает проблематика интеллектуальной собственности, защиты авторских прав и необходимости различать тексты, написанные людьми, от сгенерированных. В образовательных учреждениях сейчас остро стоит эта проблема, ведь очень важно убедиться, что работа выполнена студентом самостоятельно, а не с помощью машины. Поэтому каждая работа проверяется на рекомендованных министерством образования источниках.

Одной из важнейших проблем современности является кибербезопасность, так как сегодня ИИ активно используется в различных мошеннических схемах – создаются фальшивые новости, фэйки в

целях подрывной пропаганды в социальных сетях. Ярким примером является один из gpt-чатов, который «выдумывает» не только новости, то и рассказы, диалоги, анализируя текст на 8 млн веб-страниц. Умение распознавать подобные тексты помогает бороться с распространением дезинформации и обеспечить честные и прозрачные формы коммуникации. Это особенно важно в связи с ростом количества ботов, которые могут генерировать тысячи сообщений и вводить в заблуждение пользователей, да и в целом общество.

В этой связи упомянем и об этической стороне вопроса. Если человек читает текст и не понимает, кто является автором, то может встать вопрос о доверии, ответственности и понимания целей, с которыми создавался контент. Чтобы понять, как распознать сгенерированный текст, нужно сначала разобраться, каким образом происходит его генерация. Процесс включает в себя несколько шагов:

1. Обучение. Вначале ИИ изучает большое количество текстов – книги, статьи, новости, посты и другую информацию, созданную человеком.
2. Осознание текста. На втором этапе ИИ выявляет закономерности, образцы, структуру текста, чтобы уловить связь и запомнить какие-то шаблоны.
3. Генерация текста. Технология получает запрос в виде небольшого текстового отрывка – начала текста, обрабатывает информацию и выдает продолжение на основе полученного знания о запросе на основе огромных массивов информации по заданной теме.
4. Редактирование и оценка. Если робот «выдал» не совсем то, что желал «заказчик», чтобы улучшить качество текста, запрос дополняется, и ИИ снова генерирует результат.

При этом стоит отметить, что машина не имеет своей личности – она просто обрабатывает некое количество текстов и старается сделать что-то похожее, имитируя творчество человека. Далее рассмотрим методы распознавания сгенерированных текстов. Они включают в себя ряд методов, позволяющих определить, кто является автором текста:

1. Статистический анализ. Тексты, написанные искусственным интеллектом, часто характеризуются равномерным распределением слов и частым использованием шаблонных фраз. Синтаксические структуры могут быть слишком простые или, наоборот, очень сложными. Грубо говоря, если вы пишете очень складно, то, скорее всего, «вы» – искусственный интеллект. Для статистического анализа используются определенные модели, которые учитывают параметры, такие как частота слов, структура предложений, использование синонимов, повторяющиеся шаблоны. Все это отличает сгенерированные технологий тексты от авторски, которые в массе своей более спонтанны и имеют свои лексические изъясны, не такие логичные и стройные, а иногда даже противоречивые. То есть, проще говоря, в своих «работах» ИИ очень предсказуем. Тексты человека включают в себя: разнообразные фразы, стилистические особенности, выражения, ошибки. Человек зависим от своего настроения или характера, который не имеет наш «робот», он может писать «под настроение»: зол – очень агрессивно и кратко, влюблен – со множеством красочных эпитетов и выражений, повторами и «ненужными» отступлениями. Важным признаком является креативность текста: у машины нет творческих способностей, ИИ использует лишь уже имеющиеся знания для создания контента. Человек способен на оригинальные идеи, которые не всегда поддаются логике.

2. Машинное обучение. Машинное обучение можно встретить везде. Трудно представить себе хоть один аспект жизни, который нельзя было бы так или иначе улучшить с помощью машинного обучения. В любом деле, требующем повторения, анализа данных и сбора выходных данных, вам поможет машинное обучение. За последние несколько лет, благодаря достижениям в области вычислительных машин и повсеместному сбору данных, в этой сфере наблюдался поистине взрывной рост. В область применения машинного обучения сегодня входят системы рекомендаций, распознавание изображений, обработка текста, автопилоты в автомобилях, выявление спама, постановка медицинских диагнозов и т.д. и т.п., список можно продолжать бесконечно. Это обучение выявляет характерные признаки, присущие текстам. Для обучения таких моделей используют наборы данных, в которых есть и человеческие, и машинные тексты. После получения и обработки данных сеть может анализировать новый текст и определять его происхождение. Этот метод достигает высоких результатов в распознавании. Например, нейронные сети могут выявлять логические несоответствия или необычные

конструкции, типичные для стилистики написания текста определенного индивидуума. Они также могут выявлять повторяющиеся паттерны в использовании слов и фраз.

3. Специальные программы. Для распознавания текстов искусственного интеллекта создается все больше и больше сайтов и приложений. Но проблема заключается в том, что разные сайты в зависимости от их качества показывают разный результат. Поэтому существуют определенные стандарты, которым следуют, например, учебные заведения. Идея сама по себе хорошая, но пока программы нужно дорабатывать для более точного определения. Ведь люди начали сталкиваться с тем, что сервисы даже текст, написанный человеком, называют сгенерированным. И человеку приходится притворяться человеком, чтобы текст не засчитали как созданный ИИ.

Перейдем к сложностям распознавания. Обнаружение ИИ становится все более точным, но существуют ограничения. Во-первых, модели становятся все более изощренными. То есть они способны адаптировать стиль текста в зависимости от контекста. Например, некоторые алгоритмы сознательно встраивают в текст ошибки или используют вариативные конструкции, имитируя человека. Это затрудняет не только их нахождение, но и жизнь человека, который пишет самостоятельно статью или эссе. Ведь как бы он не старался, программа может сказать ему, что он «компьютер».

Во-вторых, тексты могут корректироваться человеком, что еще больше усложняет их распознавание. Такой подход использования алгоритма и изменения текста человеком снижает возможность его обнаружения. Человек может маскировать характерные признаки машинной генерации.

Другая сложность заключается в том, что современные ИИ-модели обучены на больших количествах данных и способны адаптироваться, что позволяет им создавать тексты, которые практически невозможно отличить от человеческих. И чем совершеннее алгоритмы, тем труднее разработчикам.

Но перспективы в данной сфере все-таки есть. Одним из таких направлений являются гибридные решения. Они будут объединять машинное обучение с человеческой интуицией и способностью к творчеству. Может быть, появятся технологии, которые будут отслеживать весь процесс создания текста – например временные метки, исправление и др. В целом же, безусловно, проблема распознавания текстов сейчас стоит очень остро и важность ее решения будет расти с развитием искусственного интеллекта. Будем надеяться, что в скором будущем появятся новые технологии, которые помогут решить данную проблему.

Стоит также отметить, что ИИ постепенно становится неотъемлемой частью многих технологий, связанных с обработкой естественного языка, поскольку современные алгоритмы уже способны создавать осмысленные тексты и интерпретировать их значение на достаточно высоком уровне. Многие исследователи считают, что развитие систем генерации и распознавания текста способствует ускоренному прогрессу в сфере автоматизации различных предпринимательских процессов, научных изысканий и даже творческих задач. При этом применение машинного обучения и нейронных сетей оказывает значительное влияние на то, как именно формируются тексты, а также на методы дальнейшего анализа их смысловых и стилистических особенностей. Трудности, с которыми сталкиваются разработчики, часто связаны не только с необходимостью обеспечить корректность и логичность машинного текста, но и с поиском эффективных способов «понимания» структуры языка алгоритмом.

В контексте современных реалий многие компании и научные организации активно инвестируют в исследования, связанные с моделями глубокого обучения, которые позволяют работать не только с синтаксисом, но и с более скрытыми смыслами. Подобная тенденция влечет за собой появление множества приложений: от чат-ботов, способных поддерживать диалог на естественном языке, до автоматизированных систем перевода, способных быстро и довольно точно интерпретировать различные варианты письменных обращений (Байтяков, 2024). В то же время наблюдается явный рост числа опубликованных исследовательских работ, посвященных сравнительному анализу методов генерации текста, а также способам улучшения распознавания сложных лингвистических структур (Харченко, 2023). В результате все больше специалистов рассматривают искусственный интеллект как универсальный инструмент, способный удовлетворить потребности в быстрой генерации контента и

одновременном его понимании, что стимулирует прогресс в когнитивной науке и вычислительной лингвистике. Постепенно стирается грань между творчеством человека и работой алгоритма, и это, судя по многим прогнозам, будет в значительной степени влиять на рынок труда будущего.

Совершенствование подобных технологий также вызывает дискуссии высокого уровня, связанные с корректностью представления информации и возможными рисками искажения смысла при автоматической обработке текста. Тем не менее невероятный потенциал систем ИИ открывает беспрецедентные возможности для создания инновационных продуктов и услуг, основанных на генерации и понимании текстов. Именно синергия человеческого творчества и мощности вычислительных моделей формирует тот самый вектор развития, который позволяет надеяться на высокую степень адаптации ИИ к меняющимся условиям окружающей среды. Все это подкрепляется гибкостью методов машинного обучения, которые легко настраиваются под различные языки и области применения, расширяя охват и глубину анализа. Так закладываются основы новой эры, где информационное взаимодействие становится все более автоматизированным, но при этом сохраняется свойственная языку человечность, отражающая культурные особенности и эмоциональное восприятие действительности.

Современные модели генерации текста опираются на огромные корпуса данных, позволяющие обучать нейронные сети определенным паттернам построения предложений, лексическим сочетаниям и множеству прочих элементов лингвистической структуры. Важнейшую роль при этом играет компонент обучения с учителем или без учителя, когда алгоритму предоставляется массив текстовой информации для выявления закономерностей, а также формирования представлений о логике построения предложений. Модели глубокого обучения, такие как трансформеры, стали революционным прорывом, благодаря которому искусственный интеллект научился улавливать контекст на более широких временных интервалах. Этот контекст позволяет моделям генерировать новые предложения, которые не только стилистически схожи с оригинальными, но и семантически адекватны. Так, например, алгоритмы могут создавать тексты, имитирующие литературные стили, научные статьи или даже юридические документы, что существенно упрощает рутинную деятельность человека. Подобные возможности мотивируют многих разработчиков исследовать применение генеративных моделей в области творческих индустрий, торговли, маркетинга и образования (Айдагулова, 2023).

Однако требовательность к вычислительным ресурсам, необходимым для обучения таких моделей, продолжает расти, поэтому не все организации могут позволить себе развернуть полноценные инфраструктуры на базе больших нейронных сетей. Несмотря на это, появляются новые подходы к дистилляции моделей, которые уменьшают их размер без потери качества, облегчая разворачивание в прикладных задачах.

Одновременно с масштабированием и оптимизацией моделей генерации стимулируется разработка алгоритмов, способных детектировать машинный текст, чтобы предотвратить злоупотребление технологиями для дезинформации или спама. На практике это приводит к тому, что появляется целый спектр инструментов, ориентированных на проверку оригинальности текстов и их классификацию по степени возможного участия машинного интеллекта. Обнаружение следов искусственного происхождения текста превращается в серьезную исследовательскую задачу, ведь язык обладает бесконечным количеством нюансов, которые не всегда поддаются формализации (Бручес, 2022). Следовательно, методы тестирования и распознавания генеративных моделей вынуждены постоянно совершенствоваться и усложняться, чтобы идти в ногу с прогрессирующими возможностями самих моделей. С этой точки зрения взаимодействие науки и рынка становится все более активным: крупные компании, занимающиеся рекламой, СМИ и предоставлением цифровых услуг, стремятся к тому, чтобы обнаружить случай несанкционированного использования алгоритмов, а университеты и лаборатории прикладывают усилия к углубленным исследованиям лингвистической природы машинного текста. Стоит также заметить, что, несмотря на высокую эффективность некоторых моделей, генерация полностью достоверной и неотличимой от человеческой речи информации пока остается сложной задачей.

Исследования генерации речи из текстового формата и обратного процесса, при котором речь распознается и преобразуется в текст, тесно связаны с прогрессом в области обработки естественного языка, поскольку и там, и там задействована лингвистическая сущность. Наиболее успешные современные решения активно используют рекуррентные нейронные сети в сочетании с механизмами внимания для анализа и синтеза звуковой информации. При такой схеме алгоритмы способны разбирать фонемные составляющие, слова, ударения и тон речи, а в результате получается реалистичная синтетическая речь, которая может сопровождаться почти человеческой интонацией. Многие крупные компании демонстрируют успехи в развитии систем распознавания голоса, способных анализировать речь самых разных людей, учитывая их индивидуальные особенности, акценты и даже дефекты произношения. В то же время все активнее обсуждается вопрос создания универсальных мультязычных систем, способных работать практически с любым языком мира. Эти усилия способствуют тому, что технологический барьер на пути глобального общения снижается, и стоит лишь немного усовершенствовать инфраструктуру, и люди получают уникальную возможность моментального перевода устной и письменной речи без особых трудностей (Бондаренко, 2023). Это касается не только личного общения, но и бизнес-процессов, при которых переговоры, подписания контрактов и презентации могут проходить в смешанном лингвистическом пространстве, более не сдерживаемом традиционными рамками владения определенным языком.

Конечно, без привлечения искусственного интеллекта к анализу семантики и контекста такие системы не могли бы достичь нынешнего уровня эффективности. Машинное распознавание подразумевает, что алгоритм должен не только выделять слова, но и понимать их синтаксические и контекстные связи, иначе теряется смысл. Существует ряд методик, позволяющих оценивать точность распознавания: от простых тестов на словарный запас до комплексных испытаний, учитывающих шумы, смешение стилей и разную громкость. Нередки ситуации, когда системы сталкиваются со сложными ударными нюансами или специфической терминологией, которая требует предварительной адаптации алгоритма. Именно поэтому тема исследований в области улучшения архитектуры нейронных сетей остается актуальной. Чем дальше будет развиваться эта сфера, тем выше шансы на появление универсальных решений, способных корректно понимать и воспроизводить речь любого человека.

Еще один аспект, неразрывно связанный с генерацией и распознаванием текстов, касается стилистической составляющей. Машине недостаточно просто сформировать правильные предложения или определить набор слов из голосового потока, необходимо научиться улавливать стилистику, эмоциональную окраску, а иногда и культурный подтекст (Пименова, 2023). Различные языки имеют множество регистров общения, жаргоны и социальные диалекты, которые могут варьироваться даже в пределах одной страны. Именно поэтому системы, ориентированные на реальное общение, обязаны учитывать разнообразие и богатство лингвистического ландшафта. Алгоритмы, натренированные на узком наборе данных, рискуют выдавать тексты, которые покажутся носителям языка излишне формальными или, напротив, пренебрежительными. В лучшем случае человек улыбнется такому курьезу, в худшем – может возникнуть социальное напряжение и недопонимание. Проблема усугубляется в случаях, когда речь идет об официальном документообороте, где малейшая неточность в передаче тонкостей формулировки способна поменять юридический смысл. В таких ситуациях оказывается важным включение в процесс контроля качества и редактуры, осуществляемой профессионалами, которые через свое экспертное знание исправляют ошибки алгоритма. С другой стороны, разрабатываются инструменты, используемые журналистами и писателями, позволяющие быстро набрасывать черновики и затем дорабатывать их под нужный стиль. Набирают популярность также приложения, в которых генерация текстов на основе стиля конкретного автора позволяет пользователям наслаждаться шутливыми результатами. И хотя эти продукты не всегда имеют высокую точность, они привлекают внимание массового пользователя к технологиям ИИ, повышая их популяризацию.

Помимо чисто лингвистических вопросов значительная часть дискуссий об искусственном интеллекте для генерации и распознавания текстов касается этических и юридических аспектов. Чем более совершенными становятся алгоритмы, тем сильнее повышается вероятность обмана – выдачи

машинно сгенерированных материалов за человеческое творчество (Киселев, 2024). Информация, распространяемая в сети Интернет, может выглядеть убедительно и реалистично, в то время как на самом деле она была создана программой, не обладающей чувством ответственности или критического осмысления. Подобная практика может использоваться для создания фейковых новостей, пропагандистских материалов или попыток дискредитировать отдельных лиц, организаций либо стран. Это провоцирует вопросы о необходимости правового регулирования и внедрения механизмов прозрачности, которые позволяли бы отслеживать источники данных и идентифицировать возможные сфальсифицированные тексты. В некоторых странах уже начинаются попытки поставить законодательные барьеры для подобных технологий, однако проблема усугубляется глобальным характером сети и разными подходами к свободе слова. Из-за этого общество часто разделяется на сторонников свободного развития ИИ и защитников более жесткого регулирования, которые указывают на риски, связанные с социальной ответственностью разработчиков.

Не менее важный аспект – защита конфиденциальности. При обучении больших языковых моделей используются огромные массивы данных, в том числе тексты, содержащие личную информацию пользователей. Потребность в эффективной анонимизации и защите баз данных встает особенно остро в тех случаях, когда речь идет о сервисах, обрабатывающих миллионы запросов в сутки. В итоге многие компании отдают приоритет комплексным исследованиям в области криптографии и дифференциальной приватности, чтобы сократить риск утечек и обеспечить репутационную безопасность.

С развитием систем генерации текста растет и интерес к вопросу о том, сможет ли искусственный интеллект когда-нибудь заменить полностью работу профессиональных авторов. Ведь если алгоритм уже сейчас способен выдавать связный текст, соответствующий определенному стилю, то стоит ли продолжать привлекать человеческих специалистов там, где можно сэкономить на зарплатном фонде и времени? На первый взгляд, эта идея кажется весьма притягательной, и часть компаний действительно внедряет роботов для создания новостных заметок или аналитических отчетов (Пименова, 2022). Тем не менее эксперты констатируют, что творческий элемент и уникальный взгляд на проблему во многом остаются за человеком. Машина, обученная на больших объемах документации, действительно может составить членораздельный и логически выверенный текст, но ей зачастую недостает глубины эмоционального осмысления и способности к оригинальному творчеству, которое выходит за рамки обучающего корпуса. Таким образом, профессиональные писатели, журналисты и копирайтеры все еще востребованы, особенно в нишах, где необходима тонкая работа с художественным словом или критический анализ ситуации.

Более вероятен вариант, при котором их роль эволюционирует: возникает синергия между ручной правкой и автоматической генерацией. Алгоритм помогает быстро подготовить черновую версию, а человек придает ей ту самую человечность, которая ценится читателями. Парадокс заключается в том, что при массовой автоматизации более сложными становятся и требования к человеческим специалистам, которым нужно разбираться в технологиях ИИ и уметь эффективно взаимодействовать с ними. В конце концов, именно человек контролирует процесс, определяя, какие данные давать алгоритму, как формулировать запросы, какие результаты принимать или отвергать и т.д..

Особое место в исследованиях занимает проблема мультимодальности, когда системы искусственного интеллекта работают не только с текстом, но и с изображениями, звуками, видео. Такой подход потенциально открывает путь к созданию универсальных алгоритмов, способных анализировать комплексную информацию и формировать более развернутые тексты с учетом визуального или аудиального контекста. Представьте, что система «смотрит» на изображение и способна не только наделять его деталями, но и придумывать целую историю, связанную с увиденным (Абрамова, 2020). Аналогичным образом можно объединять данные из видеопотока и аудиосигналов, чтобы обеспечивать всесторонний анализ происходящего, а затем превращать его в сжатое или развернутое описание. Такие технологии уже испытываются в области медицины, где врачам требуется быстрая расшифровка результатов обследований в текстовом формате, понятном для междисциплинарных консультаций. Кроме того, мультимодальные алгоритмы применяются в автоматическом создании субтитров для

фильмов и передач, а также в образовании, когда необходимо объединить видео- и аудиоинформацию для более эффективного обучения. Однако здесь встает еще более острый вопрос о надежном распознавании нюансов, поскольку мультимодальные системы должны уметь обрабатывать разные типы данных с согласованной точностью. Если алгоритм неправильно проинтерпретирует визуальную деталь, то итоговый текст будет содержать некорректные элементы, что может исказить смысл. Учитывая, что объемы цифрового контента удваиваются каждые несколько лет, мультимодальные решения становятся одним из главных направлений роста, обещая пользователям более гибкие и интеллектуально насыщенные продукты.

Одним из традиционных вызовов для технологий распознавания и генерации текста является проблема неоднозначности языков, ведь одно и то же слово или фраза могут иметь несколько различных интерпретаций в зависимости от контекста (Подтележникова, 2024). Для человека зачастую не составляет труда воспользоваться своими знаниями о реальном мире, чтобы разглядеть верный смысл, но машине приходится полагаться на статистику встречаемости слов и тщательный анализ окружающих токенов. Поэтому современные исследования все чаще обращаются к разработке моделей, учитывающих внешние мировые знания, чтобы повысить точность понимания смысла при генерации или распознавании текста. Это может включать в себя подключение баз знаний или онтологий, а также методы обучения, в которых алгоритм «читает» большое количество тематической литературы, чтобы сформировать более полное представление о предметной области. Сформировав контекст, алгоритм умеет более точно различать, когда слово, например, выступает в роли существительного, а когда – метафоры или слова с переносным значением. Так достигается более высокая адекватность итогового результата, особенно в научных и технических текстах, где неточность в терминах может существенно исказить весь смысл. Тем не менее сложность и вариативность человеческого языка такова, что всегда могут возникнуть фразы, которые даже носителю языка покажутся двусмысленными или требующими уточнения. Машине, при всем ее прогрессе, сложно решать подобные задачи без подсказки человека.

Рассматривая вопросы обучения моделей генерации и распознавания, нельзя не коснуться темы корректного отбора обучающей выборки и методов предобработки данных. В случае текстов крайне важно, чтобы данные отражали как можно более широкий спектр стилей, жанров и тематик, иначе алгоритм рискует научиться шаблонизировать и выдавать текст, неудобочитаемый для разных групп пользователей (Широких, 2023). Если же набор обучающих данных хорошо сбалансирован, система сможет обеспечить более универсальную и гибкую работу. С другой стороны, некоторые специалисты высказывают опасения, что алгоритмы могут усвоить предубеждения и стереотипы, содержащиеся в обучающем корпусе, и затем тиражировать их в своих ответах. Например, если в текстах статистически больше упоминаний о соотношении определенных профессий с мужчинами, алгоритм неосознанно начнет поддерживать подобную ассоциацию. Поэтому важным направлением остается разработка методов выявления и минимизации подобной нежелательной предвзятости, что способствует созданию более этичных и объективных систем.

Параллельно с этим встал вопрос об энергоэффективности, ведь обучение крупных моделей требует колоссальных затрат электроэнергии и вычислительных ресурсов. Организации, стремящиеся к «зеленым» технологиям, пытаются сокращать углеродный след, выбирая более экономичные методы оптимизации и архитектуры, не уступающие по качеству более громоздким моделям. Появляются специализированные аппаратные решения: графические процессоры, тензорные процессоры и нейроморфные чипы, которые лучше приспособлены для вычислительных операций в области машинного обучения (Шарапова, 2023). Все это формирует целый рынок высокопроизводительных вычислений, который одновременно поддерживает и стимулирует развитие ИИ, позволяя ему шагать в ногу с ростом задач. Впрочем, не менее успешно развиваются проекты по распределенным системам обучения, где множество устройств объединяются в глобальные сети для совместной обработки данных. Закрытым или открытым способом формируются своеобразные пулы вычислительной мощности, что особенно актуально для задач, связанных с генерацией и анализом больших объемов текстов в реальном времени.

Говоря о распознавании текста, стоит также упомянуть, что многие технологии сегодня ориентированы на работу с рукописным вводом или сканированными документами, где нужно не только корректно идентифицировать символы, но и учитывать особенности почерка автора, неидеальное качество сканирования, наклон строчной линии и другие факторы (Касатиков, 2022). Здесь на первый план выходит применение сверточных нейронных сетей, способных обучаться на разнообразных наборах изображений букв и слов. Правда, в зависимости от языка могут меняться и правила написания, что вынуждает адаптировать алгоритмы для каждой письменности. Многоязычные системы распознавания текста постепенно набирают популярность в глобальном масштабе, а вместе с ними приходят новые проблемы, например, корректная идентификация знаков, похожих между собой, но имеющих разное значение в разных письменностях. Все это свидетельствует о том, что задачи, связанные с перемещением информации из бумажного или рукописного вида во все более цифровой формат, еще долго будут оставаться актуальными и потребуют дальнейших инноваций.

Таким образом, в основе революции, которую мы наблюдаем, лежит колоссальный прогресс в области вычислительной мощности и доступности данных. Если раньше для отработки модели требовались недели или месяцы непрерывных расчетов, сегодня многие процессы можно выполнить за считанные дни или даже часы, располагая нужным оборудованием и облачными сервисами. Обучение больших языковых моделей перестает быть уделом лишь крупнейших корпораций, так как появляются инициативы по открытым наборам данных, доступным научным коллективам и стартапам. Происходит демократизация технологий, и вслед за ней растет число кросс-дисциплинарных проектов, интегрирующих алгоритмы генерации и распознавания текста в здравоохранение, финансы, логистику и промышленность. Рыночная конкуренция лишь подстегивает этот рост, вынуждая компании предлагать все новые сервисы, улучшающие пользовательский опыт и экономящие ресурсы (Абрамцов, 2023).

С другой стороны, не стоит забывать о сложностях, связанных с оценкой качества создаваемых алгоритмов. Автоматические метрики вроде BLEU, ROUGE или METEOR, используемые для оценки машинного перевода или сгенерированных текстов, хотя и дают представление о приблизительном качестве, не всегда отражают глубину смысла и стилистическую адекватность. Поэтому возникает потребность во все более совершенных методах оценки, которые могут учитывать не только лексическое совпадение, но и смысловые связи, а также соответствие контексту (Глухарев, 2024). В ряде случаев используют ручную оценку филологами, лингвистами и экспертами по тематике, что, однако, связано с большими затратами времени и возможностей человека. Возможно, в будущем мы увидим новые алгоритмы «оценки ИИ», которые сами будут способны судить о качестве текста, распознавать художественную ценность или аналитическую глубину ответа.

Поскольку мы говорим не только о генерации, но и о распознавании, важно отметить, что качественное распознавание текста напрямую влияет на успех всех последующих этапов работы с информацией. Ошибка при оцифровке документа может исказить фактологические данные, а некорректное понимание команды голосового ассистента приведет к неправильному выполнению операции. Именно поэтому компании, занимающиеся IoT и умными системами, уделяют столько внимания улучшению точности голосовых интерфейсов. Если пользователь почувствует, что голосовой помощник слишком часто ошибается или не понимает контекст, интерес к продукту резко упадет. Напротив, при высокой точности люди начинают воспринимать эти технологии как естественную часть повседневной рутины. Происходит эффект привыкания, когда взаимодействие с ИИ становится настолько обыденным, что люди перестают замечать техническую составляющую, считая ее неотъемлемой функцией окружения.

В будущих сценариях развития, вероятно, мы столкнемся с системами, которые смогут не только писать и распознавать текст, но также активно обучаться на собственном опыте в режиме реального времени. Это означает, что подобные алгоритмы начнут приспосабливаться к стилю общения конкретных пользователей, их задачам и предпочтениям. Дополнятся механизмы самокоррекции, позволяющие машине обнаруживать ошибки в своих собственных выводах и исправлять их, как если бы человек делал «черновик». В результате интеллектуальные ассистенты станут более проактивными, предлагая целевые улучшения в формулировках или варианты решения проблемы уже в ходе беседы,

а не по окончании ее (Бручес, 2022). Такой уровень интерактивности сформирует нечто вроде «языковой интроспекции», когда машина сможет объяснить, почему она выбрала тот или иной вариант, ссылаясь на внутренние представления, сформированные в ходе обучения.

Следует также ожидать все более тесной интеграции систем генерации текста с системами управления знаниями, корпоративными базами данных и экспертными системами. Если соединить большой массив структурированных данных с возможностями языковых моделей, то можно быстро формировать аналитические отчеты, резюмирующие ключевые показатели или указывающие на аномалии в больших наборах статистики. Взаимодействие с такими системами может происходить в диалоговом режиме, когда пользователь просто пишет запрос своими словами, а ИИ формирует развернутый ответ, ссылаясь на факты из базы данных. Данный подход уже меняет подходы к бизнес-аналитике и прогнозированию, значительно снижая технический барьер для использования сложных аналитических инструментов (Харченко, 2023).

Обратная сторона стремительного развития – возникновение новых видов угроз информационной безопасности. Киберпреступники могут применять генеративные модели для написания качественных фишинговых писем, вводящих в заблуждение даже достаточно опытного пользователя. Кроме того, распространяются схемы мошенничества, где с помощью синтезированной речи подделывается голос руководителя компании, дающий указания о переводе денежных средств. Такого рода атаки уже регистрируются правоохранительными органами, а компании пытаются защитить себя, усиливая проверку аутентификации. В перспективе может потребоваться повсеместное внедрение технологий цифровой подписи и биометрических датчиков, способных удостоверять личность говорящего. Тем не менее в условиях бурного роста ИИ, инструменты защиты тоже должны развиваться не менее динамично, и эта гонка обещает быть весьма длительной (Пименова, 2023).

Вопросы приватности, защищенности и качественного контроля сочетаются с необходимостью постоянно обновлять нормативно-правовую базу, которая в различных юрисдикциях сильно отстает от реальной практики применения технологий искусственного интеллекта. Законодатели только начинают осознавать, что запрет или сильное ограничение одной страны может быть легко обойдено в другой, поскольку интернет не знает границ, а цифровые сервисы мигрируют туда, где условия для них выгоднее. Отсюда возникает риск того, что единые стандарты так и не будут приняты в ближайшее десятилетие, и каждая страна станет двигаться по собственному пути. Однако глобальные организации уже ставят цель по разработке рекомендаций, которые позволят хотя бы частично согласовать общие принципы. Ведь подобно тому, как существуют нормы и правила в отношении авторского права, возможно, будут выработаны нормы и относительно происхождения контента и мониторинга генеративных систем.

Научное сообщество между тем продолжает делать ставку на открытость и обмен знаниями, что в перспективе может ускорить решение многих проблем. Публикуются сотни работ, посвященных обучению и тестированию моделей языка, их распознаванию, анализу вводных данных, методам оптимизации и многому другому (Айдагулова, 2023). Исследователи стремятся к тому, чтобы как можно больше инструментов было доступно в открытом коде, чтобы ускорить коллаборацию и повысить воспроизводимость экспериментов. Разумеется, коммерческие структуры, стремясь сохранить конкурентное преимущество, не раскрывают всех деталей и побуждают к созданию гибридных форм взаимодействия с научными лабораториями.

Параллельно гуще становится рынок стартапов, которые ищут способы реализовать прогрессивные идеи на базе искусственного интеллекта, будь то автоматическое написание писем, аналитика больших текстовых массивов, интеллектуальные чат-боты для технической поддержки или системы мониторинга социальных медиа. Инвесторы с интересом вкладываются в такие проекты, видя, что рынок растет экспоненциально и каждое успешное решение может принести серьезную прибыль. Но при этом остается фактор неопределенности, ведь не каждая разработка найдет свою нишу, а конкуренция уже становится достаточно плотной. Поэтому ключом к успеху нередко оказывается качество команды, ее научная база и техническая экспертиза в области машинного обучения.

В практическом применении генерация и распознавание текста уже показали свою пользу в системах онлайн-поддержки клиентов, где чат-боты отвечают на типовые вопросы, оставляя сложные запросы реальным операторам. Экономия времени и средств налицо, а удовлетворенность пользователей порой только повышается из-за скорости ответа. В образовании использование подобных технологий позволяет автоматизировать формирование тестов, проверку орфографии, создание учебных материалов. Некоторые университеты уже внедряют системы, способные быстро проанализировать письменные работы студентов на предмет оригинальности и стилистической грамотности (Пименова, 2023). В сфере торговли продвинутые алгоритмы формируют описания товаров, пишут рекламные объявления или даже ведут динамическую коммуникацию с клиентами, повышая конверсию продаж.

Нельзя не отметить, что в науке распознавание научных статей и генерация конспектов исследований тоже играют значительную роль. При сегодняшнем объеме публикаций учеными по всему миру обработка вручную становится невозможной, поэтому системы машинного обучения приходят на помощь, помогая находить релевантные источники, выделять ключевые точки и даже составлять рефераты (Бондаренко, 2023). Это существенно облегчает поиск информации и стимулирует междисциплинарное сотрудничество, так как исследования из смежных областей могут быть проанализированы и представлены в доступной форме.

Несмотря на все успехи, присутствует четкое понимание, что универсальный искусственный интеллект, который мог бы полностью заменить человека во всех аспектах генерации и распознавания текста, пока не существует. Мы имеем дело с инструментами, которые узко или более широко специализируются на конкретных подзадачах, и зачастую им недостает разностороннего понимания реального мира. Часто алгоритмы могут попадать в ловушки, связанные с непонятными им культурными отсылками или неочевидными контекстами, которые для человека являются частью обыденных знаний (Харченко, 2023). Это означает, что ручная корректировка, дополнение или надзор эксперта еще долго будут обязательной частью процесса.

В то же время очевидно, что ИИ уже успел существенно изменить подход к работе с текстом: создал возможность автоматизировать рутинные задачи, упростил коммуникации между людьми, говорящими на разных языках, и предоставил мощные инструменты для анализа огромных массивов данных. Перспективы дальнейшего развития кажутся безграничными, особенно учитывая синтез мультимодальных подходов и растущую вычислительную мощь. Можно предположить, что следующие поколения алгоритмов будут еще более адаптивными, способными к самообучению и самодиагностике, а значит, повысится и уровень их автономности.

Технологические гиганты уже готовы инвестировать значительные средства в любые инициативы, связанные с глубоким обучением языковых моделей, так как видят в этом ключ к развитию цифровых ассистентов, роботизированных систем, умных интерфейсов и многих других областей. В конечном итоге, выгоду получают и простые пользователи, которые смогут решать сложные задачи быстрее и эффективнее, опираясь на интеллектуальные подсказки алгоритмов. Одновременно с этим возникнут новые трудности, связанные с этикой, безопасностью, правовым регулированием и сохранением человеческой уникальности.

Заключение

Подводя итог, отметим, что использование искусственного интеллекта для генерации текста и его распознавания сегодня уже не выглядит чем-то фантастическим или далекой перспективой, а, напротив, становится частью повседневной жизни, выступая в качестве совершенно различных систем и ресурсов – от голосовых помощников до систем сопровождения бизнеса. Дискурс смещается с вопроса «Возможно ли это?» к обсуждению «Как это использовать ответственно и эффективно?». Вероятно, мы станем свидетелями быстрого роста и расширения инструментов, способных писать, говорить и понимать «как люди», но границы между человеком и алгоритмом при этом вряд ли будут стерты окончательно. Центральной фигурой, направляющей и контролирующей развитие технологий, формируя тем самым будущее информационного общества на перекрестке инноваций и гуманизма, останется Человек.

Список литературы

1. Абрамова А.В. Распознавание текста при подготовке электронных документов // Лаборатория культуры: молодежные практики и новации: мат. XVIII науч.-творч. конф. студ. СГИК. Самара, 2020. С. 11.
2. Абрамцов А.А. Искусственный интеллект по генерации текста для правовой сферы: перспективы и трудности // Новые технологии – нефтегазовому региону: мат. Межд. науч.-прак. конф. студ., асп. и мол. уч. В 2-х томах. Под ред. В.А. Чейметовой. Тюмень, 2023. С. 140-142.
3. Айдагулова А.Р. Особенности текстов, сгенерированных искусственным интеллектом // Вестник Башкирского государственного педагогического университета им. М. Акмуллы. 2023. № 4(72). С. 154-156.
4. Байтяков Н.А., Мухачев С.В. Распознавание текстов, сгенерированных искусственным интеллектом // Информационные технологии и когнитивная электросвязь: мат. X Всеросс. науч.-прак. конф. Екатеринбург, 2024. С. 11-13.
5. Бондаренко М.В. Обработка связанного текста с помощью нейросети для создания изображений // Угрозы и риски на Юге России в условиях геополитического кризиса. Достижения и перспективы научных исследований молодых ученых Юга России: мат. Всеросс. конф. с межд. уч.; XIX Ежег. мол. науч. конф. Ростов-на-Дону, 2023. С. 299.
6. Глухарев В.А. Определение методов и требований к созданию суммаризатора текста в сфере искусственного интеллекта // Ломоносовские научные чтения студентов, аспирантов и молодых ученых – 2024. Сборник материалов конференции: в 2-х томах. Архангельск, 2024. С. 155–159.
7. Касатиков Н. Н., Фадеева А. Д., Пархаев В. А. Применение алгоритмов нейронных сетей для автоматической генерации по запросу в текстовом формате // Авиация и космонавтика. Тезисы 21-ой международной конференции. Московский авиационный институт (национальный исследовательский университет). М., 2022. С. 221–222.
8. Киселев И.Н. О применении искусственного интеллекта в распознавании текстов // Вестник ВНИИДАД. 2024. № 1. С. 84–95.
9. Пикова К. Искусственный интеллект: от автоматов до нейросетей: иллюстрированная история. Пер. с англ. А. Ефимовой. М.: Синдбад, 2023. 224 с.
10. Пименова О.В., Галиев Д.В., Жиганин А.Я. Классификация текстовой информации при помощи искусственного интеллекта // Авиация и космонавтика. Тезисы 22-й Межд. конф. М., 2023. С. 166.
11. Пименова О.В., Галиев Д.В., Жиганин А.Я. Семантический анализ текстов на основе искусственного интеллекта // Авиация и космонавтика: мат. XXI Межд. конф. М.: Московский авиационный институт, 2022. С. 254–255.
12. Подтележникова Е.Н. Инструменты искусственного интеллекта для анализа языка и текста. Воронеж: ИД ВГУ, 2024. 85 с.
13. Харченко Б.В., Нестругина Е.С. Искусственный интеллект как инструмент для работы с текстами // Вестник Донецкого национального университета. Серия Г: Технические науки. 2023. № 4. С. 29-40.
14. Шарапова Е.В. Признаки текстов, сгенерированных с применением искусственного интеллекта // Машиностроение и безопасность жизнедеятельности. 2023. № 3(47). С. 58-62.

Using artificial intelligence to generate text and recognize it

Polina D. Asmanova

Student

Nizhny Novgorod State Pedagogical University named after Kozma Minin

Nizhny Novgorod, Russia

asmanovapolina14@gmail.com

ORCID 0000-0000-0000-0000

Alexander V. Ponachugin

Candidate of Economic Sciences, Associate Professor
Nizhny Novgorod State Pedagogical University named after Kozma Minin
Nizhny Novgorod, Russia
Ponachygin_AV@mininuniver.ru
ORCID 0000-0001-5518-5565

Ivan A. Galkin

Student
Nizhny Novgorod State Pedagogical University named after Kozma Minin
Nizhny Novgorod, Russia
ivan.galkin13@gmail.com
ORCID 0000-0000-0000-0000

Received 05.03.2025

Accepted 25.04.2025

Published 30.05.2025

UDC 004.9(075.8)

DOI 10.25726/y2405-6382-9286-u

EDN VPOMAS

VAK 5.8.1. General pedagogy, history of pedagogy and education (pedagogical sciences)

OECD 05.03.HA. EDUCATION & EDUCATIONAL RESEARCH

Abstract

The article discusses modern approaches to text generation using artificial intelligence (AI) and methods of their recognition. The Introduction highlights the urgency of the problem: the growing use of AI for content creation in education, media and business requires the development of tools for identifying machine texts, which is related to issues of academic integrity, copyright and the fight against misinformation. The Methods section analyzes neural network-based text generation technologies, including model training on large data sets, as well as recognition methods such as statistical analysis (uniform word distribution, template constructions), machine learning algorithms, and specialized programs. The Results demonstrate that AI texts are often characterized by predictable syntax and lack of emotional variability, however, the complexity of models (for example, the introduction of errors to simulate human style) makes it difficult to detect them. The Discussion highlights ethical and legal challenges: the risks of fraud, the need for legislative regulation and the protection of personal data. Special attention is paid to the prospects of hybrid solutions combining automatic generation with human editing, as well as the development of multimodal systems integrating text, audio and visual data. The article emphasizes that, despite the progress of AI, the key role in quality control and creative thinking remains with humans.

Keywords

artificial intelligence, text generation, text recognition, machine learning, AI ethics.

References

1. Abramova A.V. Text recognition in the preparation of electronic documents // Laboratory of culture: youth practices and innovations: mat. of the XVIII Scien. and creative conf. of stud. SGIK. Samara, 2020. p. 11.
2. Abramtsov A.A. Artificial intelligence for text generation for the legal sphere: prospects and difficulties // New technologies for the oil and gas region: mat. of the Inter. scien. and prac. conf. stud., asp. and young scientists. In 2 vols. Ed. by V.A. Cheymetova. Tyumen, 2023. pp. 140-142.

3. Aydagulova A.R. Features of texts generated by artificial intelligence // Bulletin of the Bashkir State Pedagogical University named after M. Akmulla. 2023. No. 4(72). pp. 154-156.
4. Baityakov N.A., Mukhachev S.V. Recognition of texts generated by artificial intelligence // Information technologies and cognitive telecommunications: mat. of the X Vseross. scien. and prac. conf. Yekaterinburg, 2024. pp. 11-13.
5. Bondarenko M.V. Processing linked text using a neural network to create images // Threats and risks in the South of Russia in the context of the geopolitical crisis. Achievements and prospects of scientific research of young scientists of the South of Russia: mat. of the All-Russ. conf. with inter. particip.; XIX Annual Inter. scien. conf. Rostov-on-Don, 2023. p. 299.
6. Glukharev V.A. Definition of methods and requirements for creating a text summarizer in the field of artificial intelligence // Lomonosov scientific readings of students, postgraduates and young scientists – 2024: mat. of the conf. In 2 vols. Arkhangelsk, 2024. pp. 155-159.
7. Kasatnikov N. N., Fadeeva A.D., Parkhaev V. A. Application of neural network algorithms for automatic generation on request in text format // Aviation and Cosmonautics: mat. of the 21st Inter. conf. Moscow M.: Moscow Aviation Institute, 2022. pp. 221-222.
8. Kiselev I.N. On the use of artificial intelligence in text recognition // VNIIDAD. 2024. № 1. pp. 84-95.
9. Pikover K. Artificial intelligence: from automatons to neural networks: an illustrated history. Trans. from eng. by A. Efimova. M.: Sinbad, 2023. 224 p.
10. Pimenova O.V., Galiev D.V., Zhiganin A.Ya. Classification of textual information using artificial intelligence // Aviation and cosmonautics: mat. of the 22nd Inter. conf. M., 2023, p. 166.
11. Pimenova O.V., Galiev D.V., Zhiganin A.Ya. Semantic analysis of texts based on artificial intelligence // Aviation and Cosmonautics: mat. of the XXI Inter. conf. M.: Moscow Aviation Institute, 2022. pp. 254-255.
12. Podtelevnikova E.N. Artificial intelligence tools for language and text analysis. Voronezh: Voronezh State University Publishing House, 2024. 85 p.
13. Kharchenko B.V., Nesrugina E.S. Artificial intelligence as a tool for working with texts // Donetsk National University bulletin. Series G: Technical Sciences. 2023. № 4. pp. 29-40.
14. Sharapova E.V. Signs of texts generated using artificial intelligence // Mechanical engineering and life safety. 2023. № 3(47). pp. 58-62.