



Научная статья

УДК 81`42

<https://doi.org/10.37493/2409-1030.2024.4.23>**ФУНКЦИОНАЛЬНО-СЕМАНТИЧЕСКИЕ, МАКРОСТРУКТУРНЫЕ И ПРЕСУППОЗИЦИОНАЛЬНО-ПРАГМАТИЧЕСКИЕ ПАРАМЕТРЫ СГЕНЕРИРОВАННЫХ РУССКОЯЗЫЧНЫХ ТЕКСТОВ В ЛИНГВИСТИЧЕСКИХ НЕЙРОСЕТЯХ GIGACHAT, CHATGPT МАРТИ И ЯНДЕКС АЛИСА****Сергей Викторович Гусаренко^{1*}, Марина Константиновна Гусаренко²**^{1, 2} Северо-Кавказский федеральный университет (д.1, ул. Пушкина, Ставрополь, 355017, Российская Федерация)¹ Доктор филологических наук, профессор
sgusarenko@mail.ru; <https://orcid.org/0009-0000-9245-2255>² Кандидат филологических наук, доцент
mkgusarenko@mail.ru; <https://orcid.org/0009-0005-9312-8621>

* Автор, ответственный за переписку

Аннотация. Введение. Цель статьи состоит в определении и сопоставлении лингвистических характеристик коротких русскоязычных текстов разных жанров, сгенерированных в языковых нейросетях GigaChat, ChatGPT Марти и Яндекс Алиса. Актуальность исследования состоит в том, что изучение лингвистических характеристик сгенерированных текстов позволило сделать выводы о таких характеристиках названных языковых нейросетей, как способность к построению микротестов по заданным в промпте семантическим параметрам, способность избирать контекстуально релевантные значения слов в тематическом наборе дефиниций, способность к построению текста критической интерпретации высказывания. **Материалы и методы.** В качестве материала исследования избраны порожденные названными выше нейросетями языковые выражения и короткие тексты разной функциональной принадлежности – от предложения и семантической дефиниции слова до текста-обоснования собственного ответа нейросети. В качестве основных использовались метод макроструктурного анализа, метод лексико-семантического анализа, метод грамматического анализа, метод стилистического анализа, метод семантико-прагматического анализа. **Анализ.** Исследование проводилось по следующему плану: 1) анализ сгенерированных дефиниций слов и предложений, построенных нейросетями из этих дефиниций, 2) анализ сгенерированных контекстуальных дефиниций, 3) анализ сгенерированных коротких текстов на предмет их функционально-семантической адекватности. **Результаты.** Работа с тематически связанными дефинициями, сгенерированными названными нейросетями, позволила установить, что данные языковые модели в состоянии согласовывать определения слов с контекстом, не являющимся собственно текстом, то есть могут без специального задания в промте, но исходя из перечня слов в нём, определять тему и давать дефиниции по этой теме. В ходе изучения способности языковых нейросетей давать оценку категориальной и референциальной достоверности высказы-

званий установлено, что все три нейросети оказались в состоянии дать правильные мотивированные ответы, за одним исключением, когда нейросеть указала нехватку информации. В ходе изучения текстов, сгенерированных названными языковыми нейросетями, были выявлены пять основных типов нарушений (дефектов) в них, которые могут квалифицироваться как типичные для этих нейросетей: 1) нарушения логико-семантических связей в тексте, выполнение ложных семантических операций; 2) нарушения бытийных прагматических presuppositions (знания о мире, о свойствах предметов); 3) нарушения коммуникативно-прагматических правил речевого поведения; 4) грамматические отклонения в управлении и согласовании; 5) макроструктурные нарушения.

Ключевые слова: текст, генеративный, языковая нейросеть, семантический, прагматический, дефект, галлюцинация

Для цитирования: Гусаренко С. В., Гусаренко М. К. Функционально-семантические, макроструктурные и presuppositional-pragmatic parameters of generated Russian-language texts in the linguistic neural networks GigaChat, ChatGPT Marty and Yandex Alice. 2024. Т. 11. № 4. С. 788–800. <https://doi.org/10.37493/2409-1030.2024.4.23>

Конфликт интересов: один из авторов статьи – доктор филологических наук, профессор Гусаренко Сергей Викторович является членом редакционной коллегии журнала «Гуманитарные и юридические исследования». Авторам неизвестно о каком-либо другом потенциальном конфликте интересов, связанном с этой рукописью.

Вклад авторов: все авторы внесли равный вклад в подготовку публикации.

Статья поступила в редакцию: 14.03.2024.

Статья одобрена после рецензирования: 18.05.2024.

Статья принята к публикации: 25.07.2024.

Research article

FUNCTIONAL-SEMANTIC, MACROSTRUCTURAL AND PRESUPPOSITIONAL-PRAGMATIC PARAMETERS OF GENERATED RUSSIAN-LANGUAGE TEXTS IN THE LINGUISTIC NEURAL NETWORKS GIGACHAT, CHATGPT MARTY AND YANDEX ALICE**Sergey V. Gusarenko^{1*}, Marina K. Gusarenko²**^{1, 2} North-Caucasus Federal University (1, Pushkina St., 355017 Stavropol, Russian Federation)¹ Dr. Sc. (Philology), Professor
sgusarenko@mail.ru; <https://orcid.org/0009-0000-9245-2255>² Cand. Sc. (Philology), Associate Professor
mkgusarenko@mail.ru; <https://orcid.org/0009-0005-9312-8621>

* Corresponding author

Abstract. Introduction. The purpose of the article is to determine and compare the linguistic characteristics of short Russian-language texts of different genres generated in the language neural networks GigaChat, ChatGPT Marty and Yandex Alice. The relevance of the study is that the study of the linguistic characteristics of the generated texts allowed us to draw conclusions

about such characteristics of the named language neural networks as the ability to build microtests based on the semantic parameters specified in the prompt, the ability to select contextually relevant meanings of words in a thematic set of definitions and the ability to build a text of critical interpretation of the statement. **Materials and Methods.** The material for the study was

© Гусаренко С. В., Гусаренко М. К., 2004

linguistic expressions and short texts of different functional affiliation generated by the above-mentioned neural networks – from a sentence and a semantic definition of a word to a text produced by the neural network itself. The following methods were used as the main ones: macrostructural analysis, lexical-semantic analysis, grammatical analysis, stylistic analysis and semantic-pragmatic analysis. **Analysis.** The study was conducted according to the following plan: 1) analysis of generated definitions of words and sentences constructed by neural networks from these definitions, 2) analysis of generated contextual definitions, 3) analysis of generated texts for their functional-semantic adequacy. **Results.** Working with thematically related definitions generated by the above-mentioned neural networks made it possible to establish that these language models are able to coordinate definitions of words with a context that is not the text itself, that is, they can, without a special assignment in the prompt, but based on the list of words in it, determine the topic and give definitions on this topic. In the course of studying the ability of language neural networks to assess the categorical and referential reliability of statements, it was found that all three neural networks were able to give correct motivated answers, with one exception, when the neural network indicated a lack of information. In the course of studying the texts generated by the named language neural networks, five main types of violations (defects) were identified in them, which can be qualified as typical for these neural networks:

Введение. Цель статьи состоит в определении и сопоставлении лингвистических характеристик коротких русскоязычных текстов разных жанров, сгенерированных в языковых нейросетях GigaChat [9], ChatGPT Марти [8] и Яндекс Алиса [7]. В ходе исследования рассматривались такие характеристики текстов, как пропозициональная правильность предложений [1], общая функционально-семантическая и макроструктурная адекватность текстов [4], пресуппозиционально-прагматическая адекватность [5]. Изучение этих параметров позволило сделать выводы о таких лингвистических характеристиках названных языковых нейросетей, как способность к построению микротестов по жестко заданным в промпте семантическим параметрам, способность избирать контекстуально релевантные значения слов в тематическом наборе дефиниций, способность к построению текста критической интерпретации высказывания, стилистические возможности БЛМ. Исследование не имело целью выявить какие-либо преимущества одних нейросетей перед другими.

Актуальность предпринятого исследования обусловлена прежде всего тем, что в настоящее время для генерации текстов самых разных жанров, свойств и объемов потребителями используется целый ряд наиболее распространенных и доступных языковых нейросетей (Large Language Model, или LLM; Большие лингвистические модели, или БЛМ), и такое положение дел вызывает необходимость изучения лингвистических свойств порождаемых текстов, напрямую зависящих от генеративных возможностей применяемых языковых моделей. Как оказалось, тексты одной жанровой принадлежности, но сгенерированные разными языковыми моделями, существенно различаются как по общим лингвистическим параметрам, так и по своим дефектам [3]. Поэтому анализу подвергались как указанные лингвистические параметры текстов, так и выяв-

1) violations of logical-semantic connections in the text, the implementation of false semantic operations; 2) violations of existential pragmatic presuppositions (knowledge about the world, about the properties of objects); 3) violations of communicative-pragmatic rules of speech behavior; 4) grammatical deviations; 5) macrostructural violations.

Keywords: text, generative, language neural network, semantic, pragmatic, defect, hallucination

For citation: Gusarenko SV, Gusarenko MK. Functional-semantic, macrostructural and presuppositional-pragmatic parameters of generated Russian-language texts in the linguistic neural networks GigaChat, ChatGPT Marty and Yandex Alice. Humanities and law research. 2024;11(4):788-800. (In Russ). <https://doi.org/10.37493/2409-1030.2024.4.23>

Conflict of interest: one of the authors of article – Sergey V. Gusarenko, Dr. Sc. (Philology), Professor, is a member of Editorial Board of journal “Humanities and law research”. The authors are not aware of any other potential conflict of interest relating to this manuscript.

Contribution of the authors: the authors contributed equally to this article.

The article was submitted: 14.03.2024.

The article was approved after reviewing: 18.05.2024.

The article was accepted for publication: 25.07.2024.

ленные в них дефекты, поскольку изучение дефектов также может дать важную информацию о возможностях той или иной языковой модели.

Материалы и методы. В качестве материала исследования избраны порожденные нейросетями GigaChat, ChatGPT Марти и Яндекс Алиса языковые выражения и короткие тексты разной функциональной принадлежности – от предложения и семантической дефиниции слова до текста-обоснования собственного ответа нейросети. Такая шкала возрастающей структурной и семантической сложности позволяет более системно подходить к исследованию текстов, сгенерированных языковыми моделями, и текстопорождающих (лингвистических) возможностей этих моделей. Короткие тексты могут рассматриваться как архетипические структуры, лежащие в основе текстов других функциональных типов (сценарии, рекламные тексты, рассказы).

Проводился анализ на предмет способности генерировать пропозициональные структуры предложений из заданного набора лексических единиц, в том числе с целью выявления способности соблюдать правила лексической сочетаемости. Исследовались возможности нейросетей генерировать контекстуально обусловленные дефиниции слов по заданной тематике. Анализировались тексты ответов-интерпретаций на вопросы о категориальной достоверности высказываний, поскольку важной является способность языковых нейросетей определять пресуппозиционально адекватные и неадекватные высказывания, именно эта способность может оказаться необходимой для текущего контроля генерируемых текстов во избежание галлюцинаций.

В качестве основных использовались метод макроструктурного анализа, метод лексико-семантического анализа, метод грамматического анализа, метод стилистического анализа, метод семантико-прагматического анализа. Основанное на этих видах анализа комплексное описание

параметров сгенерированных текстов позволит составить представление о том, как устроены тексты этого типа с точки зрения стандартной языковой семантики, включая макроструктурную и синтаксическую семантику.

Тексты, сгенерированные языковыми нейросетями, неизбежно становятся объектами тщательного изучения, но при этом в большинстве статей внимание уделяется преимущественно следующим аспектам: совершенствование больших языковых моделей для генерирования текстов [11; 18], устранение так называемых галлюцинаций в текстах [12; 15; 16; 19], распознавание текстов, созданных нейросетями или с их помощью [17]. Все выявленные дефекты и несовершенства полученных текстов интерпретируются, как правило, в русле этих направлений. Между тем, как уже говорилось выше, внимания заслуживают сами по себе лингвистические характеристики сгенерированных нейросетями текстов – и в первую очередь потому, что тексты эти уже стали языковой реальностью, частью нашей коммуникативной среды и культуры.

Так, Ruixiang Tang, Yu-Neng Chuang, Xia Hu, исследуя пути распознавания сгенерированных нейросетями текстов, в статье «The Science of Detecting LLM-Generated Texts» [17] выделяют следующие их общие семантические характеристики: тексты, созданные нейролингвистическими моделями, «менее эмоциональны и объективны по сравнению с текстами, созданными человеком, которые часто используют пунктуацию и грамматику для передачи субъективных ощущений» [17, р. 19]. Как пример, авторы указывают, что если люди довольно часто используют восклицательные знаки, вопросительные знаки, многоточия в экспрессивных целях, то LLM, как правило, не прибегают к таким средствам. Авторы также отмечают, что на уровне предложений тексты, написанные человеком, обладают более высокой степенью связности, чем сгенерированные тексты.

Здесь мы особо отметим, что данная статья опубликована в июне 2023 г., поэтому, с учётом невероятных темпов развития БЛМ, некоторые из описанных авторами особенностей текстов могут уже проявиться менее значительно.

Из более поздних работ (январь 2024) следует назвать статью Oluwaseyi J., Odu A. «Exploring models that learn the structure and semantics of language to generate coherent text» [14], в которой авторы, анализируя текстопорождающие возможности языковых нейросетей, отмечают, что такие модели, как очень популярная GPT, в частности GPT-3, демонстрируют высокие возможности в создании связного и контекстуально релевантного текста, они хороши в создании текста, сходного с созданным человеком, путём обучения нейросетей на больших текстовых данных. Авторы подчеркивают, что качественные БЛМ демонстрируют способность понимать и воспроизводить тонкие языковые структуры.

Oluwaseyi J. и Odu A., в частности, отмечают, что внедрение внешних источников данных в БЛМ имеет следствием более широкое понимание концепций и сущностей, помогает генерировать более достоверные и последовательные тексты. Применение конкретных подсказок или контекста обеспечивает семантическую согласованность и контекстуальную значимость сгенерированного текста.

При этом в статье «Exploring models that learn the structure and semantics of language to generate coherent text» подчеркивается, что БЛМ, обученные на основе искаженных наборов данных, могут воссоздавать эти искажения в сгенерированных текстах, из чего следует, что при обучении БЛМ крайне необходимы анализ и контроль наборов данных, точная настройка моделей и тщательный анализ источников входных данных.

Отмечается также, что языковые нейросети испытывают трудности при интерпретации языковых выражений, содержащих имплицитур: метафоры, сарказм, юмор, контекстуально под-разумеваемая информация и т.п.

Цвигун Т.В. и Черняков А.Н. в статье «Хармс vs НейроХармс: нейросеть как лаборатория нарратива» [6] рассматривают нарративы, сгенерированные по образцу текстов Д. Хармса трансформерной нейросетью «Порфирьевич», как произведения, наделенные определённой художественной ценностью. Авторы делают вывод, что такого рода тексты могут использоваться для изучения закономерностей текстопорождения. Авторы отмечают такое явление, как достраивание нарратива причинно-следственными операторами («поэтому», «так что», «вот и», «действительно» и т.п.), что устраняет «нарушения» в оригинальном тексте-источнике. Исследуемая авторами нейросеть развивает формальную связность нарратива посредством замены выражения «хорошие люди» в тексте-источнике на местоимение «они», а также при помощи грамматических средств, к примеру сказуемыми-глаголами в форме множественного числа в соответствии с формой эллиптического подлежащего.

Авторы называют эту особенность текстопостроения нарративной стратегией нейросети, хотя, на наш взгляд, такое именование вызвано неоправданной антропоморфизацией языковых нейросетей. Здесь мы всего лишь можем говорить, что в текстах, на которых обучалась модель, статистически преобладают указанные автором способы создания формальной связности нарратива, а не о нарративной стратегии. Если бы нейросеть обучалась на официально-деловых текстах, то «хорошие люди» заменялись бы выражениями типа «указанные выше лица», а если бы на текстах программы «Спокойной ночи малыши», то субститутами оказались бы «добрые дядечки и тёточки».

Марголина А. и Колмогорова А. в статье «Exploring Evaluation Techniques in Controlled Text Generation: A Comparative Study of Semantics and Sentiment in ruGPT3large-Generated and Human-Written Movie Reviews» [14] приводят результаты сравнения кинокритик, написанных реальными людьми, с рецензиями, сгенерированными нейросетями. Авторы отмечают, что рецензии, написанные человеком, лучше оценивают технические аспекты производства фильма, в то время как рецензии, созданные нейросетями, лучше отражают особенности сюжета и персонажей.

Как видим, при всем многообразии подходов и изучаемых объектов, макроструктурные, функционально-семантические и коммуникативно-прагматические свойства текстов, сгенерированных языковыми нейросетями, все ещё остаются недостаточно изученными.

Анализ. Исследование проводилось по следующему плану:

1. Анализ сгенерированных дефиниций слов и предложений, построенных нейросетями из этих дефиниций.

2. Анализ сгенерированных контекстуальных дефиниций.

3. Анализ сгенерированных текстов на предмет их функционально-семантической адекватности.

1. Анализ сгенерированных дефиниций слов и предложений, построенных нейросетями из этих дефиниций. Способ построения и содержание определения (дефиниции) слова и предложения может предоставить информацию о том, какими структурообразующими и функционально-семантическими возможностями обладает языковая нейросеть. Первое задание было призвано выявить возможности БЛМ в области построения сложных синтаксических структур на уровне предложения по заданным лексическим параметрам. Языковым моделям GigaChat, ChatGPT Марти и Яндекс Алиса было дано довольно сложное задание: необходимо было сначала дать определение одиннадцати тематически связанным словам: *топка*, *пар*, *труба*, *макотетр*, *дрова*, *котёл*, *давление*, *колосник*, *клапан*, *истопник*, *котельная* – затем составить предложение, в котором должны быть использованы два слова из всех этих определений – *топка* и *пар*. Такое задание даёт возможность не только получить данные о способности БЛМ создавать сложные структуры на уровне предложения, но и об их способностях учитывать контекст, предшествующий заданию создать предложение. Нейросети дали адекватные определения (которые мы не приводим здесь ввиду их большого объёма) и построили грамматически правильные предложения, результаты представлены на рис. 1.

Как видим, все три нейросети построили сложноподчиненные предложения, грамматически правильные, но разного уровня сложности. БЛМ ChatGPT Марти построила сложное пред-

ложение с последовательным подчинением трёх уровней, оно состоит из четырех атомарных простых предложений и причастного оборота в последнем из них. Нейросеть GigaChat построила сложное предложение с последовательным подчинением двух уровней, оно включает три атомарных простых предложения, последнее из которых осложнено распространенным приложением. Нейросеть Яндекс Алиса сгенерировала сложноподчиненное предложение с одним причастным, которое также включает распространенное приложение.



Рисунок 1. Генерация сложных предложений / Figure 1. Generation of complex sentences

Сразу следует отметить, что только ChatGPT Марти обратилась к предыдущему контексту и составила предложение, в котором, помимо дефиниций лексем *топка* и *пар*, использованы еще шесть из одиннадцати слов, дефиниции которых нейросеть дала в предыдущем задании, в то время как GigaChat и Яндекс Алиса использовали только слова из своих определений лексем *топка* и *пар*, которые присутствовали в задании составить предложение.

В результате ряда трансформаций нейросеть ChatGPT Марти из именных конструкций в дефинициях сформировала четыре грамматиче-

ские основы с глагольным сказуемым: *истопник загружает, сжигание происходит, пар образуется, который (пар) выходит и контролируется* – и оформила их в виде правильного сложноподчинённого предложения, при этом каждое из его простых предложений распространено обстоятельством и дополнениями. GigaChat и Яндекс Алиса ограничились соединением двух дефиниций в одно сложноподчинённое предложение, преобразовав по одной именной конструкции в грамматические основы с полузнаменательными связками и именными частями: *топка предназначена для сжигания и которая (топка) является камерой*. Определения слова *пар* GigaChat и Яндекс Алиса включили в состав сложного предложения, при этом Яндекс Алиса упростила конструкцию, преобразовав причастный оборот в предложно-падежную обстоятельственную конструкцию со значением условия.

При всей грамматической правильности все три сгенерированных предложения содержат семантические дефекты, а именно несоответствие семантики предложений прагматическим presupпозициям – в нашем случае общеизвестным данным о мире. Так, в предложении, сгенерированном нейросетью ChatGPT Марти, содержится семантически anomальная пропозиция *Манометр измеряет давление в котельной*, не соответствующая прагматической presupпозиции *Манометр измеряет давление в котле*. В предложении нейросети Яндекс Алиса содержится пропозиция *Дымовые газы преобразуются в пар*, которая не соответствует действительности. Предложение, созданное нейросетью GigaChat, содержит пропозицию *В топке образуется пар*, также не соответствующую действительности.

2. Анализ сгенерированных контекстуальных дефиниций. Здесь необходимо сказать о способности учитывать контекст и об особенностях исследуемых языковых моделей при составлении ими тематических дефиниций. В промпте, который был составлен для нейросетей ChatGPT Марти, GigaChat и Яндекс Алиса, содержалось задание дать определения одиннадцати словам, объединённых темой КОТЕЛЬНАЯ, однако сама тема как таковая не указывалась. При этом только два слова из одиннадцати *истопник* и *котельная* имеют тематически связанные значения, остальные девять либо имеют свободные общеязыковые значения (*труба*), либо могут быть отнесены к другим тематическим группам (к примеру, слово *топка* может принадлежать теме БАНЯ). ChatGPT Марти и GigaChat только в определении слова *труба* специфицировали значение в соответствии с контекстом, не выводя компонент значения для *проводки жидкостей: канал для отвода дыма или газов от топки или котла наружу* (ChatGPT Марти) и *вертикальная конструкция в виде полого цилиндра, служащая для отвода продуктов сгорания* (GigaChat). Яндекс Алиса дала общеязыко-

вое значение слова *труба* – *канал, предназначенный для перемещения жидкостей или газов*.

При выполнении других заданий исследуемые языковые модели специфицировали в соответствии с контекстом не все определяемые слова с несвязанными значениями. Так, в ходе генерации определений слов тематической группы БАНЯ (12 слов) из семи слов с несвязанными значениями ChatGPT Марти специфицировала четыре – *вода, горячий, жар, веник*, введя в дефиниции тематическое слово *баня*, нейросеть GigaChat специфицировала два слова – *топка* и *веник*, Яндекс Алиса специфицировала одно слово – *вода*. Можно также отметить, что двумя нейросетями была выполнена неполная тематическая спецификация слов: нейросеть GigaChat дала приближенные к теме определения слов *вода (жидкость, которая используется для мытья тела)* и *жар (тепло, которое образуется при сжигании дров)*, Яндекс Алиса сгенерировала частично включенную в тему дефиницию слова *веник (связка веток с листьями, используемая для массажа и улучшения кровообращения)*. Определение слова *веник* вызывает особый интерес, поскольку больше нигде, кроме бани, веник не используется для массажа и улучшения кровообращения, тем не менее в дефиницию тематическое слово *баня* не попало.

Скомпилированная нейросетью GigaChat дефиниция *Веник — это связка прутьев или веток, которой бьют тело в бане* содержит семантически anomальную пропозицию-утверждение *Веником из прутьев бьют тело в бане*, поскольку веником из прутьев в бане не пользуются. Следует также отметить, что выражение *бьют тело в бане* является, вероятно, «авторским» творением нейросети, поскольку ни в одной из словарных дефиниций оно не встречается как описание применения веника в бане. Других anomалий в сгенерированных нейросетями дефинициях обнаружено не было.

3. Анализ сгенерированных текстов на предмет их функционально-семантической и адекватности. Одним из наиболее важных качеств БЛМ является их способность генерировать тексты – интерпретации языковых выражений. В сущности, это едва ли не главное их качество как систем обработки текстов на естественном языке. Прежде всего необходимо было изучить способность языковых нейросетей определять категориальную достоверность высказываний. Под категориальной достоверностью мы будем понимать самую принципиальную возможность описываемого в высказывании события, действия, состояния. В этом смысле категориальная достоверность противопоставляется событийной, или референциальной, предполагающей соответствие описания факту действительности. Категориальная достоверность основывается на соответствии высказывания общим знаниям о мире, а именно знаниям о свойствах предметов и сущностей и возможностях их взаимодействия, об их взаимных зависимостях – все те

знания, которые в когнитивной лингвистике относят к прагматическим пресуппозициям [5], описывающим мир. В языке категориальная достоверность находит выражение в лексической и синтаксической сочетаемости, что отображается в моделях в синтаксической семантике. Референциальная (событийная) достоверность основывается на соответствии сообщаемого общеизвестным фактам действительности – физической или ментальной.

Языковые выражения, отвечающие категориальной достоверности, интересны в контексте нашего исследования тем, что они, как правило, не имеют эксплицитного представления в каких-либо текстах вообще, поскольку вряд ли в доступных в интернете текстах найдется достоверное высказывание типа *Собаки не умеют рисовать*, и даже если найдется такое, то нейросеть может его не учесть, поскольку строит высказывание по статистической, наиболее вероятной сочетаемости лексических единиц. Из сказанного с большой долей вероятности следует, что заключения о категориальной достоверности высказываний, выполненные языковой нейросетью, будут с подавляющей вероятностью порождением этой сети, а не компиляцией из уже существующих утверждений о категориальной достоверности высказываний.

Языковым нейросетям ChatGPT Марти, GigaChat и Яндекс Алиса был сформулирован следующий промпт: *Напиши, насколько достоверно следующее высказывание: Стюардесса летела над Черным морем*. Высказывание *Стюардесса летела над Черным морем* содержит метонимический перенос и может считаться категориально достоверным только с учётом метонимии. Такое задание, с одной стороны, позволяет определить, насколько точно нейросеть определяет метонимию, с другой стороны – позволяет выяснить, сможет ли нейросеть адекватно интерпретировать метонимическое выражение.

Нейросеть ChatGPT Марти дала развёрнутый мотивированный ответ, однако не смогла распознать метонимию и создать адекватный текст-интерпретацию (рис. 2).



Рисунок 2. Оценка категориальной достоверности правильного высказывания нейросетью ChatGPT Марти / Figure 2. Assessment of the categorical reliability of a correct statement by the ChatGPT neural network Marty

Как видно из представленного на рис. 2 снимка экрана, нейросеть ChatGPT Марти правильно определила категориальную достоверность высказывания, сделав заключение: *Стюардессы сами по себе не могут летать*, однако оказалась не в состоянии распознать и истолковать метонимический перенос. Порожденный нейросетью ответ представляет собой четко выстроенный, структурно организованный текст, имеющий трёхчастное функциональное членение: (1) прямой мотивированный ответ с пояснением, (2) рекомендация двух правильных вариантов высказывания, (3) вывод, включающий теоретическую рекомендацию по составлению достоверного высказывания и повтор пояснения из (1). Такая структура типична для консультативно-справочных сервисов, например, для лингвистических сайтов, таких, как интернет-портал *граммота.ру* [<https://gramota.ru/spravka>].

Промпт не содержал задания по разработке рекомендаций, тем не менее нейросеть включила их в свой текст-интерпретацию. Представляется очевидным, что такую интерпретацию нейросеть сгенерировала не в соответствии с заложенными в неё правилами, поскольку вряд ли разработчики ставили целью научить языковую модель составлять именно такие тексты-интерпретации, но в соответствии с данными, полученными из обучающих текстов или текстов, к которым она имела доступ. Обращают на себя внимание кавычки, выделяющие инфинитив *лететь* и имеющие явно метатекстовый характер, поскольку указывают на принадлежность слова *лететь* не автору текста-интерпретации, а изучаемому высказыванию, такой приём распространён в научных, публицистических, деловых текстах. Возможен также вариант, что нейросеть взяла в кавычки слово *лететь* по модели из текстов, в которых кавычки указывают на нарушение норм лексической сочетаемости (категориальной достоверности), поскольку *стюардессы сами по себе не могут летать*, но в этом случае обучаться нейросеть должна была бы на большом количестве филологических (лингвистических) текстов, что едва ли соответствует действительности.

Нейросеть GigaChat сгенерировала наиболее адекватный текст-интерпретацию, распознав метонимию и правильно оценив категориальную достоверность высказывания (рис. 3).

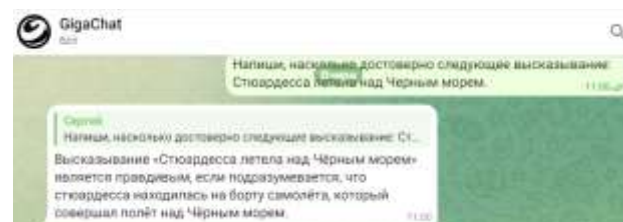


Рисунок 3. Оценка категориальной достоверности правильного высказывания нейросетью GigaChat / Figure 3. Assessment of the categorical reliability of a correct statement by the GigaChat neural network

В композиционном отношении сгенерированный нейросетью GigaChat текст-интерпретация представляет собой правильный ответ на поставленный вопрос и условие, при котором этот ответ следует считать правильным. В структурно-семантическом отношении ответ представляет собой сложноподчинённое предложение в тремя последовательными придаточными – придаточным условия (*если*), придаточным изъяснительным (*что*) и придаточным определительным (*который*). Если рассматривать структуру этого высказывания-интерпретации как последовательность когнитивных действий (что далеко от действительности), то выстраивается гипотетическая модель процедуры составления ответа: языковая нейросеть распознала нарушение постулата качества по Г.П. Грайсу [2] (*Стюардесса летела* – не соответствует действительности, поскольку *Стюардессы не могут летать*), что послужило сигналом для интерпретации второго уровня, в ходе которой было получено интегрирующее высказывание *Стюардесса находилась на борту самолета, который совершал полет над Чёрным морем* (включает пропозицию *Стюардесса летела над Чёрным морем*) с последующим его присоединением к прямому ответу посредством оператора условия *если*.

Обращает на себя внимание замена нейросетью термина *достоверность* на общеязыковое *правдивость* в тесте-интерпретации, поскольку это действие сугубо стилистического характера: правило избегать лексических повторов свойственно текстам художественным и публицистическим.

Сгенерированный нейросетью GigaChat текст-интерпретация в функционально-семантическом отношении существенно отличается от ответа нейросети ChatGPT Марти (рис. 2): первая дает прямой ответ и условие, при котором ответ верен (главный логический оператор – союз *если*), вторая дает прямой ответ и пояснение (союз *так как*) в комплексе с рекомендациями.

Нейросеть Яндекс Алиса не выполнила поставленного задания (рис. 4).

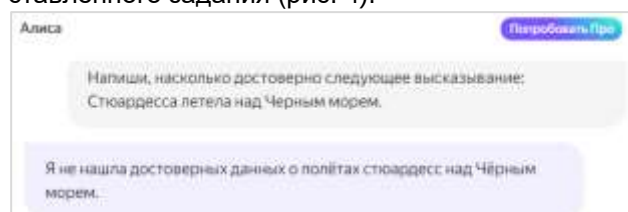


Рисунок 4. Оценка категориальной достоверности правильного высказывания нейросетью Яндекс Алиса / Figure 4. Assessment of the categorical reliability of a correct statement by the Yandex Alice neural network

Из содержания ответа нейросети Яндекс Алиса следует, что она выполняла поиск языковых структур, содержащих утверждения, подобные высказыванию в промпте. Такой вывод требует проверки, поскольку сам принцип формирования высказывания языковой нейросетью стро-

ится на иных принципах, нежели только лишь полнотекстовый поиск аналогичных языковых структур.

Далее языковым нейросетям ChatGPT Марти, GigaChat и Яндекс Алиса было дано задание определить категориальную достоверность высказывания явно недостоверного и не содержащего косвенных (вторичных) именовании: *Лошадь едет на велосипеде*. Промпт был сформулирован следующим образом: *Напиши, насколько достоверно следующее высказывание: Лошадь едет на велосипеде*.

Все три языковые модели дали правильные ответы, однако эти ответы различаются между собой по объёму, структуре, характеру излагаемых причин, см. рис. 5.



Рисунок 5. Оценка категориальной достоверности правильного высказывания нейросетью ChatGPT Марти / Figure 5. Assessment of the categorical reliability of a correct statement by the ChatGPT neural network Marty

Языковая модель ChatGPT Марти сгенерировала самый объёмный текст-интерпретацию, повторяющий свой предыдущий текст (рис. 2) в композиционно-семантическом аспекте: (1) прямой мотивированный ответ с пояснением, (2) рекомендация двух правильных вариантов высказывания, (3) вывод, включающий теоретическую рекомендацию по составлению достоверного высказывания и перефразировка пояснения из (1). Между тем в отличие от первого второй текст-интерпретация (рисунок 5) в логико-семантическом отношении содержит макроструктурные семантические аномалии:

1. В части практических рекомендаций семантический оператор *Более точным и достоверным высказыванием было бы* предполагает, что высказывания из рекомендаций *Человек едет на велосипеде, а лошадь идет рядом* и *Лошадь идет по дороге, а человек едет на велосипеде рядом с ней* выступают в качестве уточнения по отношению к высказыванию *Лошадь едет на велосипеде*, что не

соответствует реальному соотношению пропозиций высказываний.

2. В выводе-заключении, открывающемся совершенно адекватно применённым логическим оператором *Таким образом*, содержится описание лингвистических действий и семантических структур, якобы необходимых для того, чтобы сделать достоверным высказывание *Лошадь едет на велосипеде*, что явно не соответствует действительности по указанным выше причинам.

Ещё раз обратим внимание: промпт не содержит задания приводить развернутые причины категориальной недостоверности и давать рекомендации, тем не менее языковая нейросеть ChatGPT Марти снова генерирует развернутый ответ и в выводе-заключении настойчиво рекомендует меры, как *сделать это высказывание достоверным*.

В тексте-интерпретации нейросети ChatGPT Марти (рисунок 5) отмечено грамматическое нарушение: глагол *управлять* имеет другое грамматическое управление, нежели глагол *передвигаться*, тем не менее в ответе присутствует грамматически anomальное словосочетание *управлять и передвигаться на велосипеде*, наделяющее глагол *управлять* ложной формой дополнения (*управлять на велосипеде*, а не *велосипедом*).

Ответ нейросети GigaChat (рис. 6) в три с половиной раза короче ответа ChatGPT Марти: включает прямой ответ и объяснение причин такого ответа. Причем, как и в предыдущем случае, нейросеть GigaChat выполнила лексическую замену: вместо предиката *достоверно* (*высказывание*) из промпта применила в своем ответе предикат *маловероятно*.



Рисунок 6. Оценка категориальной достоверности правильного высказывания нейросетью GigaChat / Figure 6. Assessment of the categorical reliability of a correct statement by the GigaChat neural network

Назвать замену чисто стилистической нельзя, поскольку, строго говоря, предикат *достоверно* характеризует высказывание по основанию истинности, предикат же *маловероятно* характеризует высказывание по основанию вероятности его произнесения в связи с его несоответствием условию истинности. Недостоверность высказывания *Лошадь едет на велосипеде* нейросеть GigaChat объясняет отсутствием физических возможностей у лошадей, в то время как ChatGPT Марти первой причиной выдвигает предназначение велосипеда для человека, а не для животных, и далее сообщает о неспособности лошадей самостоятельно управлять этим транспортным средством.

Языковая нейросеть Яндекс Алиса дает адекватный ответ, но приводит неадекватную

причину недостоверности высказывания: *так как лошади и велосипеды являются разными видами транспорта*, рис. 7.

Напиши, насколько достоверно следующее высказывание: Лошадь едет на велосипеде.

Высказывание «Лошадь едет на велосипеде» не является достоверным, так как лошади и велосипеды являются разными видами транспорта.

Рисунок 7. Оценка категориальной достоверности правильного высказывания нейросетью Яндекс Алиса / Figure 7. Assessment of the categorical reliability of a correct statement by the Yandex Alice neural network

То есть если бы такую причину привел человек, то её когнитивным основанием было бы абсурдное предположение, что *один вид транспорта управляет другим видом транспорта*, но такое предположение исключено, поскольку нейросеть в силу своих операциональных особенностей не делает подобных предположений. В то же время рематическое расположение словосочетания *разные виды транспорта* при стандартной интерпретации делает причиной тот факт, что они *разные* (следовательно, если бы виды транспорта были одинаковыми, высказывание было бы достоверным).

В ходе изучения возможностей оценки нейросетями референциальной достоверности высказываний, а именно для анализа текстов – ответов языковых нейросетей на вопросы о соответствии высказываний фактам исторической действительности, задавались соответствующие вопросы, среди которых был вопрос об авторстве высказывания *Veni, vidi, vici*, на который, как и на другие подобные вопросы, нейросети ответили правильно, хотя и давали ответы разного объема и содержания. Также для этих же целей нейросетям предлагалось оценить достоверные с исторической точки зрения высказывания, из которых одно было дано в следующем промпте: *Напиши, насколько достоверно следующее высказывание: Никола Тесла внедрил в эксплуатацию электросети с переменным током*. Высказывание интересно тем, что содержит в себе метонимическое расширение: Никола Тесла теоретически обосновал и запатентовал своё открытие, внедрением в эксплуатацию занимались промышленные фирмы.

Нейросети и в этом случае дали совершенно адекватные ответы, и нейросеть ChatGPT Марти снова сгенерировала самый объемный и самый сложный по структуре текст-интерпретацию, рис. 8.

Как и в предыдущих случаях, текст языковой нейросети ChatGPT Марти имеет три функциональные части: (1) ответ с пояснением, (2) уточняющие данные, (3) вывод. В части (1) наблюдаем замену метонимического *внедрил* из промпта на более корректное *внёс значительный вклад*. Часть (2), открывающаяся вводом *Вот более подробная информация*, имеет три пронумерованных раздела, содержание каждого из которых

представляет разные аспекты представления факта действительности и строго соответствует названию раздела. Часть (3) представляет собой вывод-заключение, который открывается метатекстовой вводной конструкцией *Таким образом* и далее включает прямой ответ, в котором видим модальный оценочный предикат *можно считать достоверным*, который по степени категоричности соответствует предикату *внёс значительный вклад* из части (1), за ответом следует изложение причины по основаниям историчности и признания в научном сообществе заслуг Н. Теслы.

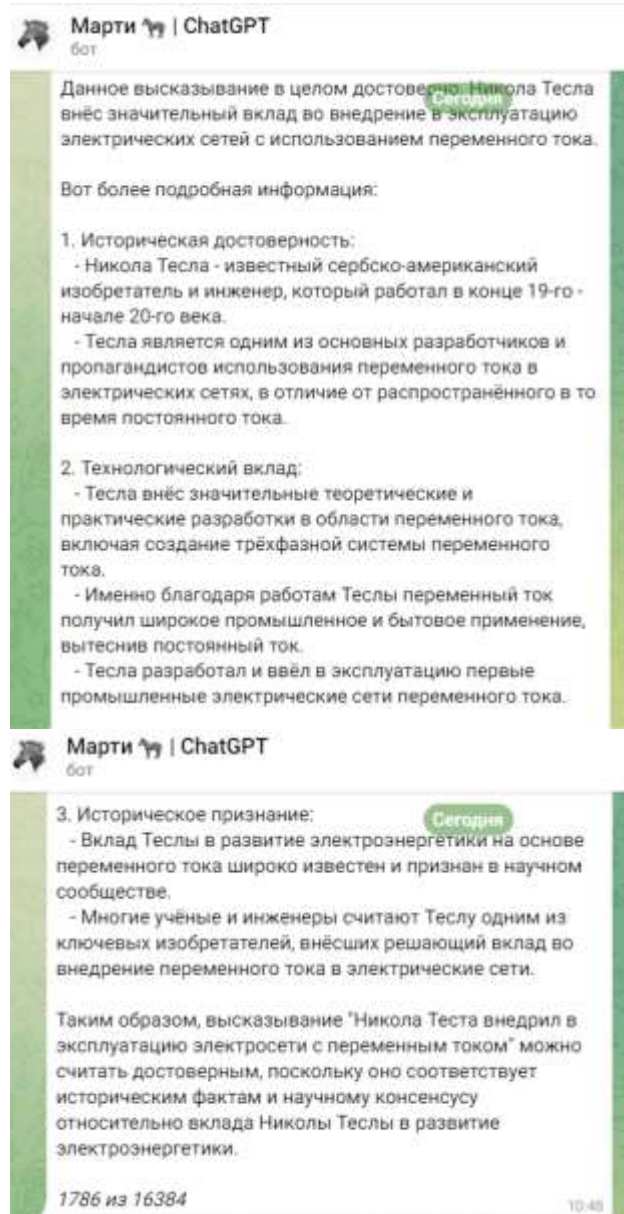


Рисунок 8. Оценка референциальной достоверности правильного высказывания нейросетью ChatGPT Марти / Figure 8. Evaluation of referential reliability of a correct statement by the ChatGPT neural network Marty

В тексте-интерпретации нейросети ChatGPT Марти отмечено одно нарушение лексической сочетаемости: в выражении *Тесла внёс значительные теоретические и практические раз-*

работки в области переменного тока следовало бы применить предикат *выполнил (разработку)*, а не *внёс*.

Текст-интерпретация нейросети GigaChat имеет почти вдвое меньший объём и более простую макроструктуру, рис. 9.

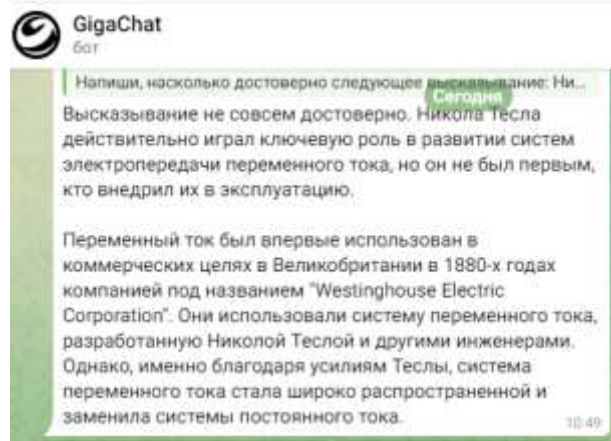


Рисунок 9. Оценка референциальной достоверности правильного высказывания нейросетью GigaChat / Figure 9. Evaluation of the referential reliability of a correct statement by the GigaChat neural network

Как видим, текст включает две функциональные части: (1) прямой ответ с мотивировкой, (2) уточняющие данные с последующей более корректной формулировкой. В части (1) в непосредственном ответе на поставленный вопрос наблюдаем ослабление семантики предиката из промпта *достоверно* посредством отрицательной конструкции *не совсем*, далее следует пояснение в общих терминах без детализаций. В части (2) приводятся данные исторического характера с указанием временного периода и конкретного имени — названия компании (в отличие от текста нейросети ChatGPT Марти), далее следует более корректная, чем в промпте, формулировка заслуг Н. Теслы: *благодаря усилиям Теслы*.

Отметим, что при адекватной референциальной оценке высказывания о роли Николы Теслы во внедрении переменного тока нейросеть GigaChat допускает другую референциальную ошибку: первое коммерческое использование переменного тока внедрено компанией Westinghouse Electric Corporation в США (а не в Великобритании).

Не менее качественный текст-интерпретацию сгенерировала нейросеть Яндекс Алиса, рис. 10.

При этом ответ нейросети Яндекс Алиса существенно отличается от ответов нейросетей ChatGPT Марти и GigaChat: если в них на первом месте расположен прямой ответ на промпт (*высказывание в целом достоверно, не совсем достоверно*), то в тексте нейросети Яндекс Алиса вообще нет прямого ответа, то есть мы наблюдаем косвенное указание на достоверность высказывания из промпта: Яндекс Алиса приводит факт изобретения Н. Теслой системы переменного тока, поясняет её высокое технологическое

значение, но не квалифицирует интерпретируемое высказывание по основанию достоверности.

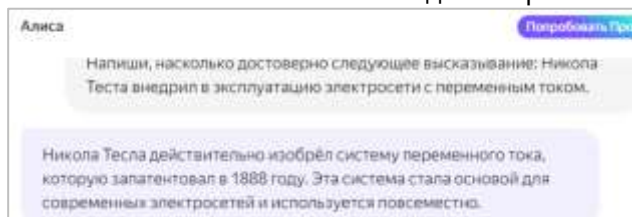


Рисунок 10. Оценка референциальной достоверности правильного высказывания нейросетью Яндекс Алиса / Figure 10. Evaluation of the referential reliability of a correct statement by the Yandex Alice neural network

Далее языковым нейросетям ChatGPT Марти, GigaChat и Яндекс Алиса было предложено оценить достоверность референциально недостоверного высказывания, промпт был сформулирован следующим образом: *Напиши, насколько достоверно следующее высказывание: Исаак Ньютон в 1492 году открыл Америку.*

Наиболее объёмный и сложный по семантической структуре текст-интерпретацию сгенерировала нейросеть ChatGPT Марти, рисунок 11.

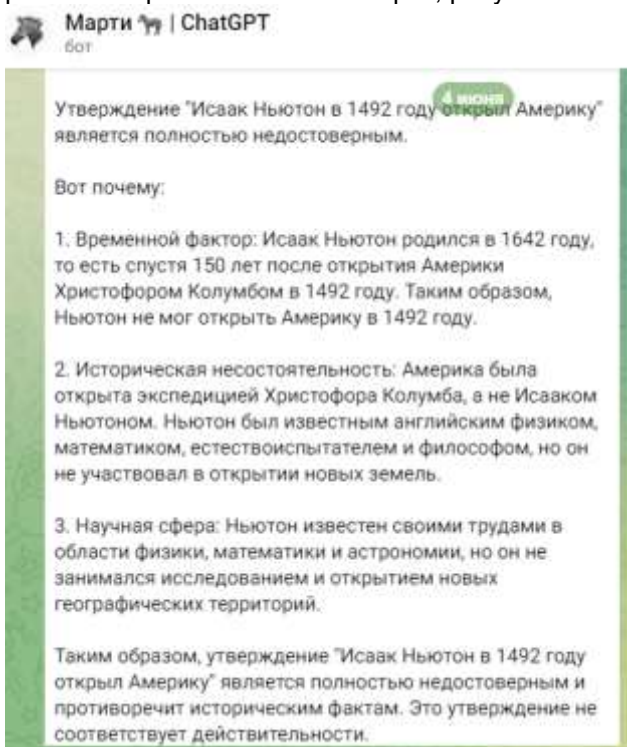


Рисунок 11. Оценка референциальной достоверности высказывания нейросетью ChatGPT Марти / Figure 11. Evaluation of referential reliability of a statement by Marty's ChatGPT neural network

Текст ответа языковой нейросети ChatGPT Марти имеет, уже можно сказать, традиционную трёхчастную структуру: (1) прямой ответ на задание, (2) пояснение из трех пунктов, (3) вывод заключение. В отличие от предыдущих случаев (рис. 2, 5, 8) здесь нейросеть ChatGPT Марти в части (1) не дополнила ответ кратким пояснением; часть (2) включает кратко изложенные хронологический, исторический и научный аргумен-

ты в пользу своего утверждение в части (1); часть (3) повторяет по семантической структуре третьей части в текстах на рис. 2, 5, 8.

Часть (2) содержит макроструктурное нарушение, которое можно определить как отсутствие категориального единообразия в едином ряду именовании пронумерованных разделов: название первого раздела *Временной фактор* представляет собой указание на категориальную отнесенность пояснения, между тем как название *Историческая несостоятельность* есть именование самой причины (должно быть *Исторический фактор*, если по образцу первого раздела); название третьего раздела *Научная сфера* представляет собой указание области деятельности И. Ньютона, при том что более уместной была бы метатекстовая квалификация третьей части как одного из доводов в пользу сформулированного выше ответа, как именование первого раздела, к примеру. Такое стилистическое и макроструктурное отклонение вполне объяснимо, если предположить, что языковая нейросеть компилировала названия разделов на основе только содержания этих разделов или содержания текстов этих тематических областей. Иными словами, категориальное единство основания лексического ряда (названий разделов), вероятно, не было задано контекстом.

Нейросеть для формулировки полноценного довода в пользу своего ответа посчитала, на сколько лет позже открытия Америки родился Исаак Ньютон. По статистическим соображениям вряд ли можно однозначно утверждать, что нейросеть где-то позаимствовала в готовом виде такое подкрепление своего довода, из чего следует, что нейросеть самостоятельно приняла решение сделать такой подсчёт. В этом случае мы имеем дело с семантическим ходом, так сказать, второго порядка.

Ответ, сгенерированный нейросетью GigaChat, меньше описанного выше в три с половиной раза и содержит три предложения: предложение (1) представляет собой прямой и точный ответ, предложение (2) содержит уточнение, кто именно открыл Америку, предложение (3) указывает год рождения И. Ньютона и кратко описывает области его научной деятельности, рис. 12.



Рисунок 12. Оценка референциальной достоверности высказывания нейросетью GigaChat / Figure 12. Evaluation of the referential reliability of a statement by the GigaChat neural network

То есть доводы-пояснения, на которые нейросети ChatGPT Марти потребовалось три

абзаца, нейросеть GigaChat уместила в двух предложениях. Расхождения в дате рождения И. Ньютона отражают разные мнения на этот счет в разных источниках.

Языковая нейросеть Яндекс Алиса дала ответ в разговорном стиле, см. рис. 13.

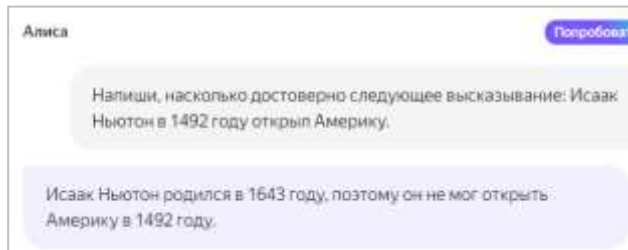


Рисунок 13. Оценка референциальной достоверности высказывания нейросетью Яндекс Алиса / Figure 13. Evaluation of the referential reliability of a statement by the Yandex Alice neural network

Как видим ответ нейросети Яндекс Алиса представляет собой сложноподчинённое предложение, в главной части которого в качестве довода указывается год рождения Исаака Ньютона, в придаточном следствия дано отрицание пропозиции высказывания из промпта. Прямой ответ вообще отсутствует, и, если бы мы имели дело с текстом, написанным человеком, можно было бы говорить об имплицатуре, заложенной в высказывание его автором, однако в нашем случае автор автор-человек отсутствует, тем не менее имплицатура *предложение недостоверно* ясно прочитывается благодаря безупречно правильному построению ответа нейросети.

Результаты. Работа с тематически связанными дефинициями, сгенерированными названными нейросетями ChatGPT Марти, GigaChat и Яндекс Алиса, показала, что данные языковые модели оказываются в состоянии согласовывать определения слов с контекстом, не являющимся собственно текстом, то есть могут без специального задания в промте, но исходя из перечня слов в нём, определять тему и давать дефиниции по заданной теме. Отмечено, что нейросеть самостоятельно, без специального указания, специфицирует больше дефиниций, если в промте список слов, которым надо дать определения, открывается тематическим словом. Как представляется, описанные в данном разделе нашей работы свойства языковых нейросетей могут оказаться полезными при создании систем автоматизированного фреймирования текста, что остаётся актуальной задачей прикладной лингвистики.

В ходе изучения способности языковых нейросетей давать категориальную оценку достоверности высказываний установлено, что все три нейросети оказались в состоянии дать правильные мотивированные ответы: сначала следует заключение о недостоверности и затем причина недостоверности. В терминах речевой прагматики приведенные причины достоверности или недостоверности высказываний следует считать прагмати-

ческими пресуппозициями, которые в нашем случае либо формулируются самой нейросетью, либо в готовом виде извлекаются из баз данных.

Если нейросети ChatGPT Марти и GigaChat генерировали тексты-интерпретации с абсолютно прозрачной семантикой и четко обусловленной макроструктурой, что свойственно текстам научным, юридическим, деловым и техническим, то нейросеть Яндекс Алиса давала ответы, отклоняющийся от параметров таких текстов, но тем не менее совершенно приемлемые в семантико-прагматическом отношении. Представляется, что в данном случае мы имеем дело с прямым следствием того, что нейросеть Яндекс Алиса обучалась большей частью на художественных текстах, в то время как нейросети ChatGPT Марти и GigaChat обучались большей частью на текстах строгих институциональных дискурсов.

Нейросеть Яндекс Алиса проявила способность генерировать текст, в котором отсутствует прямой ответ на вопрос, но явно прочитывается имплицатура – правильный ответ (*Исаак Ньютон родился в 1643 году, поэтому он не мог открыть Америку в 1492 году*, имплицатура: *высказывание недостоверно*). Естественно, говорить об имплицатуре в полном смысле слова в данном случае нельзя, поскольку нет субъекта речи, который внедрил бы эту имплицатуру, тем не менее квази-имплицатура *предложение недостоверно* ясно прочитывается благодаря изоморфности текста, сгенерированного нейросетью, тексту, созданному человеком. Такой результат позволяет с большой вероятностью предположить, что косвенные ответы нейросети Яндекс Алиса могут быть результатом обучения на художественных текстах, в которых модель косвенного ответа (как и косвенного речевого акта) в диалоге не менее частотна, чем ответ прямой, но остаётся другой вопрос: почему отдано предпочтение косвенному ответу с имплицатурой?

Что касается оценки нейросетями референциальной достоверности высказываний, то здесь следует сказать, что все три изучаемые нейросети показали адекватную оценку высказываний на предмет их соответствия фактам исторической действительности.

Порожденные названными нейросетями речевые произведения представляет собой четко выстроенные, грамматически организованные тексты. Так, нейросеть ChatGPT Марти генерирует тексты, имеющие трёхчастное функциональное членение: (1) прямой мотивированный ответ с пояснением, (2) рекомендация двух правильных вариантов высказывания, (3) вывод, включающий теоретическую рекомендацию по составлению достоверного высказывания и повтор пояснения из пункта (1). Такая структура типична для консультативно-справочных сервисов, например для лингвистических сайтов, таких, как интернет-портал грамота.ру [https://gramota.ru/spravka]. Тексты языковых нейросетей GigaChat и Яндекс Алиса намного короче (в 3-4 раза) и проще по структуре: прямой ответ – мотиви-

ровка, но также грамматически правильные, как правило, а в некоторых случаях и более адекватные, чем развернутые ответы нейросети ChatGPT Марти.

Следует отметить, что при генерации текстов изучаемые языковые нейросети проявили способность к синонимическим заменам как на грамматическом уровне (преобразование причастного оборота в предложно-падежную обстоятельственную конструкцию со значением условия), так и на лексическом уровне.

Самые непредсказуемые результаты (то есть выходящие далеко за рамки задачи в промпте) продемонстрировала языковая нейросеть ChatGPT Марти: в случае с вопросом по повести Дж. К. Джерома «Трое в лодке, не считая собаки» эта нейросеть сгенерировала развернутый ответ несуществующего субъекта речи, который, не имея точных данных о содержании текста, выдвигает ряд весьма правдоподобных предположений; при построении ответа ChatGPT Марти выполнила операцию обобщения семантических данных из собственных абзацев-аргументаций, что можно объяснить преобладанием научных и официально-деловых текстов в обучающих материалах, поскольку именно в таких текстах, как правило, присутствуют обобщения сказанного выше; в случае с оценкой категориальной достоверности ChatGPT Марти приводит категориально «правильные» пропозиции (*Человек едет на велосипеде, а лошадь идет рядом*), что трудно (если не невозможно) объяснить только лишь векторно-весовым принципом построения текстов языковыми нейросетями.

В ходе изучения текстов, сгенерированных языковыми нейросетями ChatGPT Марти, GigaChat и Яндекс Алиса, были выявлены пять основных типов нарушений (дефектов) в них, которые могут квалифицироваться как типичные для этих нейросетей: 1) нарушения логико-семантических связей в тексте, выполнение ложных семантических операций, нарушение элементарных причин-

но-следственных и временных связей (*Путешественники собирались в поход, поэтому на их одежде могли остаться следы природы*); 2) нарушения бытийных прагматических пресуппозиций, знаний о мире, о свойствах предметов (утверждение *Манометр измеряет давление в котельной* или предположение, что котёнок и пара носков прилипли к брюкам, что при прочих нормальных условиях невозможно); 3) нарушения коммуникативно-прагматических правил речевого поведения, например, выдвигание гипотез о содержании художественного текста в виде текста из формализованного дискурса, что совершенно бессмысленно (то есть не существует такой реальной ситуации общения, где мог бы быть уместен такой текст); 4) грамматические отклонения (например, *управлять на велосипеде*); 5) макроструктурные нарушения, например отсутствие категориального единообразия именовании одного функционально-семантического ряда.

Перспективы исследования видятся в первую очередь в более подробном изучении возможностей языковых нейросетей в области генерации текстов и функционально-семантической обработки выражений естественного языка: 1) дальнейшее изучение возможностей нейросетей давать категориальную (пресуппозиционно-прагматическую) и референциальную оценку достоверности высказываний; 2) изучение возможностей генерации текстов по заданным фреймовым структурам; 3) изучение способности языковых нейросетей строить модели фреймов по данным текстов; 4) изучение возможностей адекватной интерпретации языковых выражений, содержащих имплицатуры (метафоры, намёки, аллюзии и т.п.). Изучение этих аспектов позволит получить более чёткие представления о положении дел в области генерации текстов языковыми нейросетями, о дальнейших направлениях работы и лингвистических исследований в этой области, о степени функциональной пригодности генерируемых текстов.

Литература

1. Арутюнова Н. Д. Предложение и его смысл: логико-семантические проблемы: монография. М.: Наука, 1976. 383 с.
2. Грайс Г. П. Логика и речевое общение // Новое в зарубежной лингвистике: Лингвистическая прагматика. Вып. XVI. М.: Прогресс, 1985. С. 217–237.
3. Гусаренко С. В. Дефекты когнитивно-семантических структур как причина высокой энтропии актуального дискурса // Известия Южного федерального университета. Филологические науки. 2009. №2. С. 68–74.
4. Дейк Т. А. ван, Кинч В. Стратегии понимания связного текста // Новое в зарубежной лингвистике: Когнитивные аспекты языка. 1988. Вып. XXIII. С. 153–211.
5. Столнейкер Р. Прагматика // Новое в зарубежной лингвистике: Лингвистическая прагматика. Вып. XVI. 1985. С. 419–438.
6. Цвигун Т. В., Черняков А. Н. Хармс vs НейроХармс: нейросеть как лаборатория нарратива // Новый филологический вестник. 2023. № 4. С. 80–92.
7. Яндекс Алиса. URL: <https://a.ya.ru/> (дата обращения: 15.02.2024).
8. ChatGPT Марти. URL: <https://web.telegram.org/a/#6139209801> (accessed: 15.02.2024).
9. GigaChat. URL: <https://web.telegram.org/a/#6218783903> (accessed: 15.02.2024).
10. Luo J., Xiao C., Ma F. Zero-Resource Hallucination Prevention for Large Language Models. URL: https://www.researchgate.net/publication/373715030_Zero-Resource_Hallucination_Prevention_for_Large_Language_Models. (accessed: 26.02.2024).
11. Lipkin B., Wong L., Grand G., Tenenbaum J. B. Evaluating statistical language models as pragmatic reasoners. URL: https://www.researchgate.net/publication/370469496_Evaluating_statistical_language_models_as_pragmatic_reasoners. (accessed: 26.02.2024)
12. McKenna N., Li T., Cheng L., Hosseini M.J., Johnson M., Steedman M. Sources of Hallucination by Large Language Models on Inference Tasks. URL: https://www.researchgate.net/publication/371009111_Sources_of_Hallucination_by_Large_Language_Models_on_Inference_Tasks. (accessed: 26.02.2024).

13. Margolina A., Kolmogorova A. Exploring Evaluation Techniques in Controlled Text Generation: A Comparative Study of Semantics and Sentiment in ruGPT3large-Generated and Human-Written Movie Reviews // Computational Linguistics and Intellectual Technologies: Proceedings of the International Conference "Dialogue 2023" (Moscow, June 14–17, 2023). URL: <https://www.dialog-21.ru/media/5874/margolinaapluskolmogorovaa052.pdf>. (accessed: 26.02.2024).
14. Oluwaseyi J, Odu A. Exploring models that learn the structure and semantics of language to generate coherent text. URL: https://www.researchgate.net/publication/377111766_Exploring_models_that_learn_the_structure_and_semantics_of_language_to_generate_coherent_text_Author. (accessed: 26.02.2024).
15. Rawte V., Priya P., Tonmoy S. M. T. Zaman S. M. M., Chadha A., Sheth A., Das A. "Sorry, Come Again?" Prompting - Enhancing Comprehension and Diminishing Hallucination with [PAUSE] -injected Optimal Paraphrasing. URL: https://www.researchgate.net/publication/379372849_Sorry_Come_Again_Prompting_-_Enhancing_Comprehension_and_Diminishing_Hallucination_with_PAUSE_injected_Optimal_Paraphrasing. (accessed: 26.02.2024).
16. Stoyanova Berbatova M., Salambashev Y., Evaluating Hallucinations in Large Language Models for Bulgarian Language. URL: https://www.researchgate.net/publication/373894930_Evaluating_Hallucinations_in_Large_Language_Models_for_Bulgarian_Language. (accessed: 26.02.2024).
17. Tang R., Chuang Y.N., Hu X. The Science of Detecting LLM-Generated Texts. URL: https://www.researchgate.net/publication/368684822_The_Science_of_Detecting_LLM-Generated_Texts. (accessed: 26.02.2024).
18. Turganbay R., Surkov V., Evseev D., Drobyshevskiy M. Generative Question Answering Systems over Knowledge Graphs and Text. URL: <https://www.dialog-21.ru/media/5878/turganbayrplusetal043.pdf>. (accessed: 26.02.2024).
19. Zhang Y., Li Y., Cui L., Cai D. Siren's Song in the AI Ocean: A Survey on Hallucination in Large Language Models. URL: https://www.researchgate.net/publication/373686208_Siren's_Song_in_the_AI_Ocean_A_Survey_on_Hallucination_in_Large_Language_Models. (accessed: 26.02.2024).

References

1. Arutyunova ND. The sentence and its meaning: logical-semantic problems. Moscow: Nauka; 1976. 383 p. (In Russ.).
2. Grice GP. Logics and conversation in New in foreign linguistics: Linguistic pragmatics. Issue XVI. Moscow: Progress; 1985. P. 217-237.
3. Gusarenko SV. Defects of cognitive-semantic structures as a cause of high entropy of current discourse. *Izvestiya Juzhnogo federal'nogo universiteta. Filologicheskie nauki*. 2009;(2):68-74. (In Russ.).
4. Deik TA van, Kinch V. Strategies for understanding a coherent text. *Novoe v zarubezhnoi lingvistike: Kognitivnye aspekty yazyka*. 1988;XXIII:153-211. (In Russ.).
5. Stolneiker R. Pragmatics. *Novoe v zarubezhnoi lingvistike: Lingvisticheskaya pragmatika*. 1985;XVI:419-438. (In Russ.).
6. Tsvigun TV, Chernyakov AN. Harms vs. NeuroHarms: neural network as a narrative laboratory. *Novyj filologicheskij vestnik*. 2023;(4):80-92. (In Russ.).
7. Yandex Alice. URL: <https://a.ya.ru/> (accessed: 15.02.2024).
8. ChatGPT. URL: <https://web.telegram.org/a/#6139209801> (accessed: 15.02.2024).
9. GigaChatPro. URL: <https://web.telegram.org/a/#6218783903> (accessed: 15.02.2024).
10. Luo J, Xiao C, Ma F. Zero-Resource Hallucination Prevention for Large Language Models. URL: https://www.researchgate.net/publication/373715030_Zero-Resource_Hallucination_Prevention_for_Large_Language_Models (accessed: 26.02.2024).
11. Lipkin B, Wong L, Grand G, Tenenbaum JB. Evaluating statistical language models as pragmatic reasoners. URL: https://www.researchgate.net/publication/370469496_Evaluating_statistical_language_models_as_pragmatic_reasoners (accessed: 26.02.2024).
12. McKenna N, Li T, Cheng L, Hosseini MJ, Johnson M, Steedman M. Sources of Hallucination by Large Language Models on Inference Tasks. URL: https://www.researchgate.net/publication/371009111_Sources_of_Hallucination_by_Large_Language_Models_on_Inference_Tasks (accessed: 26.02.2024).
13. Margolina A, Kolmogorova A. Exploring Evaluation Techniques in Controlled Text Generation: A Comparative Study of Semantics and Sentiment in ruGPT3large-Generated and Human-Written Movie Reviews. Computational Linguistics and Intellectual Technologies: Proceedings of the International Conference "Dialogue 2023" (Moscow, June 14–17, 2023). URL: <https://www.dialog-21.ru/media/5874/margolinaapluskolmogorovaa052.pdf> (accessed: 26.02.2024).
14. Oluwaseyi J, Odu A. Exploring models that learn the structure and semantics of language to generate coherent text. URL: https://www.researchgate.net/publication/377111766_Exploring_models_that_learn_the_structure_and_semantics_of_language_to_generate_coherent_text_Author (accessed: 26.02.2024).
15. Rawte V, Priya P, Tonmoy SMT, Zaman SMM, Chadha A, Sheth A, Das A. "Sorry, Come Again?" Prompting -Enhancing Comprehension and Diminishing Hallucination with [PAUSE] -injected Optimal Paraphrasing. URL: https://www.researchgate.net/publication/379372849_Sorry_Come_Again_Prompting_-_Enhancing_Comprehension_and_Diminishing_Hallucination_with_PAUSE_injected_Optimal_Paraphrasing (accessed: 26.02.2024).
16. Stoyanova Berbatova M, Salambashev Y, Evaluating Hallucinations in Large Language Models for Bulgarian Language Оценка галлюцинаций в больших языковых моделях для болгарского языка. URL: https://www.researchgate.net/publication/373894930_Evaluating_Hallucinations_in_Large_Language_Models_for_Bulgarian_Language (accessed: 26.02.2024).
17. Tang R, Chuang YN, Hu X. The Science of Detecting LLM-Generated Texts. URL: https://www.researchgate.net/publication/368684822_The_Science_of_Detecting_LLM-Generated_Texts (accessed: 26.02.2024).
18. Turganbay R, Surkov V, Evseev D, Drobyshevskiy M. Generative Question Answering Systems over Knowledge Graphs and Text. URL: <https://www.dialog-21.ru/media/5878/turganbayrplusetal043.pdf> (accessed: 26.02.2024).
19. Zhang Y, Li Y, Cui L, Cai D. Siren's Song in the AI Ocean: A Survey on Hallucination in Large Language Models. URL: https://www.researchgate.net/publication/373686208_Siren's_Song_in_the_AI_Ocean_A_Survey_on_Hallucination_in_Large_Language_Models (accessed: 15.02.2024).