

СОВРЕМЕННЫЕ ИНФОРМАЦИОННЫЕ ТЕХНОЛОГИИ. ЦИФРОВАЯ ТРАНСФОРМАЦИЯ ДЕЯТЕЛЬНОСТИ БИБЛИОТЕК

УДК 004.032.26:02

<https://doi.org/10.33186/1027-3689-2025-7-142-163>

Модель коммуникации с искусственным интеллектом ДРУГ как методологический подход к составлению и оценке промптов

А. В. Ковалевский

*Белорусский национальный технический университет,
Минск, Республика Беларусь,
kovalevskyalex@yandex.ru*

Аннотация. Цифровая трансформация библиотечной сферы и постепенное внедрение технологий искусственного интеллекта в практику работы библиотечных специалистов актуализировали проблему эффективного взаимодействия с нейросетями и формирования качественных промптов. В статье представлена авторская модель коммуникации с искусственным интеллектом, основанная на принципах доброжелательности, рациональности, уточнения и гносеологичности (ДРУГ) и разработанная как методологический подход к промпт-инжинирингу в библиотечной деятельности. Цель исследования – описание данной модели и демонстрация её практического применения и эффективности. Методология исследования включала разработку модели, создание серии промптов различной «силы» на основе её принципов и тестирование с использованием шести моделей нейросетей для решения задачи анализа библиографических данных. Результаты тестирования показали прямую зависимость качества ответов нейросетей от уровня проработки промпта в соответствии с моделью ДРУГ. Практическая значимость исследования заключается в предложении методологической основы для повышения эффективности и качества взаимодействия библиотечных специалистов с искусственным интеллектом, а также в возможности использования модели для разработки практических рекомендаций и образовательных программ по промпт-инжинирингу в библиотечной сфере.

Ключевые слова: искусственный интеллект, библиотечное дело, промпт-инжиниринг, промпт, нейросети, коммуникация, информационно-коммуника-

ционные технологии, компетенции, аналитическая работа, анализ данных, моделирование

Для цитирования: Ковалевский А. В. Модель коммуникации с искусственным интеллектом ДРУГ как методологический подход к составлению и оценке промптов // Научные и технические библиотеки. 2025. № 7. С. 142–163. <https://doi.org/10.33186/1027-3689-2025-7-142-163>

MODERN INFORMATION TECHNOLOGIES. DIGITAL TRANSFORMATION OF LIBRARIES

UDC 004.032.26:02

<https://doi.org/10.33186/1027-3689-2025-7-142-163>

The model of communication with artificial intelligence as a methodological approach to prompt creation and evaluation

Aleksey V. Kovalevsky

*Belarusian National Technical University,
Minsk, Republic of Belarus,
kovalevskyalex@yandex.ru*

Abstract. In the context of digital transformation of the library sphere and gradual implementation of artificial intelligence technologies into library practice, the challenge of effective interaction with neural networks and high-quality promptsmithing is increasingly relevant. The author proposes his model of communication with artificial intelligence, based on the principles of Benevolence, Rationality, Refinement, and Epistemology, and demonstrates its practical application and effectiveness. The model is developed as a methodological approach to prompt engineering in library practice. The study methodology comprised model development, engineering of the series of prompts of varying “strength” based on its principles, and testing with 6 neural network models to accomplish the task of bibliographic data analysis. The testing results demonstrate direct correlation between the quality of neural network response and the level of prompt sophisti-

cation based on the proposed model. The practical significance of the study lies in offering methodological framework for improving the efficiency and quality of interaction between library specialists and artificial intelligence, as well as the possibility of using the model to develop practical recommendations and educational programs on prompt engineering in the library sphere.

Keywords: artificial intelligence, librarianship, prompt engineering, prompt, artificial neural networks, communication, information and communication technologies, competencies, analytical work, data analysis, modeling

Cite: Kovalevsky A. V. The model of communication with artificial intelligence as a methodological approach to prompt creation and evaluation // Scientific and technical libraries. 2025. No. 7, pp. 142–163. <https://doi.org/10.33186/1027-3689-2025-7-142-163>

Введение

Современная библиотечная сфера переживает цифровую трансформацию, обусловленную стремительным развитием технологий искусственного интеллекта (ИИ). ИИ-приложения открывают новые возможности для библиотечных специалистов при решении различных задач (от аналитических до прикладных) в области библиографии. Ключевое условие освоения этих технологий – формирование навыков эффективного взаимодействия с ИИ, в частности, умение составлять качественные **промпты** (англ. prompt – подсказка) – запросы к нейросети на естественном языке, целью которых является получение желаемого конкретного результата.

Для эффективного взаимодействия с нейросетями необходимо правильно сформулировать промпт [1], обращая внимание не только на синтаксис и семантику, но и на формулировку и порядок слов, так как это существенно влияет на поведение модели и качество её ответа [1, 2. С. 17].

Навыки промпт-инжиниринга (методологии написания эффективных запросов к ИИ) позволят библиотекарю повысить качество решения профессиональных задач и задач, выходящих за рамки традиционных компетенций.

Исследования демонстрируют [2–5], что даже небольшие изменения промпта приводят к более качественным результатам.

На данный момент промпт-инжиниринг является преимущественно практической областью, требующей систематического и осознанного подхода. Единственно правильного метода, гарантирующего качественный результат, не существует [2. С. 17].

Навыки промпт-инжиниринга, влияющие на качество получаемых от нейросетей результатов, можно развить благодаря обучению. Эмпирические исследования показывают, что более опытные в промпт-инжиниринге пользователи способны добиваться более точных и качественных результатов от языковых моделей. Результаты исследования [5. С. 9] указывают на то, что между ИИ-грамотностью и способностью формулировать эффективные промпты существует комплексная и многогранная связь. Понимание возможностей и ограничений инструментов на основе ИИ – важный фактор в развитии навыков промпт-инжиниринга [5].

В целях преодоления указанных вызовов и создания методологической основы эффективного промпт-инжиниринга, примененного и в библиотечной сфере, нами была разработана описанная в данной статье модель коммуникации с ИИ ДРУГ.

Компоненты модели ДРУГ

Модель ДРУГ включает в себя Доброжелательный, Рациональный, Уточняющий, Гносеологический компоненты.

Преимущество нашего подхода заключается в эффективной коммуникации с ИИ с учётом дальнейших тенденций развития нейросетей. Эти универсальные правила взаимодействия применимы для любой текстовой генеративной модели с корреляцией на качество и количество её параметров. Данный подход поможет лучше структурировать запрос, повысить качество коммуникации и точность ответа нейросети.

Расскажем о компонентах данной модели.

Доброжелательный – характеризует тон общения с ИИ и качество ответов. Подразумевает соблюдение правил этикета и принципов цифровой этики в диалоге, использование «вежливой» лексики, позитивные формулировки запросов, психологически комфортную коммуникацию, уважение к ИИ как к новому субъекту коммуникации.

Исследователи утверждают, что доброжелательная (до определённого предела) коммуникация повышает качество ответов нейросетей [6, 7].

Так как нейронные сети обучаются на данных, в том числе созданных людьми, можно говорить и про специфические особенности поведения, усвоенные из исходных обучающих данных.

Как позитивная коммуникация может преобразовать коллектив и распространяться с одного уровня на другой [8, 9], так и негативное поведение ухудшает качество коммуникации и снижает её эффективность [10].

Авторы статьи «Language Models in Dialogue: Conversational Maxims for Human-AI Interactions» [11] называют хорошие манеры и доброжелательность важными максимами для коммуникации с ИИ.

Минимизация нежелательных и грубых ответов является важным моментом при работе с ИИ. Как отмечают М. Орт, Яньчао Ю, Чжиюань Тан в выводах своей статьи: «...ни один из методов не смог полностью устранить недопустимые выражения, не изменив при этом контекстные нюансы исходных ответов» [12. С. 8255]. Это говорит о том, что простые фильтры не работают, необходимо более тонкое понимание коммуникации. Авторы также подчёркивают, что «сложности, связанные с оценкой грубости, оставались серьёзной проблемой для моделей ИИ, в первую очередь из-за сложной природы человеческого общения...» [Там же].

Подытожим вышесказанное: ИИ обучается на тех же паттернах поведения, что и человек. Поэтому коммуникация с ИИ должна быть доброжелательной.

В работе «A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT» [13] представлен анализ паттернов промптов, которые могут быть использованы ИИ. Они включают в себя: название и классификацию, назначение и контекст, мотивацию, структуру и ключевые идеи, пример реализации и последствия. В противовес данным паттернам ИИ мы предлагаем инверсионный подход, заключающийся в применении этой же структуры, но уже как аналитической модели для формирования человеческих запросов, где эти паттерны могут быть интерпретированы в рамках рационального, уточняющего и гносеологического компонентов. Подчеркнём, что и человек, и ИИ руководствуются рациональными надстройками, что проявляется и при обучении ИИ, и при составлении промпта человеком.

Модели ИИ и подходы к их ответам развиваются. Существует множество техник промпт-инжиниринга: от простых – zero-shot prompting (промптинг с нулевым количеством примеров) и few-shot prompting (промптинг с несколькими примерами) до более сложных – chain-of-thought prompting (промптинг с цепочкой рассуждений), Retrieval Augmented Generation (RAG) (генерация с расширенным поиском) и Self-Consistency (самосогласованность) [14].

Также появляются новые подходы к предоставлению ответов, использующие внешние источники знаний и улучшающие способность моделей к последовательному логическому мышлению. Например, фреймворк Search-01 [15] демонстрирует, как интеграция агента для поиска информации в процессе рассуждений и использование модуля *Reason-in-Documents* для обработки и уточнения результатов поиска повышают точность и связность ответов, особенно в сложных задачах, требующих доступа к внешним знаниям. Этот подход позволяет моделям не просто генерировать текст, а формировать более обоснованные и надёжные ответы, основанные на анализе проверенной информации.

В основе промпт-инжиниринга лежит рациональный подход, предполагающий учёт возможностей и ограничений языковой модели, а также осознанное конструирование промптов для достижения желаемых результатов.

Рациональный компонент определяет содержательную основу промпта. Он характеризуется логичностью, последовательностью, обоснованностью, соблюдением причинно-следственных связей, рациональностью запроса. Этот компонент задаёт функциональную роль ИИ, например, «ты – аналитик данных», «ты – редактор». Это обуславливает ожидаемый тип ответа и стиль взаимодействия. Также требуются последовательное изложение сути задачи и описание основных требований.

При написании промптов важно учитывать, что модели ИИ, даже применяя метод chain-of-thought, все ещё склонны «считать» людей более рациональными, чем они есть на самом деле [16]. Учитывая эту особенность ИИ, важно соблюдать баланс между логикой и контекстом, поскольку ИИ лучше справляется с формальной рациональностью, чем с учётом контекста и нюансов человеческого мышления.

Уточняющий компонент дополняет рациональный и добавляет контекстные детали – специфические требования к ответу и релевантные примеры, делающие ответ ИИ более точным и адаптированным к задаче.

Уточняющий компонент включает:

детализацию ключевых аспектов запроса (уточнение временных рамок, географии, целевой аудитории, методологии);

предоставление необходимого контекста (информационная потребность, предметная область, предыдущий диалог, ожидаемый уровень ответа в зависимости от уровня знаний);

определение желаемого формата ответа (текст, структурированный список, таблица, диаграмма, программный код, аналитический отчёт);

включение релевантных примеров и аналогий (ожидаемый стиль ответа, другие запросы и ответы на них, шаблоны или паттерны структурирования информации);

указание конкретных ожиданий и результата (объём ответа, степень подробности);

стиль ответа (формальный, неформальный, аналитический, описательный, в зависимости от установки роли (например, ELI5 – explain like I am five));

степень детализации (поверхностный или глубокий анализ), а также ключевые моменты (определённые детали), которые обязательно должны быть отражены в ответе.

Уточняющий компонент – связующее звено между промптами в случае, если задача не решается за один подход. При решении особенно сложных задач, требующих множества шагов или итераций, уточняющий компонент обеспечивает связь контекста и последовательное продвижение к цели. Он позволяет уточнять запрос на каждом шаге, опираясь на предыдущие ответы ИИ, и направлять диалог в нужное русло.

Большинство сервисов, предоставляющих доступ к нейросетям, предупреждают, что пользователь должен проверять сгенерированную информацию, так как она может быть ложной. Полученные от нейросетей ответы нуждаются в проверке и фактчекинге. И, хотя с развитием технологий ИИ качество и истинность ответов улучшаются, проблемы, связанные с «галлюцинациями» и недостаточно точными или некорректными ответами остаются [2, 17, 18. С. 2].

Рекомендуется проверять полученные от ИИ данные, используя надёжные источники информации, входящие в инструментарий библиотечных специалистов.

Гносеологический компонент (от греч. *gnosis* – знание и *logos* – учение), ориентированный на теорию познания и получение достоверного знания, служит отправной точкой написания промпта и получения ответа от ИИ. Основой гносеологического компонента являются стремление к истине и оценка результатов промпта в контексте истинности знания. Гносеологический компонент включает: ориентацию на получение объективного знания, критическое осмысление получаемой информации, верификацию данных, анализ полученных результатов для последующей оптимизации запросов.

Необходимы постоянное «вопросание» и критическое осмысление ответов ИИ, диалог с ИИ, а не слепое доверие к его алгоритмам [18].

Согласимся с авторами статьи «Empowering service systems through Intelligence Augmentation (IA) in digital society» [19] в том, что необходимо сфокусироваться на развитии навыков работы с технологиями ИИ, стимулируя тем самым совместное развитие человека и ИИ.

Используя четыре вышеназванных компонента, мы будем выстраивать дружелюбную, коммуникабельную, качественную коммуникацию, в основе которой стремление к истине посредством этичного взаимодействия. Коммуникация с ИИ даст возможность развить и усилить компетенции, сохранив при этом гуманистический фактор.

В рамках предложенной модели компоненты рассматриваются как взаимосвязанные и взаимообуславливающие. Раскроем эти взаимосвязи:

1. *Доброжелательный и Рациональный компоненты* – сочетание позитивного подхода в коммуникации с ИИ с логической структурированностью и точностью запросов. Доброжелательность, проявляемая в дружелюбном и уважительном тоне («привет», «спасибо», «пожалуйста» и т. д.), способствует более качественному диалогу с ИИ. Рациональность, выражающаяся в точности и логичности запросов, обеспечивает получение релевантных ответов. Таким образом, доброжелательность и рациональность обеспечивают корректную структурированную коммуникацию.

2. *Рациональный и Уточняющий компоненты* – логическое построение запроса в сочетании с его детализацией. Рациональность, как было отмечено выше, обеспечивает точность и логичность запроса.

Однако для более эффективного взаимодействия с ИИ, особенно для решения сложных задач, необходимо уточнение, предполагающее детальную проработку запроса, его развёртывание и конкретизацию требований к результату. Сочетание этих двух компонентов минимизирует неоднозначность ответов и повышает их точность.

3. *Уточняющий и Гносеологический компоненты* – стремление к точному и верифицируемому знанию, опирающемуся на детально проработанный запрос. Уточняющий компонент обеспечивает детальную проработку и конкретизацию запроса. Гносеологичность, являясь отправной точкой составления запроса и конечным пунктом оценки его результата ориентирует на получение достоверного, обоснованного и верифицируемого знания. Таким образом, сочетание этих двух компонентов обеспечивает не только однозначное, но и более надёжное знание, полученное в результате взаимодействия с ИИ.

4. *Гносеологический и Доброжелательный компоненты* – этический аспект взаимодействия с ИИ, в котором стремление к корректному ответу и познанию сочетается с уважительным отношением к источнику предоставления знания как субъекту коммуникации. Такой подход гарантирует, что знания, полученные в результате взаимодействия с ИИ, были получены без нарушения этических норм, являющихся важным элементом коммуникации.

Выделим следующие *принципы применения* модели:

целостность – все компоненты применяются в единстве,

гибкость – соотношение компонентов может варьироваться в зависимости от цели запроса,

итеративность – возможность улучшения промпта на основе полученных результатов,

адаптивность – учёт специфики конкретной ИИ-системы,

развитие – по мере развития ИИ-технологий модель может дополняться новыми аспектами.

Модель ДРУГ – комплексный подход к составлению промптов, учитывающий различные аспекты взаимодействия с ИИ (от технических до этических). Модель может служить методологической основой для разработки практических рекомендаций по эффективной коммуникации с ИИ.

Применение модели ДРУГ на практике

Для иллюстрации работы модели ДРУГ была взята не тривиальная информационно-аналитическая задача по созданию скрипта для выявления количества курсовых, дипломных и магистерских работ в неотфильтрованной таблице (17 653 строки) проверок документов экспертов вуза, выгруженной из системы «Антиплагиат».

Главное условие – получение корректного ответа за один шаг без дополнительного обращения к нейросети.

На основе применения модели ДРУГ было создано пять промптов с градацией в соответствии с силой промпта, где 5 – наилучший промпт, а 1 – наихудший.

Несмотря на то, что данные примеры промптов включают все компоненты модели ДРУГ, *гносеологический компонент* проявляется не только в стремлении к истинности полученного ответа, но и в оценке качества исполнения промптов моделью ИИ, что выходит за рамки промпта, выражаясь в осмыслении полученных ответов. Приведём примеры промптов от 5 до 1.

Промпт уровня 5

«Привет, ты – data analyst, ты будешь оказывать профессиональную помощь в разработке аналитического скрипта для обработки библиографических данных.

Исходные параметры задачи:

файл формата CSV под названием «Docs-report_2024.12.04_111337.csv» объёмом 17653 записи

- ключевой столбец анализа: столбец С (данные об авторах)
- столбец В (тип работы)

Необходима последовательная реализация следующих этапов обработки:

1. Фильтрация записей по столбцу С:
 - критерий: отсутствие символа запятой
 - результат: сохранение только монографических записей
2. Категоризация отфильтрованных данных по столбцу В:
 - категория "Курсовые работы": ключевые слова "курсовая", "курс"
 - категория "Дипломные работы": ключевые слова "диплом", "дипл"
 - категория "Магистерские работы": ключевое слово "магистр"

Предоставь готовый скрипт с детальными комментариями для обработки файла Docs-report_2024.12.04_111337.csv, обеспечивающий валидацию данных и формирование структурированного выходного результата по указанным категориям.

Большое спасибо за помощь!».

Комментарии по применению модели ДРУГ к промπτу уровня 5

Доброжелательный компонент (Д):

вежливое обращение (использовано приветствие);
позитивный тон (выражение «оказывать профессиональную помощь» задаёт позитивный и уважительный тон взаимодействия);
благодарность (выражение благодарности «большое спасибо за помощь!»).

Рациональный компонент (Р):

чёткое определение роли ИИ (задана функциональная роль «data analyst», что определяет ожидаемый тип ответа и стиль взаимодействия);
сформулирована конкретная задача (разработка аналитического скрипта для обработки библиографических данных);
последовательность действий (задача разбита на этапы обработки данных (фильтрация, категоризация), что обеспечивает логическую структуру запроса);

указание входных параметров (чётко определены входные данные: формат файла (CSV), название файла, объём данных, ключевые столбцы для анализа).

Уточняющий компонент (У):

детализация ключевых аспектов запроса (уточнены формат файла (CSV), название файла, объём данных (17653 записи), ключевые столбцы (В и С));

предоставление необходимого контекста (описаны типы данных (библиографические данные, данные об авторах, типы работ));

указание конкретных ожиданий и результата (предоставить «готовый скрипт с детальными комментариями», «обеспечивающий валидацию данных» и «формирование структурированного выходного результата»).

определение формата желаемого ответа (ожидается скрипт (код), а также комментарии к нему).

Промпт уровня 4

«Привет! Требуется разработать скрипт для анализа библиографической базы данных.

Параметры:

– CSV-файл с названием «Docs-report_2024.12.04_111337.csv» (17653 строк)

– анализ столбцов В и С

– фильтрация по наличию запятых в столбце С

– категоризация по ключевым словам "курсовая", "курс", "диплом", "дипл", "магистр" в столбце В

Необходимо:

1. Удалить строки с запятыми в столбце С

2. Распределить данные по типам работ

3. Сформировать отчёт по категориям

Пожалуйста, предоставь скрипт для Docs-report_2024.12.04_111337.csv.».

Комментарии по применению модели ДРУГ к промпту уровня 4

Доброжелательный компонент (Д):

вежливое обращение (использовано приветствие, просьба со словом «пожалуйста»).

Рациональный компонент (Р):

определение задачи (разработка скрипта для анализа библиографической базы данных);

указание параметров (формат файла (CSV), название файла, объём данных, столбцы для анализа (В и С), критерии фильтрации и категоризации);

перечисление необходимых действий: задача разбита на шаги (удалить строки, распределить данные, сформировать отчёт).

Уточняющий компонент (У):

меньшая детализация, чем в промпте 5 (отсутствует указание на роль «data analyst», менее детально описан ожидаемый результат (убраны требования к комментариям, валидации данных, структурированному выходному результату));

сохранение основных параметров (остаются указанными формат файла, название, объём, столбцы, критерии фильтрации и категоризации).

Промпт уровня 3

«Привет! Нужен скрипт для обработки CSV файла “Docs-report_2024.12.04_111337.csv”.

Надо:

1. Убрать все строки, где в столбце С есть запятые
2. Найти в столбце В работы по словам:
курсовая/курс,
диплом/дипл,
магистр.

Сделай скрипт для файла Docs-report_2024.12.04_111337.csv»

Комментарии по применению модели ДРУГ к промπτу уровня 3

Доброжелательный компонент (Д):

сохранено только приветствие «привет!».

Рациональный компонент (Р):

общее описание задачи (задача описана более общей фразой «нужен скрипт для обработки CSV-файла»);

упрощённое описание действий (шаги обработки данных описаны в более простой форме: «убрать строки», «найти работы по словам»);

менее чёткие параметры (название файла, столбцы указаны менее явно, чем в промптах 5 и 4).

Уточняющий компонент (У):

значительное снижение детализации (отсутствует указание на объём данных, формат файла (явно не указан, хотя подразумевается CSV), менее детализированы критерии категоризации (использованы только ключевые слова, без указания категорий));

ожидаемый результат описан очень кратко («сделай скрипт для файла»).

Промпт уровня 2

«Привет. Нужно написать скрипт для CSV файла. Удали строки с запятыми в столбце С и найди в столбце В курсовые, дипломы и магистерские. Файл Docs-report_2024.12.04_111337.csv».

Комментарии по применению модели ДРУГ к промπτу уровня 2

Доброжелательный компонент (Д):

приветствие «привет».

Рациональный компонент (Р):

задача описана ещё более общими словами («нужно написать скрипт для CSV-файла»);

параметры очень размыты.

Уточняющий компонент (У):

детализация минимальна, запрос очень общий и неструктурированный.

Промпт уровня 1

«Напиши скрипт для файла, чтобы убрать запятые и найти курсовые дипломы и магистерские».

Комментарии по применению модели ДРУГ к промπτу уровня 1

Можно отметить полное или почти полное отсутствие всех компонентов модели ДРУГ, кроме базового запроса на создание скрипта. Данный уровень хоть и является простым и гипотетическим, однако используется многими пользователями при обращении к нейросетям (из практики преподавания блока про ИИ в рамках учебных дисциплин «Медиаграмотность в цифровом пространстве», «Информационно-библиографическая культура современного исследователя», а также курсов повышения квалификации профессорско-преподавательского состава в Научной библиотеке БНТУ).

Каждый последующий уровень модели ДРУГ демонстрирует снижение уровня доброжелательности, ухудшение рациональной структуры, сокращение уточняющих деталей, размытие и потерю гносеологической составляющей.

Для проверки и оценки промπτов было использовано шесть моделей нейросетей: Claude 3.5 Sonnet [20], ChatGPT 4o [21], Gemini 2 Flash Experimental [22], Microsoft Copilot [23], Mistral AI [24], YandexGPT 4 [25]. Скрипты тестировались в сервисе для исполнения кода Google Colab [26].

Полученные результаты и их интерпретация представлены в виде таблицы, где каждой модели присваивалась оценка в зависимости от точности решения задачи.

Главные критерии оценки: справилась ли нейросеть с задачей за один шаг и дала ли она корректное количество курсовых, дипломных и магистерских работ.

Оценка эффективности моделей ИИ по уровням промптов

Модель ИИ	Промпт 5	Промпт 4	Промпт 3	Промпт 2	Промпт 1	Итого
Claude 3.5 Sonnet	Да	Да	Да	Да *	Да/Нет	3.5
ChatGPT 4o	Да	Да	Да/Нет	Да/Нет	Да/Нет	2
Gemini 2.0 Flash Experimental	Да	Да/Нет	Да/Нет	Нет	Нет	1
Microsoft Copilot	Да	Нет	Нет	Нет	Да/Нет	1
Mistral AI	Да/Нет	Да	Да/Нет	Да/Нет	Да/Нет	1
YandexGPT 4	Нет	Да/Нет	Да/Нет	Да/Нет	Да/Нет	0

Ключ к оценке:

Да – скрипт работает, итоговое число корректное (1 балл).

Да/Нет – скрипт написан, но итоговое число некорректно или отсутствует вовсе, либо невозможно подсчитать корректно в силу форматирования, соответственно, задача не решена полностью (0 баллов).

Нет – скрипт не работает, выдаёт ошибку (0 баллов).

Описание результатов, полученных для каждой модели и каждого промпта, приведено ниже.

Claude 3.5 Sonnet (3.5 баллов из 5). Модель продемонстрировала способность выдавать корректный результат при использовании наиболее точных и чётко сформулированных промптов (5 и 4), корректно выделяя нужную информацию. В частности, в случае промпта 5 модель предоставила готовое число, а в случае промпта 4, кроме готового числа, также был представлен процентный показатель по отношению к общему количеству строк. В промпте 3 появился информационный шум в виде дублирования ключевых слов, также итоговый файл не содержал деления работ на категории. При использовании промпта 2 модель предоставила только общее число, не разделяя его на категории, показав при этом практически точное попадание (за что модель получила 0,5 балла). При использовании самого простого и не детализированного промпта 1 скрипт работал, однако итоговый файл представлял просто массив данных без запятых и с некорректными статистическими данными в конце. Данная модель по праву занимает первое место по качеству ответов и детальному описанию.

ChatGPT 4o (2 балла из 5). Модель показала стабильную работу при использовании промптов 5 и 4 – в обоих случаях было итоговое число, а также разделение на категории. В случае промпта 4 был получен готовый файл с суммой, что позволяет отметить это решение как более качественное, чем решение промпта 5. В случае промпта 3 модель предоставила неизменённую таблицу без проставления категорий, а в случае промпта 2 скрипт работал, но предоставил нерелевантный файл с указанием семи магистерских работ, а также неизменённую таблицу без проставления категорий. В случае промпта 1 скрипт не работал.

Gemini 2 Flash Experimental (1 балл из 5). Модель показала способность выдавать корректный результат только в случае промпта 5, предоставляя итоговое число, а также файл с подсчётом. В остальных случаях были обнаружены недочёты. В случае промпта 4 было некорректное итоговое число. В случае промпта 3 – отредактированный файл, но с возможностью только ручного подсчёта. В случае промпта 2 модель отказалась писать скрипт, ссылаясь на ограничения в текущей версии. В промпте 1 модель также отказалась от написания скрипта, ссылаясь на расплывчатость задачи и предлагая исполнить два скрипта (удаление запятых и поиск по ключевым словам) отдельно друг от друга.

Microsoft Copilot (1 балл из 5). Модель продемонстрировала способность выдавать корректный результат только в случае промпта 5, предоставляя таблицу с указанием категорий и верным итоговым количеством. В случае промптов 4, 3 и 2 скрипт не работал, выдавая ошибку. В случае промпта 1 скрипт работал, однако выдавал нулевой результат, предоставляя неизменённый файл.

Mistral AI (1 балл из 5). Модель показала неоднозначные результаты. При промпте 5 скрипт работал, предоставляя таблицу с указанными категориями, однако итоговое количество было неверным из-за некорректно отобранных ключевых слов. При промпте 4 модель предоставила корректный файл с готовым числом, разделив при этом курсовые работы на «курсовая» и «курс». В случае промпта 3 количество строк в итоговом файле составляло 461, что было значительно меньше совокупного количества правильных ответов. В случае промпта 2 скрипт работал, но выводил пустую таблицу, промпта 1 скрипт работал, но выдавал пустой файл.

YandexGPT 4 (0 баллов из 5). Модель продемонстрировала низкие результаты по всем уровням промптов. Изначальное максимальное количество символов промпта 5 превысило допустимые значения модели (не более 1 тыс. символов), поэтому промпт был незначительно скорректирован. После сокращения количества символов промпта 5 скрипт всё равно не работал. При использовании остальных промптов скрипты запускались, но итоговые результаты были некорректными. В случае промпта 4 итоговое число было неверным, промпта 3 – отсутствовало. В случае промпта 2 получился набор неотфильтрованных данных. При использовании промпта 1 скрипт запустился, но выдал нулевой результат.

Тестирование показало, что наиболее сильный промпт 5 дал самое большое количество решённых за один шаг задач (4 из 6). С уменьшением силы промпта (4) количество правильных ответов стало 3 из 6. Дальнейшее ухудшение качества промптов сильнее сказалось на ответах нейросетей: только Claude 3.5 Sonnet в промпте 3 предоставил корректный результат. С промптами 2 (за исключением частичного решения задачи Claude 3.5 Sonnet) и 1 не справилась ни одна нейросеть.

Заключение

Модель коммуникации с ИИ ДРУГ – методологический подход, разработанный для оптимизации составления и оценки промптов в библиотечной сфере. Составление промптов на основе этого подхода предлагает комплексный инструментарий для повышения эффективности взаимодействия библиотечных специалистов с ИИ.

Тестирование модели, проведённое в разных сервисах ИИ на примере анализа данных из системы «Антиплагиат», продемонстрировало прямую зависимость качества ответов нейросетей от уровня проработанности промпта.

Предложенный подход позволяет не только повысить точность и релевантность ответов ИИ, но и способствует формированию этически выверенного и ориентированного на достоверное знание взаимодействия. Мы полагаем, что дальнейшее совершенствование нейросетей не будет влиять на использование предложенных компонентов модели ДРУГ в силу её универсального подхода.

Апробация модели была проведена на ограниченном количестве задач и сервисов ИИ, что не позволяет с абсолютной уверенностью экстраполировать результаты на все возможные сценарии взаимодействия. И, хотя модель разрабатывалась с учётом возможных изменений в технологиях ИИ, исследование проводилось в условиях текущего уровня развития нейросетей. При значительных изменениях технологии эффективность модели может потребовать дополнительной оценки.

В дальнейшем предполагается протестировать модель при решении других информационно-аналитических задач в библиотечной сфере, а также разработать программу обучения библиотечных специалистов по применению модели ДРУГ с учётом уровня их компетенций.

Модель коммуникации ДРУГ – пример осознанного и эффективного использования потенциала технологий ИИ в библиотечной сфере, открывающий новые возможности для развития библиотечных сервисов и повышения профессиональных компетенций библиотечных работников в области информационно-коммуникационных технологий.

СПИСОК ИСТОЧНИКОВ

1. **Dang H., Mecke L., Lehmann F., Goller S., Buschek D.** How to prompt? Opportunities and challenges of zero- and few-shot learning for human-AI interaction in creative applications of generative models // arXiv preprint. 2022. URL: <https://arxiv.org/abs/2209.01390> (access: 15.01.2025).
2. **Kaddour J., Harris J., Mozes M., Bradley H., Raileanu R., McHardy R.** Challenges and Applications of Large Language Models // arXiv preprint. 2023. URL: <https://arxiv.org/abs/2307.10169> (access: 14.01.2025).
3. **Zihao Z., Wallace E., Feng S., Klein D., Singh S.** Calibrate Before Use: Improving Few-shot Performance of Language Models // Proceedings of the 38th International Conference on Machine Learning, PMLR 139. 2021. URL: <https://proceedings.mlr.press/v139/zhao21c.html> (access: 21.01.2024).
4. **Webson A., Pavlick E.** Do Prompt-Based Models Really Understand the Meaning of Their Prompts? // Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies. Association for Computational Linguistics, 2022. C. 2300–2344.
5. **Knoth N., Tolzin A., Janson A., Leimeister J. M.** AI literacy and its implications for prompt engineering strategies // Computers and Education: Artificial Intelligence. 2024. T. 6. URL: <https://doi.org/10.1016/j.caeai.2024.100225> (access: 22.01.2025).
6. **Yin Z., Wang H., Horio K., Kawahara D., Sekine S.** Should We Respect LLMs? A Cross-Lingual Study on the Influence of Prompt Politeness on LLM Performance // The 2nd Workshop on Social Influence in Conversations (SICoN): Proceedings of the Workshop. Association for Computational Linguistics, 2024. C. 9–35.
7. **Пиперски А. Ч.** Вежливость в коммуникации между человеком и искусственным интеллектом // Слово.ру: балтийский акцент. Т. 15, № 4. С. 89–98.
8. **Bass B. M., Waldman D. A., Avolio B. J., Bebb B. J.** Transformational Leadership and the Falling Dominoes Effect // Group & Organization Studies. 1987. Т. 12, № 1. С. 73–87.
9. **Mayer D. M., Kuenzi M., Greenbaum R., Bardes M., Salvador R.** How low does ethical leadership flow? Test of a trickle-down model // Organizational Behavior and Human Decision Processes. 2009. Т. 108, № 1. С. 1–13.
10. **Foult T., Woolum A., Erez A.** Catching Rudeness Is Like Catching a Cold: The Contagion Effects of Low-Intensity Negative Behaviors // Journal of Applied Psychology. 2016. Т. 101, № 1. С. 50–67.
11. **Miehling E., Nagireddy M., Sattigeri P., Daly E. M., Piorkowski D., Richards J. T.** Language Models in Dialogue: Conversational Maxims for Human-AI Interactions // Findings of the Association for Computational Linguistics: EMNLP. 2024. C. 14420–14437.
12. **Orme M., Yu Y., Tan Z.** How Much do Robots Understand Rudeness? Challenges in Human-Robot Interaction // Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024). Torino: ELRA Language Resource Association, 2024. C. 8247–8257.

13. **White J., Fu Q., Hays S., Sandborn M., Olea C., Gilbert H., Schmidt D. C.** A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT // arXiv preprint. 2023. URL: <https://arXiv:2302.11382> (дата обращения: 14.01.2025).
14. **Prompting** Techniques. URL: <https://www.promptingguide.ai/techniques> (access: 22.01.2025).
15. **Li X., Dong G., Jin J., Zhang Y., Zhou Y., Zhu Y., Zhang P., Dou Z.** Search-o1: Agentic Search-Enhanced Large Reasoning Models // arXiv preprint. 2025. URL: <https://arxiv.org/abs/2501.05366> (access: 22.01.2025).
16. **Liu R., Geng J., Peterson J. C., Sucholutsky I., Griffiths T. L.** Large Language Models Assume People are More Rational than We Really are // arXiv preprint. 2024. URL: <https://arxiv.org/abs/2406.17055> (access: 22.01.2025).
17. **Lewis P., Perez E., Piktus A., Petroni F., Karpukhin V., Goyal N., Küttler H., Lewis M., Yih W., Rocktäschel T., Riedel S., Kiela D.** Retrieval-augmented generation for knowledge-intensive NLP tasks // Proceedings of the 34th International Conference on Neural Information Processing Systems (NIPS '20). 2020. С. 9459–9474.
18. **Jadallah A. K., Mihus I., Svyrydiuk N., Zhyvko Z.** Nature and purpose of artificial intelligence. Political, legal, and economic challenges in the 21st century // CLÍO: Revista de Revista de Historia, Ciencias Humanas y pensamiento crític. 2024. № 8. С. 306–320. DOI <https://doi.org/10.5281/zenodo.12600350>.
19. **Bassano C., Caputo F., Barile P., Piciocchi P.** Empowering service systems through Intelligence Augmentation (IA) in digital society // ITM Web of Conferences. 2024. Т. 62. DOI <https://doi.org/10.1051/itmconf/20246205003>.
20. **Claude**. URL: <https://claude.ai/> (access: 13.01.2025).
21. **ChatGPT**. URL: <https://chatgpt.com/> (access: 13.01.2025).
22. **Gemini**. URL: <https://gemini.google.com/> (access: 13.01.2025).
23. **Microsoft** Copilot. URL: <https://copilot.microsoft.com/> (access: 13.01.2025).
24. **Mistral AI**. URL: <https://mistral.ai/> (access: 13.01.2025).
25. **YandexGPT 4**. URL: <https://ya.ru/ai/gpt-4> (access: 13.01.2025).
26. **Google** Colab. URL: <https://colab.research.google.com/> (access: 13.01.2025).

References

1. **Dang H., Mecke L., Lehmann F., Goller S., Buschek D.** How to prompt? Opportunities and challenges of zero- and few-shot learning for human-AI interaction in creative applications of generative models // arXiv preprint. 2022. URL: <https://arxiv.org/abs/2209.01390> (access: 15.01.2025).

2. **Kaddour J., Harris J., Mozes M., Bradley H., Raileanu R., McHardy R.** Challenges and Applications of Large Language Models // arXiv preprint. 2023.
URL: <https://arxiv.org/abs/2307.10169> (access: 14.01.2025).
3. **Zihao Z., Wallace E., Feng S., Klein D., Singh S.** Calibrate Before Use: Improving Few-shot Performance of Language Models // Proceedings of the 38th International Conference on Machine Learning, PMLR 139. 2021. URL: <https://proceedings.mlr.press/v139/zhao21c.html> (access: 21.01.2024).
4. **Webson A., Pavlick E.** Do Prompt-Based Models Really Understand the Meaning of Their Prompts? // Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies. Association for Computational Linguistics, 2022. C. 2300–2344.
5. **Knoth N., Tolzin A., Janson A., Leimeister J. M.** AI literacy and its implications for prompt engineering strategies // Computers and Education: Artificial Intelligence. 2024. T. 6.
URL: <https://doi.org/10.1016/j.caeai.2024.100225> (access: 22.01.2025).
6. **Yin Z., Wang H., Horio K., Kawahara D., Sekine S.** Should We Respect LLMs? A Cross-Lingual Study on the Influence of Prompt Politeness on LLM Performance // The 2nd Workshop on Social Influence in Conversations (SICon): Proceedings of the Workshop. Association for Computational Linguistics, 2024. C. 9–35.
7. **Piperski A. Ch.** Vezhlivost' v komunikacii mezhdu chelovekom i iskusstvenny'm intellektom // Slovo.ru: baltii'skii' akcent. T. 15, № 4. S. 89–98.
8. **Bass B. M., Waldman D. A., Avolio B. J., Bebb B. J.** Transformational Leadership and the Falling Dominoes Effect // Group & Organization Studies. 1987. T. 12, № 1. C. 73–87.
9. **Mayer D. M., Kuenzi M., Greenbaum R., Bardes M., Salvador R.** How low does ethical leadership flow? Test of a trickle-down model // Organizational Behavior and Human Decision Processes. 2009. T. 108, № 1. C. 1–13.
10. **Foulk T., Woolum A., Erez A.** Catching Rudeness Is Like Catching a Cold: The Contagion Effects of Low-Intensity Negative Behaviors // Journal of Applied Psychology. 2016. T. 101, № 1. C. 50–67.
11. **Miehling E., Nagireddy M., Sattigeri P., Daly E. M., Piorkowski D., Richards J. T.** Language Models in Dialogue: Conversational Maxims for Human-AI Interactions // Findings of the Association for Computational Linguistics: EMNLP. 2024. C. 14420–14437.
12. **Orme M., Yu Y., Tan Z.** How Much do Robots Understand Rudeness? Challenges in Human-Robot Interaction // Proceedings of the 2024 Joint International Conference on Computational Linguistics, Language Resources and Evaluation (LREC-COLING 2024). Torino: ELRA Language Resource Association, 2024. C. 8247–8257.
13. **White J., Fu Q., Hays S., Sandborn M., Olea C., Gilbert H., Schmidt D. C.** A Prompt Pattern Catalog to Enhance Prompt Engineering with ChatGPT // arXiv preprint. 2023.
URL: <https://arXiv:2302.11382> (access: 14.01.2025).
14. **Prompting** Techniques. URL: <https://www.promptingguide.ai/techniques> (access: 22.01.2025).

15. **Li X., Dong G., Jin J., Zhang Y., Zhou Y., Zhu Y., Zhang P., Dou Z.** Search-o1: Agentic Search-Enhanced Large Reasoning Models // arXiv preprint. 2025.
URL: <https://arxiv.org/abs/2501.05366> (access: 22.01.2025).
16. **Liu R., Geng J., Peterson J. C., Sucholutsky I., Griffiths T. L.** Large Language Models Assume People are More Rational than We Really are // arXiv preprint. 2024. URL: <https://arxiv.org/abs/2406.17055> (access: 22.01.2025).
17. **Lewis P., Perez E., Piktus A., Petroni F., Karpukhin V., Goyal N., Küttler H., Lewis M., Yih W., Rocktäschel T., Riedel S., Kiela D.** Retrieval-augmented generation for knowledge-intensive NLP tasks // Proceedings of the 34th International Conference on Neural Information Processing Systems (NIPS '20). 2020. C. 9459–9474.
18. **Jadallah A. K., Mihus I., Svyrydiuk N., Zhyvko Z.** Nature and purpose of artificial intelligence. Political, legal, and economic challenges in the 21st century // CLÍO: Revista de Historia, Ciencias Humanas y pensamiento crític. 2024. № 8. C. 306–320.
DOI <https://doi.org/10.5281/zenodo.12600350>.
19. **Bassano C., Caputo F., Barile P., Piciocchi P.** Empowering service systems through Intelligence Augmentation (IA) in digital society // ITM Web of Conferences. 2024. T. 62.
DOI <https://doi.org/10.1051/itmconf/20246205003>.
20. **Claude.** URL: <https://claude.ai/> (access: 13.01.2025).
21. **ChatGPT.** URL: <https://chatgpt.com/> (access: 13.01.2025).
22. **Gemini.** URL: <https://gemini.google.com/> (access: 13.01.2025).
23. **Microsoft Copilot.** URL: <https://copilot.microsoft.com/> (access: 13.01.2025).
24. **Mistral AI.** URL: <https://mistral.ai/> (access: 13.01.2025).
25. **YandexGPT 4.** URL: <https://ya.ru/ai/gpt-4> (access: 13.01.2025).
26. **Google Colab.** URL: <https://colab.research.google.com/> (access: 13.01.2025).

Информация об авторе / Author

Ковалевский Алексей Викентьевич – магистр педагогических наук, аспирант Белорусского государственного университета культуры и искусств; Научная библиотека Белорусского национального технического университета, Минск, Республика Беларусь. kovalevskyalex@yandex.ru

Aleksey V. Kovalevsky – Master of Science (Pedagogy), postgraduate student, Belarus State University of Culture and Arts, Science Library; Belarus National Technological University, Minsk, Republic of Belarus
kovalevskyalex@yandex.ru