

БОЛЬШЕ
ВОЗМОЖНОСТЕЙ
КАЖДЫЙ ДЕНЬ

ПРОЩЕ,
УМНЕЕ,
ЭФФЕКТИВНЕЕ

ИИ, который работает на вас

Как использовать искусственный
интеллект в работе и жизни,
чтобы экономить время,
зарабатывать больше
и достигать целей



Идеи



Контент



Аналитика



Задачи

*Ваш
личный
ИИ-помощник
уже здесь*



КИМ ВОЛАН

Ким Волан

ИИ, который работает на вас

«Автор»

2026

Волан К.

ИИ, который работает на вас / К. Волан — «Автор», 2026

Вы уже пользуетесь ChatGPT или его аналогами — печатаете вопросы, получаете ответы, иногда впечатляетесь. Но между этим и реальной экономией часов каждую неделю — огромный разрыв. Эта книга — про то, как его закрыть. 25 глав без воды про «промт-инжиниринг ради промт-инжиниринга»: как на самом деле работают языковые модели и что из этого следует на практике; как встроить ИИ в почту, документы, планирование и обучение; как применить его в маркетинге, продажах, аналитике, найме, финансах, юридических вопросах, управлении проектами и CRM; как построить своих ИИ-помощников и не потерять голову от мультимодальных инструментов. Каждая глава заканчивается разделом «Частые ошибки» и конкретными шагами на неделю. Технологии продолжают меняться — принципы в этой книге переживут любое обновление интерфейса.

© Волан К., 2026

© Автор, 2026

Содержание

От автора	5
Часть 1. Фундамент: как думать об ИИ, а не просто им пользоваться	6
Глава 1. От игрушки к инструменту	6
Глава 2. Как на самом деле работают языковые модели	8
Глава 3. Анатомия хорошего запроса	10
Глава 4. Контекст — главный рычаг качества	12
Часть 2. Личная продуктивность с ИИ	14
Глава 5. Почта и коммуникация	14
Глава 6. Работа с документами и знаниями	16
Конец ознакомительного фрагмента.	17

Ким Волан

ИИ, который работает на вас

От автора

Почти все, кто читает эту книгу, уже пользовались ChatGPT, Claude или похожим ассистентом. Попросили написать письмо, сократить текст, придумать пару идей для поста. И почти все на этом остановились — потому что дальше начинается разрыв между «я умею печатать вопросы в чат» и «я встроил ИИ в то, как устроена моя работа». Этот разрыв — не про технические навыки программиста. Он про мышление: как задавать не один хороший запрос, а выстраивать вокруг ИИ повторяемый процесс, который экономит часы каждую неделю, а не разово впечатляет один раз.

Эта книга написана не для тех, кто никогда не открывал чат с нейросетью, и не для разработчиков, которые хотят обучать собственные модели. Она для тех, кто уже освоился с базовым использованием и хочет перейти на следующий уровень: от разового «напиши мне текст» к системному использованию ИИ в переписке, документах, планировании, маркетинге, продажах, аналитике, найме и автоматизации рутины.

Разговоры об искусственном интеллекте последних лет колеблются между двумя крайностями: либо ИИ представляют как волшебную кнопку, решающую любую задачу без усилий, либо как модную игрушку, годную разве что для смешных картинок. Обе крайности одинаково бесполезны на практике. Реальная выгода от ИИ в работе появляется не от веры в его всемогущество и не от скепсиса, а от конкретного, методичного встраивания инструмента в конкретные, повторяющиеся задачи — там, где рутина отнимает время, которое можно направить на то, что действительно требует человеческого суждения.

Книга разделена на четыре части. Первая — про фундамент: как на самом деле работают языковые модели на уровне, достаточном для практики, и что отличает посредственный запрос от того, что реально экономит время. Вторая — про личную продуктивность: почту, документы, планирование, обучение. Третья — про то, как ИИ встраивается в конкретные бизнес-функции: маркетинг, продажи, аналитику, автоматизацию, найм. Четвёртая — про переход от пользователя готовых инструментов к архитектору собственной системы: персональные ассистенты, мультимодальные инструменты, проверка качества и понимание границ того, что ИИ не должен решать за человека.

В каждой главе — раздел «Частые ошибки», собранный из типичных провалов при внедрении ИИ в реальную работу, и «Что сделать на этой неделе» — конкретные шаги, а не абстрактные советы. Технологии, о которых пойдёт речь, продолжают меняться и после выхода этой книги — новые модели и интерфейсы появляются каждые несколько месяцев. Поэтому акцент сделан не на конкретных кнопках конкретных интерфейсов, а на принципах, которые остаются рабочими вне зависимости от того, какая именно модель стоит за чат-окном в момент чтения.

Часть 1. Фундамент: как думать об ИИ, а не просто им пользоваться

Глава 1. От игрушки к инструменту

Почему большинство людей используют ИИ на десять процентов возможностей

Типичный сценарий знакомства с ИИ выглядит так: человек открывает чат, задаёт один вопрос, получает ответ, впечатляется или разочаровывается — и либо забывает про инструмент на неделю, либо использует его время от времени для разовых мелких задач вроде «сократи этот текст» или «предложи заголовок». На этом уровне ИИ действительно похож на игрушку: забавную, иногда полезную, но никак не меняющую структуру рабочего дня.

Разница между этим уровнем и системным использованием — не в объёме знаний о технологии, а в смене самого вопроса, который человек себе задаёт. Вместо «что ИИ может сделать интересного» продуктивный вопрос звучит так: «какая из моих регулярных задач отнимает время непропорционально своей сложности, и можно ли делегировать эту задачу инструменту, освободив время для того, что действительно требует моего суждения». Это смещение фокуса с возможностей инструмента на структуру собственной работы — ключевой шаг, который пропускает большинство пользователей.

Полезное упражнение на старте — провести честный аудит недели: выписать все задачи, отнимающие больше получаса, и отметить, какая их часть является рутинной обработкой информации — чтением, обобщением, черновым написанием, структурированием — а какая требует уникального суждения, знания контекста конкретной ситуации или личных отношений с людьми. Первая категория задач — основной кандидат на делегирование ИИ; вторая почти всегда должна остаться за человеком, даже если ИИ формально способен сгенерировать правдоподобный черновик.

Одна из причин, почему люди застревают на поверхностном уровне использования, — восприятие ИИ как поисковика с более вежливым тоном. Поисковик отвечает на разовый запрос и забывает о нём в следующую секунду. Современный ассистент, использованный правильно, способен удерживать сложный, многошаговый контекст задачи на протяжении длинного диалога, уточнять детали, предлагать альтернативы и дорабатывать результат по итерациям — но эта возможность остаётся невостребованной, если пользователь каждый раз открывает новый чат вместо того, чтобы развивать один диалог вглубь конкретной задачи.

Второй барьер — ожидание идеального результата с первой попытки. На практике первый ответ ИИ на сложный запрос — это почти всегда черновик, отправная точка для диалога, а не финальный продукт. Пользователи, разочаровавшиеся в инструменте после одной неудачной попытки, часто просто не дошли до второго и третьего шага уточнения, на которых обычно и раскрывается реальная польза: «это близко, но перепиши раздел про цифры более сдержанно» работает значительно лучше, чем повторный запрос с нуля в новом чате.

Третий барьер — использование ИИ исключительно для генерации нового контента, при полном игнорировании его способности обрабатывать уже существующую информацию: анализировать длинные документы, находить противоречия в собственных заметках, структурировать хаотичный поток мыслей в связный план. Генерация — лишь часть возможностей; не менее ценная часть — понимание, обобщение и структурирование того, что уже есть, но разбросано по разным файлам, письмам и заметкам.

Переход от игрушки к инструменту не требует технической подготовки — требуется смена привычки: прежде чем взяться за рутинную задачу вручную, на несколько секунд оста-

новиться и спросить себя, можно ли начать её с помощью ИИ хотя бы в черновом виде. Эта привычка, повторённая за несколько недель на разных типах задач, сама подсказывает, где инструмент даёт реальную экономию времени, а где привычные ручные методы по-прежнему быстрее и надёжнее.

Возьмём конкретный пример: менеджер, тратящий по сорок минут каждый понедельник на составление сводки для руководства из разрозненных отчётов по проекту за неделю. После аудита задач по методу, описанному выше, он заметил, что сама сводка — рутинная компиляция уже известных фактов, а не творческая работа. Делегировав первый черновик ИИ, он сократил время до десяти минут на финальную проверку и корректировку тона — освободив полчаса в неделю, которые эта же задача отнимала на протяжении многих месяцев без единой попытки её пересмотреть.

Частые ошибки

Использовать ИИ только для разовых, не связанных между собой мелких задач, не пытаясь выстроить вокруг него повторяемый процесс

Ожидать идеального результата с первой попытки и бросать диалог вместо того, чтобы уточнить и доработать ответ

Ограничиваться генерацией нового текста, игнорируя способность ИИ обрабатывать и структурировать уже существующую информацию

Что сделать на этой неделе

Проведите аудит одной рабочей недели и выпишите все задачи дольше получаса, отделив рутинную обработку информации от задач, требующих личного суждения

Выберите одну рутинную задачу из списка и в течение недели делегируйте её первый черновик ИИ вместо выполнения с нуля вручную

В следующем сложном запросе не открывайте новый чат при неидеальном первом ответе — уточните и доработайте результат минимум два раза подряд

Загрузите в диалог с ИИ один уже существующий документ или набор заметок и попросите структурировать или обобщить его вместо того, чтобы просить что-то написать с нуля

Глава 2. Как на самом деле работают языковые модели

Минимум теории, достаточный для практики

Чтобы эффективно работать с языковой моделью, не нужно разбираться в математике нейронных сетей — но полезно понимать несколько практических следствий того, как она устроена, потому что именно из них вытекают правила хорошей работы с инструментом. Языковая модель предсказывает наиболее вероятное продолжение текста на основе того, что она увидела при обучении и что находится в текущем диалоге. Она не «знает» факты в человеческом смысле — она воспроизводит статистически правдоподобные образцы языка, что объясняет как её впечатляющие сильные стороны, так и характерные слабости.

Первое практическое следствие — модель одинаково уверенно звучит и когда она права, и когда ошибается. В отличие от человека, который обычно интонационно или явно сигнализирует неуверенность, языковая модель по умолчанию формулирует и точный факт, и выдуманную деталь в одинаково гладком, убедительном тоне. Это явление принято называть галлюцинацией — не потому что модель «обманывает» осознанно, а потому что генерация правдоподобного текста и генерация фактически верного текста для неё — разные, не всегда совпадающие задачи. Практический вывод: любой факт, число или ссылку, критически важные для решения, стоит проверять отдельно, а не принимать на веру только потому, что ответ звучит уверенно.

Второе следствие — модель работает в пределах ограниченного «контекстного окна»: определённого объёма текста, который она способна удерживать в рамках одного диалога одновременно. Чем длиннее и разнообразнее становится диалог, тем выше риск, что модель начнёт упускать детали, упомянутые в начале разговора, особенно если между ними прошло много несвязанных сообщений. Практический вывод: для сложных, многошаговых задач полезно периодически явно напоминать модели ключевые условия задачи, а не полагаться на то, что она помнит всё сказанное несколькими десятками сообщений ранее.

Третье следствие — качество ответа сильно зависит от качества и конкретности предоставленного контекста, а не только от формулировки самого вопроса. Модель не имеет доступа к невысказанным предположениям пользователя: если не указана целевая аудитория текста, желаемый тон, ограничение по длине или конкретная цель материала, модель вынуждена угадывать эти параметры на основе общих, усреднённых образцов — и результат получается таким же усреднённым, обезличенным, не подходящим для конкретной ситуации.

Четвёртое следствие касается разницы между моделями, оптимизированными для разных задач: одни лучше подходят для творческого, более свободного генерирования текста, другие — для строгого следования инструкциям и работы с фактической точностью, третьи специально настроены на программирование или математические рассуждения. Использование модели не по её сильной стороне — например, ожидание точных вычислений от модели, оптимизированной под творческий текст, — системная причина разочарования, которую легко спутать с общей ненадёжностью технологии, хотя на деле проблема в выборе инструмента под задачу.

Пятое следствие — способность модели рассуждать пошагово заметно улучшается, если явно попросить её объяснить ход мысли перед итоговым ответом, а не выдавать готовый вывод сразу. Промежуточное рассуждение не только облегчает проверку логики человеком, но и статистически повышает точность самого итогового ответа модели на сложных, многошаговых задачах — эффект, хорошо задокументированный на практике и полезный для любой задачи, требующей не одного факта, а цепочки связанных выводов.

Понимание этих пяти следствий не превращает работу с ИИ в точную науку, но заметно снижает частоту разочарований: пользователь, знающий про ограничение контекстного окна,

не удивляется, когда модель «забывает» условие, сказанное час назад в длинном диалоге; пользователь, знающий про склонность к уверенным галлюцинациям, привычно перепроверяет факты, а не разочаровывается в инструменте после первой обнаруженной ошибки.

На практике эти пять следствий проявляются вместе в одной типичной ситуации: сотрудник просит модель посчитать точную сумму по сложной формуле в середине длинного диалога о маркетинговой стратегии. Модель, оптимизированная для беглого текста, а не точных вычислений, в таком контексте статистически чаще ошибается в арифметике, чем специализированный калькулятор или табличный редактор, — не потому что «не старается», а потому что вычисление точных чисел структурно отличается от генерации правдоподобного текста, для которой она в первую очередь построена.

Частые ошибки

Принимать уверенно звучащий ответ модели за факт без проверки, особенно в задачах с конкретными цифрами, датами или ссылками

Вести очень длинный диалог без напоминания ключевых условий задачи и удивляться, когда модель их «забывает»

Использовать модель не по её сильной стороне — например, ждать точных вычислений от модели, заточенной под творческие тексты

Что сделать на этой неделе

В следующем важном запросе явно попросите модель сначала объяснить ход рассуждения, а затем дать итоговый ответ

Возьмите один недавний ответ ИИ с конкретными фактами или цифрами и перепроверьте их в независимом источнике

В длинном рабочем диалоге раз в 15–20 сообщений кратко напоминайте модели ключевые условия исходной задачи

Определите, для каких из ваших задач стоит использовать разные модели или режимы вместо одного универсального инструмента на все случаи

Глава 3. Анатомия хорошего запроса

За пределами базового prompt engineering

Большинство вводных материалов про запросы к ИИ ограничиваются советом «будьте конкретны» — верным, но недостаточным для практики. Конкретный, но плохо структурированный запрос всё равно даёт посредственный результат. Хороший запрос для рабочей задачи обычно содержит несколько отдельных элементов, каждый из которых решает свою функцию, а не сливается в один расплывчатый абзац текста.

Первый элемент — роль или рамка, в которой должна работать модель: не просто «напиши текст», а « выступи в роли редактора делового издания, который проверяет материал перед публикацией». Роль задаёт модели набор критериев и словарь, которым она будет пользоваться при генерации ответа, значительно сужая диапазон возможных, но нежелательных вариантов ответа.

Второй элемент — контекст задачи: для кого предназначен результат, в какой ситуации он будет использован, какие ограничения уже существуют. Формулировка «напиши email клиенту» и формулировка «напиши email клиенту, который уже трижды переносил встречу и, судя по последнему сообщению, раздражён задержкой проекта» дают принципиально разные по качеству и уместности результаты, потому что вторая версия даёт модели реальную картину ситуации вместо абстрактного шаблона.

Третий элемент — явный формат ожидаемого результата: список, таблица, связный текст определённой длины, структура с подзаголовками. Модель без указания формата выбирает наиболее статистически типичный для похожих запросов вариант — что не всегда совпадает с тем, что реально нужно в конкретной рабочей ситуации, и заставляет тратить время на переформатирование уже после получения содержательно верного, но неудобно оформленного ответа.

Четвёртый элемент — примеры желаемого стиля или тона, там, где это применимо. Модель значительно точнее воспроизводит специфический голос бренда или личный стиль письма, если ей предоставлен образец — один-два реальных примера уже написанного текста, — чем при попытке описать этот стиль абстрактными прилагательными вроде «дружелюбно, но профессионально», которые каждый интерпретирует по-разному, включая саму модель.

Пятый элемент — явные ограничения и то, чего делать не следует: не использовать определённые слова, не превышать заданный объём, не включать конкретные темы. Отрицательные инструкции работают менее надёжно, чем позитивные указания того, что нужно сделать, поэтому по возможности стоит переформулировать ограничение в позитивную форму: вместо «не пиши слишком формально» — «пиши в разговорном, но профессиональном тоне, как будто объясняешь коллеге за чашкой кофе».

Шестой, часто пропускаемый элемент — явный критерий, по которому сам пользователь поймёт, что результат его устраивает. Формулировка задачи с заранее продуманным критерием успеха («результат должен убедить занятого руководителя прочитать письмо целиком за тридцать секунд») даёт модели ориентир для самопроверки ответа перед выдачей, а пользователю — чёткое основание для решения, дорабатывать ответ дальше или принимать его как есть.

Собранные вместе, эти шесть элементов — роль, контекст, формат, стиль, ограничения, критерий успеха — не обязаны присутствовать в каждом запросе целиком: для простых задач достаточно двух-трёх из них. Но для повторяющихся, важных рабочих задач продуманный запрос, включающий большинство этих элементов, экономит несколько раундов уточнений и заметно повышает качество результата с первой попытки — экономия, которая окупает несколько лишних секунд на формулировку запроса многократно за счёт сокращённого числа последующих исправлений.

Сравним на практике: запрос «напиши пост про наш новый продукт» и запрос « выступи в роли автора блога компании, напиши пост на 200 слов для подписчиков email-рассылки, которые уже знакомы с брендом, в разговорном тоне, как в примере ниже, с акцентом на практическую пользу продукта, без превосходных эпитетов вроде революционный или уникальный». Второй запрос почти наверняка даёт результат, пригодный к публикации после лёгкой правки, тогда как первый — лишь отправную точку, требующую нескольких раундов уточнений, чтобы прийти к тому же результату.

Частые ошибки

Ограничиваться советом «быть конкретнее», не структурируя запрос по отдельным функциональным элементам

Формулировать ограничения в отрицательной форме («не делай так») вместо позитивного указания желаемого результата

Не задавать формат ожидаемого ответа и тратить время на переформатирование содержательно верного, но неудобно оформленного результата

Что сделать на этой неделе

Возьмите один регулярный тип запроса и перепишите его, явно указав роль, контекст, формат и критерий успеха

Для следующей задачи, где важен конкретный стиль, приложите к запросу один-два реальных примера текста в желаемом тоне вместо словесного описания стиля

Просмотрите последние отрицательные инструкции («не делай X»), которые вы давали ИИ, и переформулируйте их в позитивную форму

Сформулируйте для одной текущей задачи явный критерий успеха и включите его прямо в текст запроса

Глава 4. Контекст — главный рычаг качества

Как кормить модель нужной информацией

Разница между посредственным и по-настоящему полезным ответом ИИ чаще всего объясняется не мастерством формулировки самого вопроса, а объёмом и качеством контекста, которым располагает модель в момент ответа. Модель без контекста о конкретной компании, продукте или ситуации вынуждена отвечать на основе общих, усреднённых по всему интернету представлений — и результат получается формально грамотным, но обезличенным, будто написанным для абстрактной компании, а не для конкретной.

Практический способ радикально повысить релевантность ответов — заранее подготовить и держать под рукой набор базовых контекстных материалов, к которым можно быстро обращаться: краткое описание компании или проекта, основные факты об аудитории или клиентах, примеры уже одобренных текстов, ключевые термины и внутренний словарь, которым принято пользоваться. Разовая инвестиция времени в подготовку такого набора окупается на каждом последующем запросе, где не приходится заново объяснять базовые вещи о контексте бизнеса.

Отдельного внимания заслуживает работа с уже существующими документами как источником контекста. Вместо пересказа своими словами содержания отчёта, переписки или набора заметок эффективнее сразу предоставить модели сам исходный материал и попросить работать непосредственно с ним — так модель опирается на точные формулировки и цифры оригинала, а не на возможно неточный пересказ, где часть деталей могла быть потеряна или искажена при устной передаче.

Важный практический навык — умение явно указывать модели, какому источнику доверять больше, если предоставленная информация противоречит её общим знаниям. Если во внутреннем документе компании фигурируют конкретные цифры или названия, отличные от общедоступной информации о рынке, стоит явно попросить модель ориентироваться именно на предоставленный документ — иначе есть риск, что модель незаметно подменит специфичные данные более распространёнными, но неверными в конкретном случае формулировками.

Контекст полезен не только в положительном смысле — не менее важно явно очерчивать границы того, что модели не следует предполагать. Например, при составлении коммерческого предложения стоит явно указать, какие детали ещё не согласованы и не должны звучать в тексте как окончательное решение, — иначе модель, стремясь дать связный и уверенный текст, может представить непроверенные предположения как согласованные факты.

Существует разумный предел объёма контекста, за которым отдача начинает снижаться: избыточно длинный, плохо структурированный набор фоновой информации может, наоборот, размыть фокус модели на действительно важных деталях задачи среди второстепенных подробностей. Практический ориентир — предоставлять ровно столько контекста, сколько потребовалось бы для введения в курс дела нового, толкового, но незнакомого с ситуацией коллеги, не больше и не меньше.

Для повторяющихся типов задач полезно один раз выстроить структурированный шаблон контекста — например, фиксированный набор вопросов о продукте, аудитории и цели, — который заполняется перед каждым похожим запросом. Такой шаблон превращает подготовку контекста из творческого, каждый раз заново придумываемого процесса в быструю, почти механическую процедуру, что особенно ценно при частом повторении однотипных задач в течение рабочей недели.

Показательный случай: отдел поддержки клиентов долгое время получал от ИИ ответы, формально корректные, но не учитывающие, что компания предлагает три разных тарифных плана с разными условиями возврата. Проблема решилась не более умной моделью, а про-

стым добавлением в контекст запроса краткой таблицы условий по каждому тарифу — после этого точность ответов выросла заметно, подтверждая, что узкое место было не в возможностях модели, а в объёме предоставленной ей информации о специфике бизнеса.

Частые ошибки

Пересказывать своими словами содержание документов вместо того, чтобы предоставить модели сам исходный материал

Предоставлять избыточно длинный, неструктурированный контекст, размывающий фокус модели на действительно важных деталях

Не указывать явно, каким конкретным данным доверять больше, когда предоставленная информация расходится с общими знаниями модели

Что сделать на этой неделе

Подготовьте базовый набор контекстных материалов о вашей компании, продукте или проекте, который можно быстро прикладывать к запросам

В следующей задаче предоставьте модели исходный документ целиком вместо пересказа его содержания своими словами

Для одного повторяющегося типа задачи создайте фиксированный шаблон вопросов-контекста, который будете заполнять перед каждым похожим запросом

Проверьте один из последних важных ответов ИИ на предмет того, не выдал ли он непроверенное предположение за согласованный факт

Часть 2. Личная продуктивность с ИИ

Глава 5. Почта и коммуникация

Как ИИ снимает рутину переписки, не обезличивая её

Электронная переписка — одна из первых областей, где системное использование ИИ даёт быструю и легко измеримую экономию времени: большинство рабочих писем следуют ограниченному числу повторяющихся сценариев — подтверждение встречи, вежливый отказ, запрос информации, последующее напоминание, — и именно повторяемость делает их идеальным кандидатом для делегирования черновика инструменту.

Наиболее эффективная практика — не просить ИИ написать письмо с нуля по общему описанию ситуации, а предоставлять ему полную цепочку предыдущей переписки и просить составить ответ, учитывающий уже сказанное обеими сторонами. Ответ, составленный без учёта предыдущих сообщений в цепочке, рискует повторить уже обсуждённые детали или, наоборот, упустить важный нюанс, о котором собеседник уже сообщил ранее в переписке.

Отдельная, менее очевидная функция ИИ в работе с почтой — не генерация, а анализ входящих сообщений: обобщение длинной цепочки писем в несколько предложений сути, выделение конкретных вопросов, требующих ответа, среди второстепенных деталей, определение эмоционального тона сообщения для калибровки ответа. Такое использование особенно ценно при работе с длинными, многодневными перепиской по сложным проектам, где легко упустить важную деталь, затерявшуюся среди десятков сообщений.

Ключевая граница, которую стоит соблюдать при делегировании переписки, — важные, эмоционально чувствительные или стратегические сообщения (конфликтные ситуации, серьёзные отказы, деликатные личные темы) заслуживают личного, а не автоматического сгенерированного ответа, даже если черновик от ИИ используется как отправная точка для формулировок. Получатель такого сообщения интуитивно чувствует разницу между текстом, продуманным конкретно для него, и обезличенным, пусть грамматически безупречным шаблоном.

Практичный рабочий процесс для регулярной переписки выглядит так: черновик стандартного ответа генерируется ИИ на основе цепочки переписки и заданного тона, затем человек за тридцать секунд просматривает черновик, вносит правки, специфичные для конкретной ситуации и отношений с получателем, и только после этого отправляет письмо. Такой процесс сохраняет скорость автоматизации, но не жертвует персонализацией там, где она действительно важна для отношений с конкретным человеком.

Для регулярно повторяющихся типов писем — приветственное письмо новому клиенту, стандартный ответ на частый вопрос, шаблон еженедельного отчёта — имеет смысл один раз выстроить с ИИ качественный, проверенный шаблон, а не генерировать каждое подобное письмо заново с нуля. Такой шаблон, доработанный до нужного качества один раз, экономит время при каждом следующем использовании и снижает риск непоследовательности тона между разными письмами одного типа.

Календарное планирование и координация встреч — ещё одна область, где ИИ снимает рутину: составление вежливых предложений времени с учётом часовых поясов, формулировка переноса встречи, обобщение результатов состоявшейся встречи в структурированные заметки с чёткими следующими шагами. Здесь особенно ценна способность ИИ быстро превращать неструктурированные, торопливо записанные заметки во время встречи в чёткий, готовый к рассылке участникам протокол с распределением задач.

Пример из практики: руководитель отдела продаж, получающий по несколько писем в день с похожими вопросами от разных клиентов о сроках поставки, выстроил простой процесс

— вставлять входящее письмо в заранее настроенный шаблон запроса вместе с актуальной таблицей сроков, получать черновик ответа за секунды и тратить оставшееся время на личные штрихи, а не на формулировку с нуля. За месяц это сэкономило заметную часть рабочего дня, ранее уходившую на почти идентичные по содержанию, но каждый раз заново написанные ответы.

Частые ошибки

Генерировать ответ на письмо без предоставления модели полной цепочки предыдущей переписки

Отправлять сгенерированный ИИ черновик деликатного или конфликтного письма без личной доработки под конкретного получателя

Каждый раз заново генерировать типовые письма вместо того, чтобы один раз выстроить проверенный шаблон для повторяющегося сценария

Что сделать на этой неделе

Выберите один тип регулярно повторяющегося письма и составьте вместе с ИИ качественный шаблон для многократного использования

В следующем ответе на цепочку переписки предоставьте модели все предыдущие сообщения, а не только последнее

После следующей рабочей встречи попросите ИИ превратить ваши черновые заметки в структурированный протокол с распределением задач

Определите один тип сообщений (конфликтные, деликатные, стратегические), которые вы никогда не будете отправлять без личной доработки черновика

Глава 6. Работа с документами и знаниями

Саммари, анализ и структурирование информации

Объём текстовой информации, которую современный сотрудник обязан прочитать, чтобы оставаться в курсе своей работы, — отчёты, исследования, длинные цепочки переписки, протоколы встреч — давно превысил разумные временные рамки рабочего дня. Способность ИИ быстро обрабатывать большие объёмы текста и выделять из них суть — одна из самых прямых и измеримых форм экономии времени среди всех возможных применений технологии.

Конец ознакомительного фрагмента.

Текст предоставлен ООО «Литрес».

Прочитайте эту книгу целиком, [купив полную легальную версию](#) на Литрес.

Безопасно оплатить книгу можно банковской картой Visa, MasterCard, Maestro, со счета мобильного телефона, с платежного терминала, в салоне МТС или Связной, через PayPal, WebMoney, Яндекс.Деньги, QIWI Кошелек, бонусными картами или другим удобным Вам способом.