

| Марина Лямкина |

НЕ ЕСТ, ^{← ему не нужно}
НЕ СПИТ, ^{не устаёт →}
НЕ ПЛАЧЕТ. ^{{ без эмоций}

НО ИНОГДА
ВРЁТ *

Как жить и работать
рядом с искусственным
интеллектом и не перестать
думать самому

* данные могут
быть убедительными,
но не всегда —
правдивыми.
Проверяйте.

Марина Лямкина

Не ест, не спит, не плачет. Но иногда врёт

<https://litres.ru/74335328>

SelfPub; 2026

Аннотация

ИИ уже умеет за минуты делать то, на что раньше уходили часы. Но почему свободного времени от этого не становится больше? Почему работа, сделанная за пять минут, не обязательно стоит дешевле? Что происходит с навыком, если хороший результат всё чаще можно получить вместе с системой? И кто отвечает, когда убедительный ответ оказывается неверным?

Марина Лямкина пишет об искусственном интеллекте не как о наборе сервисов и кнопок, а как о новом участнике обычной человеческой жизни и работы. В книге — личный опыт, реальные рабочие наблюдения, исследования и бизнес-ситуации: от усталости и обучения до профессиональной ценности, управления и ответственности.

Это книга не о том, как передать мышление нейросети. Она о том, как остаться думающим человеком рядом с инструментом, который умеет очень многое — и при этом иногда врёт.

Для руководителей, предпринимателей, специалистов и всех, кто уже работает с ИИ или только пытается понять, какое место он должен занимать в их жизни и работе.

Содержание

ПРОЛОГ	5
ГЛАВА 1	11
ГЛАВА 2	19
ГЛАВА 3	32
Конец ознакомительного фрагмента.	37

Марина Лямкина

Не ест, не спит, не плачет. Но иногда врёт

ПРОЛОГ

Да, эту книгу я пишу с ИИ

Я ругаюсь с искусственным интеллектом почти каждый день. Иногда довольно жёстко.

«Очень плохо».

«Нет. Вообще не это».

«Я просила не сокращать».

«Зачем ты опять переписал то, что я не просила переписывать?»

«Верни предыдущую версию».

Бывает, бросаю мышку. Потом поднимаю её, потому что она мне всё-таки нужна, и продолжаю работать.

С человеком я бы так не разговаривала. Во-первых, это невежливо. Во-вторых, человек после пятнадцатого «нет, опять не то» может начать сомневаться в выбранной профессии, во мне и в человечестве в целом. Искусственный интеллект не обижается, не увольняется, не пишет подруге, что

заказчик токсичный, и не сообщает, что сможет вернуться к задаче в четверг после обеда. Через секунду он отвечает: «Понял. Исправляю».

Правда, иногда он ничего не понял. Но это выяснится чуть позже.

Поэтому у меня есть любимое определение искусственного интеллекта: **не ест, не спит, не плачет. Но иногда врёт.**

Примерно так и выглядят наши отношения. Я могу полчаса объяснять ему, почему результат плохой, а через десять минут искренне восхищаться тем, что он сделал.

И именно с ним я пишу эту книгу. Я этого не скрываю. Более того, мне кажется довольно странным писать книгу о жизни рядом с искусственным интеллектом, делая вид, что никакого искусственного интеллекта в процессе её создания не было.

Был. Ещё как.

Сначала я говорю. Иногда буквально наговариваю в телефон всё, что думаю по какой-нибудь теме, без структуры, красивых вступлений и заботы о будущей литературной судьбе текста. Могу повторить одну мысль трижды, передумать на середине предложения, сказать «блин», вспомнить историю, забыть, к чему её рассказывала, потом вернуться и добавить ещё одну.

Получается не глава. **Получается голова человека в текстовом виде.**

Потом начинается работа. Мы ищем в этом потоке главную мысль. Я решаю, что хочу оставить, что кажется важным, в чём сама пока не уверена, где нужен российский контекст, а где моё личное наблюдение необходимо проверить исследованиями.

ИИ собирает структуру. Я читаю и говорю: «Слишком правильно». Он переписывает. Я читаю ещё раз: «Теперь слишком примитивно». Потом прошу вернуть одну фразу, убрать другую, потому что она звучит как текст нейросети, или вообще вернуться на несколько версий назад.

После этого мы ищем исследования, проверяем цифры и смотрим, где мой опыт совпадает с данными, а где я слишком быстро сделала вывод из нескольких случаев. Иногда исследования подтверждают первоначальную мысль, иногда делают её интереснее, а иногда выясняется, что всё устроено совсем не так, как мне казалось.

И это, пожалуй, одна из лучших частей работы. Мне не нужна книга, которая докажет, что я с самого начала была права. Мне нужна книга, после которой мы чуть лучше поймём, что вообще с нами происходит.

Мне не нравится, когда текст, написанный с помощью ИИ, выдают за полностью человеческий. Я довольно быстро чувствую такую прозу — не потому, что обладаю сверхспособностью распознавать нейросети, а потому, что она часто удивительно одинаковая.

Всё очень правильно и очень структурно. Каждая мысль

имеет три подпункта. Каждый подпункт чему-нибудь способствует. В современном мире всё стремительно развивается. Важно отметить важность важного. А в конце обязательно выясняется, что искусственный интеллект — это не просто инструмент, а что-нибудь значительно более великое.

Читаешь десять страниц и понимаешь: текст есть. Человека нет.

Вот этого я не хочу.

ИИ может работать вместе со мной, но я не хочу позволять ему разговаривать вместо меня. Поэтому эта книга стала ещё и экспериментом: можно ли использовать искусственный интеллект очень много и при этом не потерять собственный голос?

Если я надиктовала мысль, потом двадцать раз её переделала, попросила найти исследования, выбросила половину найденного, изменила вывод, вернула свою старую формулировку, снова переписала, проверила источники и в конце подписалась под результатом своим именем — кто написал текст?

Для меня ответ довольно понятный. Авторство — это не количество букв, которые человек лично набрал пальцами. Это мысль, опыт, выбор, критерии, позиция и ответственность за то, что в итоге осталось на странице.

ИИ может предложить двадцать вариантов. Девятнадцать я выброшу. Иногда выброшу все двадцать, а иногда после них придумаю двадцать первый.

Вот так мы и работаем.

Почему я вообще решила, что могу написать книгу об искусственном интеллекте?

Я не создаю языковые модели, не занимаюсь фундаментальными исследованиями машинного обучения и не собираюсь объяснять устройство трансформеров человеку, который совершенно спокойно жил без этого знания до сегодняшнего дня. И точно не знаю об ИИ всё. По-моему, сейчас вообще довольно рискованно заявлять, что знаешь об ИИ всё: пока договоришь эту фразу, может выйти новая модель.

Но несколько лет я очень много наблюдаю за другим местом — за точкой, в которой искусственный интеллект сталкивается с обычным человеком.

С руководителем, предпринимателем, преподавателем, сотрудником компании. С человеком, который впервые открыл нейросеть и написал: «Привет. Что ты умеешь?» С тем, кто попробовал один раз, получил ерунду и теперь уверен, что искусственный интеллект бесполезен. И с человеком, который уже плохо помнит, как раньше работал без него.

Чем больше я на всё это смотрю, тем меньше меня интересует вопрос «на что способен искусственный интеллект?». Возможности всё равно изменятся быстрее, чем эта книга доедет до читателя.

Мне интереснее, что происходит с нами, когда рядом появляется инструмент, способный за минуты сделать то, на что раньше уходили часы, дни, деньги и участие других лю-

дей. Как меняется работа, цена профессионализма, обучение, доверие к собственным знаниям, отношения между руководителем и сотрудником — и почему человек, наконец получивший возможность делать быстрее, далеко не всегда получает больше свободного времени.

Модели и сервисы будут меняться быстрее, чем эта книга доедет до читателя. Пока вы её читаете, какие-то функции уже могли появиться, а то, чем сегодня все восхищаются, через несколько месяцев может стать совершенно обычным.

Поэтому я не хочу писать книгу о расположении кнопок. Кнопки проиграют времени.

А вот способность понять, чего ты хочешь от результата, скорее всего, проживёт дольше. Как и вопрос, кто отвечает, если ИИ ошибся, а решение принял ты. Как и умение отличить красивый ответ от хорошего и иногда вовремя спросить: «Подождите. А мы вообще зачем это делаем?»

Поэтому здесь будут модели, исследования, цифры и конкретные примеры, но книга всё-таки не про них. Она про нас — людей, которые внезапно получили очень мощного нового работника.

Он не ест, не спит, не плачет.

И, к сожалению, не снимает с нас ответственность.

ГЛАВА 1

Я обещала никогда не работать по специальности

Есть вещи, которые Вселенная, видимо, записывает, поэтому нужно осторожнее с обещаниями.

Когда я закончила Алтайский государственный технический университет, я была совершенно уверена в одном: никогда не буду работать по специальности.

Учиться мне было тяжело и, что хуже, мне было неинтересно. «Информационно-измерительная техника и технологии» звучит намного солиднее, если рассказывать об этом спустя много лет. В процессе обучения моя главная мысль чаще звучала примерно так: «Господи, когда это закончится?»

Закончила. Через множество препятствий. Выдохнула и решила, что теперь моя жизнь наконец-то пойдёт в совершенно другую сторону.

И она пошла. Просто сделала очень большой круг.

Мой диплом был связан с разработкой интеллектуальной системы для решения финансовой задачи. Я тогда уже работала в банке кредитным специалистом, и идея заключалась в том, чтобы система могла оценивать кредитоспособность человека на основании известных параметров.

Сегодня мне очень хочется красиво сказать: «Представляете, уже в 2012 году я занималась искусственным интеллектом». Но здесь мой сегодняшний внутренний фактчекер поднимает руку и говорит: «Марина, спокойно».

Поэтому скажу честнее. Тогда я работала над интеллектуальной системой. Насколько технически корректно относить конкретно ту дипломную работу к тому, что сегодня мы называем искусственным интеллектом, нужно отдельно смотреть по её архитектуре и методам.

В 2012 году я точно не сидела с ощущением: «Вот он, мой первый шаг в великое будущее ИИ». Я хотела защитить диплом. Это значительно более правдивое описание происходящего.

Но сама ирония мне сейчас нравится: человек несколько лет учится технической специальности, обещает больше никогда к ней не возвращаться, а потом проходит десять лет — и значительная часть его работы снова оказывается связана с технологиями.

Если бы мне тогда это рассказали, я бы, скорее всего, расстроилась.

Между дипломом и сегодняшней работой с искусственным интеллектом, вообще-то, поместилась довольно большая профессиональная жизнь. Я больше семнадцати лет работаю в продажах и управлении, восемь лет руководила филиалом федеральной IT-компании, потом к этому добавились предпринимательство, бизнес-наставничество и обуче-

ние взрослых.

То есть к моменту, когда генеративный ИИ начал все-рьёз входить в мою работу, я пришла к нему далеко не с чистым листом. У меня уже были сотрудники, клиенты, планы продаж, управленческие задачи, процессы, отчёты, обучение людей, собственный бизнес и приличное количество ситуаций, в которых что-нибудь нужно было понять, придумать, исправить или объяснить другому человеку.

Сейчас мне кажется, что именно поэтому ИИ так быстро занял большое место в моей работе. Он попал не в одну профессию, а сразу в несколько слоёв опыта, которые к тому моменту уже успели накопиться.

Следующая встреча с ИИ была гораздо менее интеллектуальной. Примерно в 2022 году он для меня выглядел как развлечение: приложение, в котором можно посмотреть, как я выглядела бы мужчиной, человеком другой национальности или ещё кем-нибудь, кем совершенно не собиралась становиться.

Это было весело. Иногда страшновато, иногда похоже, чаще не очень.

Я не думала: «Перед нами технология, которая радикально изменит мой способ работы». Я думала: «Прикольно».

И мне кажется, очень многие познакомились с ИИ примерно так. Не через стратегическую сессию, доклад McKinsey или заседание совета директоров. Через баловство.

А потом игрушка внезапно начала работать.

Для меня перелом произошёл в 2023 году, когда я стала осознанно использовать ИИ для текста. Вот там был настоящий вау-эффект.

Я довольно быстро думаю. Гораздо хуже у меня иногда получается аккуратно упаковывать всё, что я уже успела подумать. В голове мысль понятная, внутри неё уже пять связей, есть пример, вывод, ещё один вывод и раздражение по поводу вывода. Нужно только каким-то образом превратить всё это в текст, который поймёт другой человек.

И вдруг появляется инструмент, которому можно отдать мысль почти в сыром виде, а через несколько секунд получить структуру. Не всегда хорошую, не всегда мою, но уже что-то, с чем можно работать.

Для меня это было почти физическим ощущением ускорения — как будто между мыслью и результатом кто-то убрал несколько лестничных пролётов.

Потом произошло то, что с хорошими технологиями происходит довольно быстро: я привыкла.

Это вообще удивительная человеческая способность. Сегодня происходит чудо, а через месяц мы раздражены, что чудо работает медленно. То, что в 2023 году вызывало восторг, сегодня я могу отвергнуть словами «слишком примитивно». То, что несколько лет назад казалось фантастическим качеством изображения, теперь легко выглядит устаревшим.

В работе с ИИ это ощущается особенно странно: не успеваешь как следует привыкнуть к новой планке, как она становится нормой. А потом появляется следующая.

У меня есть ещё одна личная особенность, которая в сочетании с ИИ дала довольно опасный результат. Я много лет живу с установкой: «Хочешь сделать хорошо — сделай сам».

Я понимаю, что управленческие консультанты сейчас где-то почувствовали боль. Да, я знаю: так нельзя. Нужно делегировать, строить процессы, доверять людям, развивать команду.

Всё это прекрасно звучит.

Но моя первая реакция на незнакомую задачу всё равно обычно не «я этого не умею», а скорее: «Так, а что нужно, чтобы разобраться?»

До появления генеративного ИИ эта особенность и так регулярно приводила меня в чужие профессиональные территории. С ним территория резко расширилась. Я стала использовать ИИ для текстов, исследований, аналитики, презентаций, визуалов, методических материалов, простых цифровых решений — и постепенно обнаружила, что барьер между «я этого никогда не делала» и «давайте хотя бы попробуем» стал гораздо ниже.

Подробно о последствиях этой новой самостоятельности мы ещё поговорим дальше. Здесь для меня важнее другое: вместе с ощущением новых возможностей довольно быстро пришло понимание, что доступ к инструменту не делает ме-

ня специалистом во всех областях, куда я теперь могу пойти.

И однажды я на этом вполне наглядно обожглась.

Мне понадобилось подготовить юридические документы для собственной организации. Я пошла привычным путём: информация рядом, ИИ рядом — что может пойти не так?

Выяснилось — форма.

Я исходно пошла не в ту сторону и, поскольку сама не обладала достаточной юридической компетенцией, не заметила ошибку на входе. В результате мы могли сколько угодно прекрасно работать внутри неправильной постановки задачи.

Цена ошибки в деньгах была нулевая. Цена во времени — вполне реальная.

Вывод оказался неприятно простым: ИИ не компенсирует компетенцию, которой у тебя нет, если без этой компетенции ты не способен заметить, что изначально задал неправильную задачу.

Иногда он помогает новичку очень быстро продвинуться. Иногда позволяет сделать то, что раньше было недоступно. А иногда очень эффективно помогает идти не туда.

Наверное, именно поэтому сейчас я всё меньше верю в разговоры о том, что главное — научиться писать хорошие промпты.

Промпт важен. Но хороший запрос не спасёт, если вы не понимаете, чего хотите, какие данные имеют значение, как проверить результат и где заканчивается ваша собственная

компетенция.

Можно очень профессионально попросить ИИ построить лестницу к стене. Он построит — быстро, аккуратно, может даже с таблицей этапов. Проблема обнаружится, когда выяснится, что дверь была с другой стороны.

Сегодня я иду к ИИ практически с любым вопросом: рабочим, исследовательским, творческим, бытовым. Могу обсуждать бизнес-процесс, через полчаса — структуру тренинга, потом текст, таблицу, цифровой продукт, а вечером — выбор школы для ребёнка.

Это уже не отдельный инструмент в моей работе. Скорее слой, который лежит поверх огромного количества задач.

Поэтому мне всё меньше нравится вопрос: «Для чего вы используете ИИ?» Слишком долго отвечать.

Гораздо интереснее другой: **для чего я его не использую?**

И ответы пока есть.

Я не хочу отдавать ему окончательный выбор там, где решение имеет серьёзные последствия. Можно спросить: «Какие у меня варианты?», «Какие риски я не учла?», «Поспорь со мной», «Почему этот вариант может быть плохим?»

Даже вопрос «Разводиться мне или нет?» может стать началом разговора с собой. Но решение всё равно придётся принимать человеку.

ИИ может расширить поле зрения. **Руль я пока оставляю себе.**

За несколько лет работы с ним я точно стала быстрее и смелее в незнакомых задачах. Стала требовательнее к результату и гораздо спокойнее отношусь к первой попытке: если она дешёвая, можно просто попробовать и посмотреть, что получится.

Но вместе с этим появилась более интересная вещь. Чем больше ИИ умеет делать для меня, тем чаще приходится задумываться уже не о его возможностях, а о собственных границах.

Где я действительно научилась новому, а где просто научилась получать результат с хорошим помощником? Где можно довериться, а где мне не хватает знаний даже для нормальной проверки? Что стоит делать самой, потому что теперь технически могу, а что всё-таки разумнее оставить человеку, который действительно понимает предмет?

В 2023 году меня больше всего интересовало: «Что он ещё умеет?» Сейчас значительно чаще возникает другой вопрос: «Что постоянная жизнь рядом с ним начинает менять во мне?» Но мой путь к ИИ оказался довольно быстрым. У многих знакомство начинается с другой фразы: «Мне это не нужно». И она заслуживает более внимательного ответа, чем «вы просто сопротивляетесь новому».

ГЛАВА 2

Мне это не нужно

Фраза «мне искусственный интеллект не нужен» встречается мне регулярно.

Иногда её произносит человек, который действительно пока не понимает, зачем ему ИИ. Иногда тот, кто уже попробовал, получил какой-то совершенно бессмысленный ответ и решил, что второй попытки технология не заслуживает. Кто-то опасается за данные, кто-то уверен, что его профессия слишком специфична, а кто-то просто не хочет осваивать ещё одну новую штуку, потому что новых штук за последние годы и без того было достаточно.

И чем больше я работаю с разными людьми, тем меньше мне хочется объединять их всех словом «сопротивление».

Очень удобное слово. Особенно если ты сам уже активно пользуешься ИИ. Ты сидишь такой прогрессивный, освоивший будущее, а вокруг почему-то люди сопротивляются неизбежному. Дальше довольно легко объяснить всё ленью, страхом нового, возрастом или знаменитым «мы всегда так делали».

Проблема только в том, что реальность обычно сложнее. Иногда человек не видит пользы. Иногда ему показали плохой инструмент. Иногда руководство действительно

предлагает автоматизировать то, что сначала неплохо было бы просто привести в порядок. Иногда человек уже видел уверенную ошибку ИИ и совершенно справедливо спрашивает, кто будет отвечать, если такая ошибка попадёт в реальную работу.

А иногда он вообще не сопротивляется. Уже пользуется ИИ, просто никому об этом особо не рассказывает.

Последнее особенно интересно, если посмотреть на российские данные.

В 2025 году Росстат впервые включил в Обследование рабочей силы отдельный блок вопросов о навыках работы с искусственным интеллектом. Масштаб исследования хороший: 314 тысяч респондентов, результаты репрезентативны для 46,7 миллиона занятых в российской экономике. По этим данным, 37,5% работающих россиян сообщили, что уже обладают навыками работы с ИИ. При этом только 4,9% считают такие навыки необходимыми для своей нынешней работы. [2]

Мне очень нравится эта статистика, потому что она довольно сильно портит привычный корпоративный сюжет.

По этому сюжету сначала компания понимает, что наступила эпоха искусственного интеллекта. Потом руководитель принимает стратегическое решение, HR организует обучение, сотрудников торжественно ведут в будущее, а они по дороге сопротивляются.

В реальности сотрудник вполне может уже несколько ме-

сяцев пользоваться нейросетью, пока компания ещё проводит совещание на тему «нужно ли нам внедрять искусственный интеллект».

Исследователи ВШЭ отдельно отмечают, что у значительной части занятых фактический уровень владения ИИ превышает требования рабочего места. Есть и совершенно прекрасная для российской реальности оговорка: использование ИИ нередко происходит в «серой зоне» — например, через личные устройства — и никак не отражено в должностных инструкциях.

То есть ИИ у вас, возможно, уже внедрён.

Просто вы его пока официально не внедряли.

Здесь, правда, начинается статистическая комедия.

На вопрос «Сколько российских компаний уже используют ИИ?» вполне можно получить два ответа: **86% и 4,8%**.

И оба будут взяты не из случайного телеграм-канала.

В исследовании НАФИ и компании «Технологии Доверия», проведённом в октябре 2025 года среди 1001 руководителя бизнеса, 86% респондентов сообщили, что в их организациях уже применяются инструменты искусственного интеллекта. [5]

По оценкам ИСИЭЗ НИУ ВШЭ, технологии ИИ используют около 4,8% обследованных крупных и средних организаций; субъекты малого предпринимательства в эту статистику не входят. Среди организаций с численностью более 500 работников показатель достигает 14,9%, а в группе до

100 работников — 4,1%. [3]

Можно выбрать ту цифру, которая больше нравится, и написать красивый заголовок:

«86% российского бизнеса уже с ИИ!»

Или:

«95% компаний ещё не внедрили искусственный интеллект!»

Обе версии будут выглядеть убедительно. Именно поэтому с цифрами иногда полезно работать чуть дольше пяти минут.

Исследования измеряют разные вещи, используют разные выборки и разные определения. В одном случае руководитель отвечает, применяются ли в компании ИИ-инструменты. Это может быть доступный генеративный сервис для отдельных задач. В другом применяется более строгая статистическая классификация технологий ИИ на уровне организации.

Для меня здесь важна даже не конкретная цифра, а разница между двумя фразами:

«У нас люди пользуются нейросетями».

И:

«У нас ИИ встроен в рабочий процесс».

Это вообще не одно и то же.

Если сотрудник попросил нейросеть переписать письмо клиенту, ИИ в компании уже появился. Но если завтра этот сотрудник уволится, вместе с ним вполне может исчезнуть и

всё «внедрение».

Если каждый работает через свой личный аккаунт, по своим правилам, никто не понимает, какие данные туда можно загружать, результаты нигде не фиксируются, а качество никто не проверяет, это тоже использование ИИ.

Только я бы пока не спешила называть это цифровой трансформацией.

Я сама довольно долго повторяла фразу: **«ИИ заменит тех, кто игнорирует ИИ»**.

Она мне до сих пор нравится. Хорошая, короткая, бодрая. На лекции работает прекрасно.

С книгой сложнее, потому что здесь можно не торопиться.

Если воспринимать эту фразу буквально, я не могу подтвердить, что всё будет именно так. Российская статистика пока точно не показывает картину, в которой компании массово увольняют всех, кто не освоил нейросеть. Среди крупных и средних организаций, уже использующих технологии ИИ, 15% сообщили о снижении численности работников в связи с внедрением, а 63% сказали, что влияния на численность персонала не увидели. [4]

Это не означает, что риска замещения нет, и тем более не позволяет предсказать, что произойдёт через несколько лет. Но формулировка «научись пользоваться ИИ или потеряешь работу» сейчас выглядит слишком простой.

Мне гораздо интереснее другой вопрос: **что будет с человеком, который продолжает выполнять за два дня**

то, что его коллега научился делать за несколько часов?

Не обязательно увольнение. Возможно, изменится ценность отдельных навыков, ожидания работодателя, количество задач, которые способен вести один человек, или сама должность. А возможно, именно в этой работе ИИ окажется почти бесполезен.

Угроза не обязательно выглядит как робот, который утром приходит и садится на ваш стул. Иногда она выглядит как коллега с тем же рабочим днём, но другим набором возможностей.

Вот эту версию я бы воспринимала серьёзнее.

При этом российские данные вообще плохо поддерживают красивую историю о массовой панике перед ИИ.

В июле 2026 года ВЦИОМ получил довольно спокойную картину: 67% россиян считают, что искусственный интеллект стоит применять только в некоторых сферах, 19% поддерживают его широкое использование и лишь 9% выступают против применения ИИ в принципе. Одновременно 77% считают необходимой маркировку продуктов, изображений и решений, созданных с помощью ИИ. [6]

То есть позиция большинства выглядит не как «запретить всё», а скорее:

«Пользоваться можно. Но давайте понимать, где, как и кто за это отвечает».

Мне это очень близко, потому что чем больше я сама ра-

ботаю с ИИ, тем меньше мне нравится идея «давайте использовать его везде». Везде — плохой критерий. Если процесс плохо работает без ИИ, добавление технологии само по себе его не исправит: прежде чем автоматизировать, я хочу понять, что именно происходит в процессе и зачем. ИИ прекрасно умеет ускорять работу, в том числе ту, которую вообще не нужно было делать.

Есть ещё один тип скепсиса, который я считаю совершенно рациональным.

Человек уже попробовал нейросеть. Написал вопрос. Получил уверенный ответ. Ответ оказался неправильным.

На этом эксперимент закончился.

Я понимаю такую реакцию гораздо лучше, чем позицию человека, который не открывал ни одной нейросети, но точно знает, что всё это ерунда.

Проблема только в том, что плохой первый опыт очень легко превращается в вывод о технологии целиком. Это примерно как однажды заказать невкусную пиццу и решить, что итальянская кухня не получилась.

Хотя противоположная крайность мне нравится ещё меньше: получить один хороший ответ и немедленно решить, что теперь можно доверять всему.

Иногда скепсис исчезает не после лекции о будущем рынка труда, а после одной нормально выбранной задачи.

Я видела людей, которые довольно равнодушно слушали рассказ о возможностях ИИ, пока мы не переходили к их

собственной работе.

Не «ИИ умеет анализировать данные», а «давайте возьмём вашу таблицу».

Не «нейросеть может помочь с коммуникацией», а «покажите письмо, на которое вы вчера потратили сорок минут».

Не «с помощью ИИ можно сделать цифровой продукт», а «давайте попробуем собрать ваш первый вариант».

В этот момент технология перестаёт быть темой новостей и становится способом решить раздражающую задачу.

Мне вообще кажется, что один из худших способов внедрить ИИ в компании — объявить:

«С понедельника мы используем искусственный интеллект».

Где? Зачем? Для чего? Как поймём, что стало лучше? Какие данные можно туда загружать? Кто проверяет результат?

Если человек не знает ответов, вполне нормально, что его не охватывает технологический восторг.

Гораздо лучше работает другой путь: вот конкретная задача, вот как она выполняется сейчас, вот что мы попробовали изменить, вот результат, вот где ИИ ошибся, а вот что всё равно оставили человеку.

После этого скепсис обычно становится гораздо предметнее. Человек уже не спорит с «искусственным интеллектом вообще». Он говорит:

«Вот здесь полезно. А вот сюда я бы его не пускал».

Мне кажется, это и есть нормальное взросление отноше-

ния к технологии.

Возраст здесь тоже оказался намного интереснее стереотипа.

По российским данным молодые действительно чаще обладают ИИ-навыками, особенно среднего и продвинутого уровня. Среди работников с высшим образованием они тоже встречаются заметно чаще. Но базовый уровень освоения распределён гораздо ровнее: ВШЭ отдельно отмечает, что для начала использования популярных генеративных сервисов высокий уровень формального образования сам по себе не является критическим условием. [1]

У меня есть очень любимый пример из собственной практики. Меня пригласили провести занятие для группы, где средний возраст участников был примерно 65–70 лет.

Я до сих пор помню людей, которые смотрели в телефон и доставали лупу, чтобы разглядеть текст на экране. И при этом продолжали формулировать запрос, получать ответ, пробовать ещё раз, задавать вопросы.

Для меня такие случаи важны не потому, что они доказывают отсутствие возрастных различий. Различия есть, статистика их показывает.

Они доказывают другое: возраст не даёт права заранее решить за конкретного человека, что он не освоит технологию.

Иногда самый технически уверенный участник пробует один раз, получает неинтересный результат и теряет интерес. А человек с лупой продолжает.

Есть ещё замечательная категория людей: они против ИИ, но пользуются им настолько хорошо, что способны с его помощью аргументировать собственную позицию.

У меня был случай, когда сотрудники подготовили с помощью ИИ ответ о том, почему ИИ — это плохо.

По-моему, это практически идеальная иллюстрация происходящего.

Инструмент вообще не обязан разделять ваш восторг по поводу собственного существования.

Можно попросить его написать десять причин внедрить ИИ, а через минуту — двадцать причин этого не делать. Можно попросить доказать, что профессия исчезнет, а потом — что она станет только востребованнее.

Если поставить задачу достаточно уверенно, можно получить убедительный текст почти под любую заранее выбранную позицию.

И это ещё одна причина, почему я не считаю умение **получить хороший ответ от нейросети** главным новым навыком.

Получить ответ становится всё проще. Сложнее понять, почему вы задали именно этот вопрос и что собираетесь делать с ответом дальше.

Когда я смотрю на российские исследования, мне вообще кажется, что главный разрыв сейчас проходит не между теми, кто «за ИИ», и теми, кто «против».

Он проходит между разными уровнями использования.

Можно открыть нейросеть раз в месяц, попросить поздравление с днём рождения и совершенно честно сказать: «Я использую ИИ».

Можно каждый день делать с ней тексты и картинки.

Можно встроить её в конкретный рабочий процесс.

А можно изменить сам процесс, потому что часть операций теперь выполняется иначе.

Во многих опросах всё это помещается в один ответ: «Да, используем».

ВШЭ как раз формулирует один из ключевых вызовов российского рынка как переход от массового использования доступных ИИ-сервисов к профессиональной интеграции ИИ в рабочие процессы.

Для меня это очень точное различие.

Мы довольно быстро научились спрашивать.

Теперь придётся учиться перестраивать работу.

А это уже гораздо сложнее, потому что здесь недостаточно купить подписку, провести мастер-класс и раздать сотрудникам список из пятидесяти промптов. Нужно решить, что именно мы хотим изменить, какие данные используем, где границы допустимого, кто проверяет результат и кто отвечает за ошибку.

И главное — зачем человеку вообще менять привычный способ работы.

Без ответа на последний вопрос всё остальное иногда пре-

вращается в очень дорогой корпоративный кружок по интересам.

Поэтому я теперь гораздо осторожнее отношусь к фразе «люди сопротивляются ИИ».

Иногда сопротивляются. Иногда боятся. Иногда просто не хотят разбираться. Но иногда задают совершенно правильные вопросы, понимают процесс лучше человека, который решил его автоматизировать, уже используют ИИ быстрее собственной компании или вполне рационально не видят в нём пользы для конкретной задачи.

Мне вообще не нужно, чтобы люди любили искусственный интеллект.

Я сама периодически бросаю из-за него мышку.

Гораздо полезнее спросить:

«Что конкретно здесь станет лучше, если мы добавим ИИ?»

Если ответа нет — возможно, пока не надо ничего добавлять.

Если ответ звучит как «потому что сейчас все внедряют» — тем более.

А если можно показать, что раньше человек час собирал информацию, а теперь получает за пять минут черновик и тратит оставшееся время на решение, что раньше здесь терялись данные, а теперь не теряются, что вместо ручной проверки ста строк специалист смотрит только десять спорных, — разговор становится совсем другим.

Не про любовь к технологиям, а про работу.

И, возможно, именно поэтому ИИ уже довольно глубоко проник в нашу жизнь, пока компании всё ещё спорят о том, внедрён он у них или нет.

Сотрудник просто открыл нейросеть и начал пользоваться.

Начальство узнает чуть позже.

ГЛАВА 3

Почему я устала, если всё сделал он?

Когда человек всё-таки находит задачу, где ИИ действительно полезен, возникает следующий сюрприз: сэкономленное время далеко не всегда становится свободным. У начинающих пользователей я довольно часто слышу одну и ту же мечту:

«Вот сейчас я отдам ему всю рутину и наконец-то освобожу время».

Хороший план.

У меня не сработал.

Я совершенно точно стала делать многие вещи быстрее. То, на что раньше мог уйти день, иногда занимает час. Где-то вместо нескольких часов получается двадцать минут. Какие-то задачи вообще появились в моей работе только потому, что теперь я могу достаточно быстро сделать их сама.

Свободного времени почему-то больше не стало.

Наоборот, я стала делать ещё больше.

Если раньше идея могла спокойно умереть на этапе «для этого нужно найти человека, объяснить, договориться, дождаться результата», сейчас появляется гораздо более опасная мысль:

«А давай попробуем прямо сейчас».

И мы пробуем.

Потом ещё вот это. Раз уж получилось, можно проверить следующую идею, провести дополнительное исследование, собрать презентацию или наконец сделать то, что несколько месяцев лежало в папке «когда-нибудь».

В какой-то момент хочется спросить: а где, собственно, та часть истории, в которой искусственный интеллект освобождает мне время?

Наверное, он его освобождает.

Просто я тут же его занимаю.

И это, как мне кажется, важная разница: **экономия времени и уменьшение нагрузки — совсем не одно и то же.**

У меня был случай на обучении, который очень хорошо это показал.

Целый отдел примерно два дня работал над одной задачей. Потом мы с этими же людьми разобрали её с помощью ИИ и получили результат примерно за пять минут.

Это не научный эксперимент и не доказательство того, что искусственный интеллект ускоряет любую работу в сотни раз. Там была конкретная задача, конкретные люди и конкретная ситуация.

Но самое интересное произошло не с задачей.

С людьми.

Часть действительно обрадовалась: представляете, сколь-

ко времени теперь можно не тратить на эту работу.

У другой части возникла совершенно иная мысль:

«Так. Если теперь эта работа занимает пять минут, что будет с оставшимся временем? Мне же просто дадут ещё задач?»

С точки зрения бизнеса это даже выглядит логично. Если производительность выросла, освободившийся ресурс можно направить на другую работу.

С точки зрения сотрудника арифметика выглядит иначе.

Вчера у меня была одна задача. Сегодня я научился делать её в четыре раза быстрее. Завтра у меня будет четыре задачи.

Где моя выгода?

И вот здесь разговор про чудесное высвобождение времени становится намного сложнее.

В мировых исследованиях уже видно, что ИИ действительно способен снижать нагрузку, но происходит это далеко не автоматически.

В большом исследовании OECD малого и среднего бизнеса в семи странах примерно треть компаний, использующих генеративный ИИ, сообщила, что нагрузка сотрудников снизилась. Но 11,8% сказали обратное — нагрузка выросла. Причём результат заметно различался по отраслям: например, среди опрошенных компаний в сельском хозяйстве рост нагрузки отметили 38,5%, в образовании — 29,6%. OECD отдельно подчёркивает, что само исследование не позволяет установить причины этих различий. [11]

Мне нравится эта статистика хотя бы потому, что она ломает простую формулу «внедрили ИИ — стало легче»: у кого-то стало легче, у кого-то — нет.

А кто-то, возможно, просто получил очень эффективный инструмент для того, чтобы работать ещё больше.

Интересно и то, что у индивидуальных предпринимателей и компаний из одного человека снижение собственной нагрузки в этом исследовании отмечалось чаще, чем у более крупных компаний. Исследователи осторожно предполагают, что одной из причин может быть больший контроль такого человека над собственным рабочим временем.

Вот это мне особенно понятно.

Потому что существует огромная разница между:

«ИИ сэкономил мне два часа, и я решила, что делать с этими двумя часами».

И:

«ИИ сэкономил компании два часа моего времени, и компания решила, что я буду делать с этими двумя часами».

Технология одна.

Ощущение от неё совершенно разное.

В российской практике этот вопрос исследован хуже, но уже есть интересные сигналы.

Например, в исследовании ИСИЭЗ НИУ ВШЭ, основанном на 30 глубинных интервью с ведущими российскими учёными, главным положительным эффектом называли экономию времени за счёт передачи ИИ технических и вспомо-

гательных задач. При этом среди гипотетических рисков респонденты отмечали рост потока низкокачественных публикаций, который способен увеличить нагрузку на научное сообщество из-за необходимости всё это читать и рецензировать. [7]

Один человек сэкономил время, быстро создав текст.

Другой потратил время, проверяя то, что первый быстро создал.

В масштабе системы экономия уже выглядит не настолько очевидной.

Мне кажется, с ИИ мы ещё не раз увидим такой перенос работы. Он убирает один вид труда и создаёт другой. Меньше времени нужно на первый черновик — больше внимания может понадобиться на проверку. Не нужно вручную искать двадцать вариантов — зато нужно выбрать один из двадцати предложенных. Можно быстро получить анализ — но всё равно придётся понять, можно ли на его основании принимать решение.

ИИ умеет очень быстро производить информацию.

Человеческая способность её переваривать почему-то не увеличилась с той же скоростью.

И это начинает чувствоваться.

Я долго думала, почему иногда после нескольких часов плотной работы с ИИ устаю совершенно нормально, хотя значительную часть работы вроде бы выполнял не я.

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.