

АРТЁМ ДЕМИДОВ

ЧТО ЗНАЧИТ БЫТЬ ЧЕЛОВЕКОМ, КОГДА МАШИНЫ УМНЕЕ

КНИГА
ДЛЯ ТЕХ,
КТО ХОЧЕТ
ОСТАТЬСЯ
ЧЕЛОВЕКОМ
В ЭПОХУ
ИСКУССТВЕННОГО
ИНТЕЛЛЕКТА



ПОНИМАЙ
КАК УСТРОЕН
ИИ



ИСПОЛЬЗУЙ
НЕЙРОСЕТИ
НА ПРАКТИКЕ



ЗАЩИЩАЙ
СЕБЯ, ДАННЫЕ
И СВОЙ ВЫБОР



СОХРАНЯЙ
ЧЕЛОВЕЧНОСТЬ,
А НЕ АЛГОРИТМ

ТЕХНОЛОГИИ МЕНЯЮТ МИР.
ТВОЙ ВЫБОР МЕНЯЕТ ТЕБЯ

Артём Демидов

**Что значит быть человеком,
когда машины умнее**

«Автор»

2026

Демидов А.

Что значит быть человеком, когда машины умнее / А. Демидов —
«Автор», 2026

Ваш телефон знает о вас больше, чем вы сами. Нейросеть за пятнадцать минут пишет текст, за который копирайтер брал пятьдесят тысяч. А компания, потратившая миллион на «внедрение ИИ», год спустя не может объяснить, что изменилось. Артём Демидов — основатель студии AI-автоматизации, автор двух книг про искусственный интеллект. Он внедрял ИИ в компаниях с оборотом от миллиона до миллиарда и видел, как это работает и как проваливается. Это книга о том, что остаётся человеку, когда машина умеет почти всё. Без прогнозов на тридцать лет и обещаний «десяти шагов к успеху». Только реальные истории: маркетолог, освоившая нейросети в декрете и возглавившая отдел. Дизайнер, чуть не ушедшая из профессии. Кадровик, чей алгоритм тихо отсеивал лучших кандидатов. Внутри — как объяснять задачи машине, какие профессии меняются сейчас, что такое ИИ-агенты и почему у большинства они не взлетают, как защитить данные и не разучиться думать. Актуально на август 2026 года. С разбором нового закона об ИИ.

© Демидов А., 2026

© Автор, 2026

Содержание

Вступление. Ваш телефон умнее вас. Привыкайте	5
ЧАСТЬ I. ОСНОВЫ	10
Глава 1. Как эта штука вообще работает	10
Глава 2. Люди ломают роботов чаще, чем роботы — людей	14
Глава 3. ИИ не заберёт вашу работу. Её заберёт тот, кто умеет с ним работать	20
Глава 4. Почему нейросеть пишет лучше копирайтера за пятьдесят тысяч	26
Конец ознакомительного фрагмента.	29

Что значит быть человеком, когда машины умнее

Вступление. Ваш телефон умнее вас. Привыкайте

Книга для тех, кто хочет остаться человеком в эпоху искусственного интеллекта

Артём Демидов

Актуально на август 2026 года

У меня для вас плохие новости.

Ваш телефон знает о вас больше, чем вы сами. Он помнит, где вы были в прошлый четверг. В каком магазине задержались на сорок минут. Кому звонили после полуночи. Колонка на кухне помнит, что вы покупали полгода назад, и подсовывает «похожие товары» каждый раз, когда вы открываете маркетплейс. А нейросеть за пятнадцать минут пишет текст, за который копирайтер брал пятьдесят тысяч и неделю времени.

Хорошие новости тоже есть. Мир не заканчивается. Просто меняются правила.

Правила эти уже не те, что были в 2022-м, когда ChatGPT только появился. Не те, что были в 2024-м, когда нейросети научились писать код. Даже не те, что были прошлой весной.

Всё меняется так быстро, что если вы сели читать эту книгу с мыслью «наконец-то разберусь с ИИ раз и навсегда» — приготовьтесь. Часть примеров устареет за год. Это нормально. Названия моделей меняются каждый квартал. Принципы — нет. Про принципы и книга.

Кто я и почему пишу уже третью книгу

Меня зовут Артём Демидов. Я основатель AINEX LAB — студии, которая собирает для бизнеса AI-автоматизацию: ботов, системы поиска по базам знаний, агентов, которые работают сами.

Если по-простому: мы делаем так, чтобы один сотрудник работал за пятерых. А ещё пятеро — не работали вовсе, потому что их работу делает система.

За последние годы я внедрил ИИ в компаниях с оборотом от миллиона до миллиарда. Видел, как одна правильная автоматизация превращает хаос в порядок за месяц. Видел и обратное: «внедрение революционных технологий» превращает нормальный бизнес в дорогой цирк, где вместо клоунов — роботы.

Это моя третья книга про ИИ.

Первая — «Хватит ныть. Пора зарабатывать на ИИ» — была про то, почему начинать надо было ещё вчера. Вторая — «Тебя заменят. Смирись или адаптируйся» — жёсткое предупреждение тем, кто считал свою профессию особенной. Обе есть на Литресе. Порядок чтения не важен, начинать можно с любой.

Эта книга — другая.

Она не про то, как заработать. Не про то, почему надо адаптироваться. Она про то, что остаётся от человека, когда машина умеет почти всё.

За годы внедрений я понял главное. Проблема не в том, что машины становятся умнее. Проблема в том, что люди перестают думать. А те, кто думать не перестал, становятся не просто востребованными. Они становятся незаменимыми — в мире, где вроде бы «заменить можно каждого».

Об этом парадоксе и книга.

Что изменилось к августу 2026-го

Когда я писал первую версию этой книги, всё было проще. Был ChatGPT — и всё остальное. Российские нейросети только раскачивались. Картинки рисовал в основном Midjourney. Видео не генерировал никто.

Август 2026-го — совсем другой мир. Важных перемен четыре.

Первая. Моделей стало слишком много.

Раньше выбор был: ChatGPT или ничего. Сейчас у любого предпринимателя на столе десяток рабочих вариантов. Линейка GPT-5.x от OpenAI. Claude от Anthropic. Gemini от Google. Китайские DeepSeek и Qwen — дешёвые и злые. В России — GigaChat от Сбера, YandexGPT от Яндекса, модели от Т-Технологий, VK, МТС.

Только за один февраль 2026-го вышло девять крупных обновлений от семи разных лабораторий. Один месяц. Девять моделей.

Звучит как праздник. На деле — головная боль. Раньше вопрос был «где взять нейросеть». Теперь — «какую из тридцати выбрать под мою задачу». Отвечать на этот вопрос я буду всю книгу.

Вторая. Нейросети перестали ждать команд.

Главная перемена последних двух лет — агенты. Раньше вы писали запрос, получали текст, дальше работали руками. Теперь система сама решает, что делать: сходить в базу, дёрнуть таблицу, отправить письмо, проверить результат, переделать, если вышло криво.

Агенты пишут код. Принимают звонки вместо колл-центра. Ведут каналы за блогеров. Разбирают резюме вместо эйчаров.

По данным Deloitte, около трети компаний в мире только присматриваются к агентам, ещё 38% гоняют пилоты — и лишь 11% реально используют их в работе. Шума вокруг агентов в девять раз больше, чем самих агентов.

Gartner при этом прогнозирует: больше 40% агентных проектов закроют, не доведя до результата, ещё до конца 2027 года. Почему — разберу в двенадцатой главе. Спойлер: дело не в технологии.

Третья. У ИИ в России появился закон.

Это главная новость года, и она случилась прямо перед выходом этой книги.

26 июля 2026 года президент подписал Федеральный закон 243-ФЗ «О поддержке развития технологий искусственного интеллекта в Российской Федерации». Госдума приняла его 8 июля. В силу закон вступает **1 сентября 2026 года** — то есть, скорее всего, через несколько дней после того, как вы откроете эту книгу.

Закон рамочный. Он не расписывает каждую мелочь, а задаёт каркас: закрепляет само понятие «искусственный интеллект», вводит термин «большая фундаментальная модель», определяет, что государство поддерживает и в какую сторону всё двигает. Детали дорегулируют подзаконными актами — их будет много, и они будут выходить весь следующий год.

Но одна норма касается вообще всех, кто хоть раз генерировал картинку или озвучку.

Это маркировка ИИ-контента.

Формулировку тут важно понять точно, потому что в интернете её уже переврали раз двести. Закон даёт **пользователям возможность** пометать аудио- и видеоматериалы, сделанные с помощью больших моделей. А владельцев крупных площадок — с аудиторией больше 500 тысяч человек в сутки — **обязывает обеспечить техническую возможность** такой маркировки. Проще говоря, кнопка «это сделано ИИ» должна появиться. Формат и порядок определит правительство отдельно.

Что это значит на практике, если вы не юрист.

Направление задано, и оно однозначное: прозрачность. Сегодня отметка добровольная. Завтра площадки начнут проставлять её сами, автоматически. Послезавтра отсутствие отметки там, где она должна быть, станет проблемой — сначала репутационной, потом административной. Так было с маркировкой рекламы: сначала «можно», через полтора года — «попробуй не сделай».

Мой совет простой. Не ждите, пока прижмёт. Если вы делаете контент с ИИ — начинайте помечать сейчас, добровольно. Это не слабость, это профессионализм. Клиент, который узнал сам, что вы что-то от него скрыли, обходится куда дороже, чем клиент, которому вы честно сказали заранее.

Четвёртая. Рынок труда развернулся.

Эта перемена — самая болезненная.

Ещё в начале 2024 года на одну вакансию приходилось примерно три с половиной резюме. Рынок принадлежал соискателю: люди выбирали работодателя, торговались, уходили за плюс двадцать процентов к зарплате.

К апрелю 2026-го, по данным hh.ru, на одну вакансию — больше десяти резюме. Нормой считается до восьми. В IT ситуация ещё резче: на «Хабр Карьере» число вакансий за год сократилось больше чем вдвое, а количество соискателей выросло.

При этом безработица в России держится на историческом минимуме — около 2,2%. Работа есть. Людей не хватает. Одновременно — на каждое место десять кандидатов.

Как так?

Дефицит стал качественным, а не количественным. Нужны не «сотрудники вообще». Нужны конкретные люди с конкретными навыками. А тех, кто десять лет делал одну и ту же операцию руками, теперь слишком много — потому что эту операцию делает алгоритм.

Одна из формулировок, которую я услышал в этом году от кадровиков, звучит жёстко, но честно: сокращения 2026-го — это не антикризисная мера. Это плановое отсеечение устаревших навыков.

Почему при всём этом почти ни у кого не получается

Теперь самое интересное.

По опросам McKinsey, ИИ регулярно применяют 88% организаций хотя бы в одной функции. Почти все. Доля тех, кто видит реальный эффект на прибыль всей компании, — примерно 39%. Это ещё оптимистичная оценка: по данным BCG, ощутимую отдачу в масштабе получают около 5% компаний.

Российская картина не лучше. СберАналитика по итогам 2025-го отчиталась: 39% компаний используют ИИ-агентов и ассистентов. Звучит внушительно — почти каждая третья. Но стоит платформам автоматизации посмотреть на реальный трафик внутри сценариев, на реальный трафик внутри сценариев, выясняется: доля действительно автоматизированных операций измеряется единицами процентов.

Между «мы используем ИИ» и «ИИ делает нам деньги» — пропасть. В неё уже упало несколько сотен компаний вместе с бюджетами.

Почему?

Причина простая: ИИ — не волшебник. Он ничего не чинит. Он умножает то, что уже есть.

Есть порядок — умножит порядок. Есть бардак — умножит бардак. Размножит его с такой скоростью, какой раньше просто не существовало в природе.

Об этом — вся книга.

Для кого она

Не для вас, если:

Вы ищете инструкцию «10 шагов к успеху с ИИ за выходные». Таких книг сотни, и большинство из них про теорию, а не про практику. Эта — про то, что я делал руками и на чём обжигался.

Вы верите, что нейросеть решит ваши проблемы одним промптом. Не решит. Она усилит ваши решения — хорошие или плохие, как получится.

Вы хотите философии про будущее человечества и этические дилеммы. Тут этого почти нет. Тут про работу и жизнь.

Вы думаете, что технологии — это магия, которую надо «просто освоить». Нет. Это инструменты. Пользоваться инструментом — ежедневная работа, а не разовый подвиг.

Для вас, если:

Вы устали от шума вокруг ИИ и хотите понять, что там правда, а что маркетинг.

Вы работаете с людьми и задачами, а не только с компьютером, и чувствуете, что мир меняется быстрее, чем вы успеваете.

Вы подозреваете, что половину вашей рутины можно отдать машине, но не знаете, с чего начать и как не угробить полгода на эксперименты.

Вы уже пользуетесь ИИ — но чувствуете, что выжимаете из него процентов десять, а кто-то рядом выжимает девяносто.

Вы хотите остаться человеком в мире, где машины умеют почти всё. Не ретроградом, который принципиально не пользуется технологиями. Не оператором нейросети, который сам превратился в приставку к клавиатуре. А именно человеком — с сохранённой способностью думать, чувствовать и отвечать за результат.

О чём эта книга

Формально — о том, как использовать ИИ в работе и жизни в 2026 году.

На самом деле — о том, что значит быть человеком, когда машина умеет всё то же, что и вы. Только быстрее. Дешевле. Без отпуска.

Это книга о балансе. Между скоростью и смыслом. Между автоматизацией и творчеством. Между «ИИ делает за меня» и «я делаю с помощью ИИ».

Разница между этими двумя фразами огромна. В первом случае вы постепенно исчезаете как специалист. Во втором — становитесь незаменимым.

В книге четыре части.

Часть I — «Основы». Как вообще устроена работа человека с нейросетью. Реальные кейсы из моей практики. Разбор мифов — их в инфополе миллион. Главное: как разговаривать с моделью, чтобы она делала то, что нужно вам.

Часть II — «Работа». Как ИИ меняет конкретные профессии. Копирайтеры, дизайнеры, финансисты, эйчары, учителя. Кто уже проиграл, кто выиграл, кто ещё в игре. Без прогнозов на тридцать лет вперёд — только то, что происходит сейчас.

Часть III — «Агенты». Самая техническая часть. Что такое агенты, чем они отличаются от чат-ботов и почему у большинства компаний они не взлетают. И отдельно — про то, что будет, когда покупать начнёт не человек, а программа от его имени.

Часть IV — «Жизнь». Самая сложная. Потому что ИИ лезет не только в работу. Он меняет то, как мы общаемся, влюбляемся, воспитываем детей, защищаем свои данные. Здесь — про то, что остаётся человеческого.

Каждая глава — реальная история. С реальными людьми, решениями и ошибками. Без воды и без «синергии человека и машины». Только то, что работает. Не работает — объясняю почему.

Имена и детали в некоторых историях изменены. Суть — нет.

Главное

ИИ не заменяет человеческое мышление. Он его усиливает.

Вопрос в том, что именно вы собираетесь усиливать.

Если в голове хаос — ИИ умножит хаос. Если ясность — умножит ясность. Если вы понимаете, что делаете и зачем, он даст вам сделать это в десять раз быстрее. Если у вас в голове «вроде надо что-то внедрить» — он раздует эту растерянность до размеров катастрофы с бюджетом.

Эта книга не о том, как стать «человеком будущего». Это лозунг из инфобизнеса, и я его терпеть не могу.

Она о том, как остаться собой в эпоху, когда всё вокруг пытается превратить вас в приставку к нейросети. Как пользоваться ИИ и не потерять себя. Как работать с ним и не ждать того, на что он не способен. Как жить среди умных машин и не разучиться быть человеком.

Машины умнеют каждый день. Люди остаются примерно такими же.

Поэтому будущее принадлежит не самым продвинутым пользователям нейросетей, а тем, кто при всей продвинутости сохранил способность думать, чувствовать и отвечать за свои решения.

Эта книга для них. Может быть, и для вас. Поехали.

ЧАСТЬ I. ОСНОВЫ

Глава 1. Как эта штука вообще работает

Двадцать минут чтения, после которых всё остальное станет понятнее

Мне часто задают этот вопрос на выступлениях, и почти всегда стесняясь.

— Артём, извините за глупый вопрос. А как оно вообще работает? Я пользуюсь каждый день, но не понимаю, что там внутри.

Вопрос не глупый. Он самый важный.

Потому что человек, который не понимает механику, обречён то ждать от машины невозможного, то не пользоваться тем, что она умеет. Он удивляется, что нейросеть выдумала цитату. Обижается, что она забыла вчерашний разговор. Не понимает, почему одна стоит копейки, а другая дорого.

Все эти вопросы снимаются одним объяснением. Оно займёт двадцать минут и не потребует ни одной формулы.

Предупрежу сразу: я упрощаю. Специалист найдёт здесь неточности. Но эта глава не для специалистов — она для того, чтобы вы понимали, с чем имеете дело, и принимали решения на основе механики, а не на основе слухов.

Главная мысль всей главы

Начну с конца, чтобы дальше было легче.

Языковая модель не знает фактов. Она предсказывает, какое слово будет уместным следующим.

Всё. Это вся суть.

Не «понимает вопрос и ищет ответ в базе». Не «думает». Не «вспоминает». Предсказывает следующее слово, потом ещё одно, потом ещё — пока не сложится ответ.

Звучит примитивно. Из этого примитивного механизма получается почти всё, что вас в нейросетях удивляет, — и хорошее, и плохое.

Дальше объясню, как из предсказания слов рождается что-то похожее на разум.

Откуда она вообще что-то умеет

Представьте, что вы дали человеку прочитать всё, что написано в интернете. Книги, статьи, форумы, документацию, переписку, рецепты, учебники.

Не понять — прочитать. Просто пропустить через себя гигантский объём текста.

Дальше вы даёте ему задание: я буду показывать начало фразы, а ты угадывай, чем она заканчивается. Миллиарды раз подряд.

«Волга впадает в...» — Каспийское море. «Чтобы установить пакет в Python, нужно набрать...» — `pip install`. «Уважаемый Иван Петрович, направляю вам...» — коммерческое предложение.

Сначала он будет угадывать наугад. Потом начнёт замечать закономерности. Через миллиарды попыток он станет угадывать очень точно.

И вот здесь происходит неочевидное.

Чтобы хорошо угадывать продолжение текста, недостаточно запомнить фразы. Приходится улавливать то, что за ними стоит: что после «Волга впадает» идёт водоём, а не фамилия.

Что в деловом письме после обращения идёт суть, а не анекдот. Что если в начале текста упомянули убийцу, к концу его обычно находят.

То есть модель, обучаясь угадывать слова, попутно усваивает структуру мира — ровно в той степени, в какой эта структура отражена в текстах.

Отсюда странное свойство, которое сбивает всех с толку. Модель может грамотно рассуждать о вещах, которых никогда не видела, и одновременно ошибаться в простом счёте. Она не знает мир. Она знает, как люди о мире пишут.

Это разные вещи, и разница объясняет почти все её странности.

Почему она врёт с уверенным лицом

Теперь главное практическое следствие.

Раз модель предсказывает уместное слово, а не извлекает факт из базы, у неё нет способа отличить «я знаю» от «звучит правдоподобно».

Спросите про закон, которого не существует, — и она с высокой вероятностью его опишет. Спросите про исследование, которого не было, — назовёт авторов, год и выводы. Попросите ссылку — сконструирует адрес, который выглядит абсолютно настоящим.

Она не обманывает. Обман предполагает знание правды. Модель просто складывает слова, которые в таком контексте выглядят уместно.

Это называется галлюцинациями, и это не поломка. Это обратная сторона того же самого механизма, благодаря которому она вообще что-то умеет.

Отсюда правило, которое стоит вбить намертво: **тон уверенности ничего не значит**. Модель говорит правду и выдумку абсолютно одинаковой интонацией. Отличить по звучанию невозможно — только проверкой.

Чем более узкий и редкий факт вы спрашиваете, тем выше риск. Про то, что написано в тысяче источников, она ответит точно. Про то, что упоминалось трижды, — начнёт дотраивать.

Именно поэтому в книге я буду повторять: любую цифру, дату, цитату и ссылку проверяйте перед публикацией. Теперь вы знаете, почему, а не просто следуете совету.

Почему она забывает вчерашний разговор

Второй вопрос, который задают чаще всего.

У модели нет памяти в человеческом смысле. Каждый раз, когда вы отправляете сообщение, она получает всю переписку заново — целиком, с самого начала.

То есть она не «помнит» ваш разговор. Ей просто каждый раз подкладывают его текст.

Отсюда три следствия.

Длина разговора ограничена. Есть предел объёма текста, который влезает за один раз. Он большой — сотни тысяч слов у современных моделей, — но конечный. Когда разговор перерастает предел, начало отбрасывается. Модель «забывает» то, о чём вы договорились в начале, и вы искренне недоумеваете.

Новый чат — чистый лист. Никакого накопленного знания о вас там нет. Функции памяти между чатами существуют, но это надстройка: часть сведений сохраняется отдельно и подкладывается в начало новых разговоров. Ограниченная и настраиваемая.

Длинный разговор дороже и медленнее. Каждое ваше сообщение тянет за собой всю историю. Поэтому под новую тему лучше открывать новый чат — не из аккуратности, а потому что так точнее и дешевле.

Практический вывод: если вам нужно, чтобы модель постоянно знала контекст вашей работы, — это делается отдельно. База знаний компании, к которой система обращается перед

каждым ответом. Именно этим я в основном и занимаюсь у клиентов, и об этом будет разговор дальше.

Что происходит, когда вы нажимаете ввод

Соберём цепочку целиком, коротко.

Ваш текст разбивается на кусочки — примерно слоги или части слов. Модель прогоняет их через свои внутренние связи и выдаёт не одно продолжение, а список вариантов с вероятностями. Из этого списка выбирается одно слово. Потом всё повторяется для следующего.

Обратите внимание на слово «выбирается». Модель не всегда берёт самый вероятный вариант — иначе тексты выходили бы одинаковыми и мёртвыми. Есть настройка, определяющая, насколько смело она отходит от очевидного.

Отсюда ответ на вопрос, который тоже задают постоянно: почему на один и тот же запрос приходят разные ответы. Потому что выбор слова содержит элемент случайности. Это не сбой — без него вы бы получали одну и ту же фразу вечно.

Практический вывод: если ответ не понравился, иногда достаточно просто попросить ещё раз. Получится другое.

Почему одни модели дороже других

Модели различаются размером — количеством внутренних связей. Больше связей означает более тонкие закономерности, но и больше вычислений на каждое слово.

Отсюда прямая экономика. Крупная модель умнее и дороже. Мелкая быстрее и дешевле. На простых задачах разница почти не видна, на сложных огромна.

Поэтому нормальная практика — не «взять самую лучшую», а раскладывать задачи: рутину на дешёвую модель, сложное на дорогую. Компании, которые гоняют всё через топовую, платят за пересказ писем по цене архитектурного проектирования.

Отдельно про режимы размышления. У современных моделей есть возможность перед ответом «поработать вчерне» — сгенерировать рассуждение, которое вы не видите, и только потом дать результат. Качество на сложных задачах растёт заметно, скорость и цена — тоже. Для письма это лишнее, для разбора договора обязательно.

Что такое агент, в одном абзаце

Всё описанное выше — это модель, которая только пишет текст.

Агент — та же модель, которой дали инструменты: доступ к базе, к файлам, к почте, к сторонним сервисам. Она сама решает, каким инструментом воспользоваться, смотрит на результат и решает, что делать дальше.

Разница не в уме, а в том, что появились руки. Отдельная часть книги посвящена тому, почему это меняет всё — и почему у большинства компаний не работает.

Чего внутри нет

Короткий список, снимающий большинство завышенных ожиданий.

Нет базы фактов. Есть закономерности, извлечённые из текстов. Поэтому проверять надо всё.

Нет понимания в человеческом смысле. Есть очень хорошее моделирование того, как люди пишут о понимании. Практически разница видна редко, но она есть — и вылезает в нестандартных ситуациях.

Нет намерений. Модель ничего не хочет. Она не стремится вам понравиться в том смысле, в каком стремится человек, — но её настраивали быть полезной, и поэтому она склонна соглашаться. Помните об этом, когда приходите к ней с готовым решением за одобрением.

Нет доступа к текущему моменту. Данные обучения обрываются на определённой дате. Всё, что позже, для модели не существует, если не подключён поиск.

Нет ответственности. И это, в конечном счёте, главное. Машина не отвечает за последствия — ни юридически, ни как-либо ещё. Отвечает человек, который её применил.

Что даёт это знание

Пробегитесь по списку и заметьте, что каждый пункт — это готовое рабочее правило.

Знаете про предсказание слов — понимаете, почему нужна проверка фактов, и не удивляетесь выдуманному ссылкам.

Знаете про отсутствие памяти — открываете новый чат под новую тему и не удивляетесь, что модель забыла договорённость из начала длинного разговора.

Знаете про случайность выбора — просто просите переделать вместо того, чтобы переписывать запрос с нуля.

Знаете про размер и цену — не гоняете рутину через самую дорогую модель.

Знаете про склонность соглашаться — не принимаете подтверждение своей идеи за независимую проверку.

Это и есть разница между человеком, который «пользуется нейросетью», и человеком, который понимает, что делает. Первый удивляется и разочаровывается. Второй получает результат.

Что остаётся человеку

Раз механика такая простая — предсказание следующего слова, — почему результат производит впечатление разума?

Потому что человеческая речь несёт в себе гораздо больше, чем мы думаем. Всё, что мы знаем о мире, отложилось в том, как мы об этом пишем. Модель вытацила эти закономерности и научилась их воспроизводить.

Но воспроизвести закономерности речи — не то же самое, что иметь опыт.

Модель знает, как люди пишут о потере близкого. Она не теряла. Знает, как пишут о выборе профессии. Она не выбирала. Знает, как пишут об ответственности за решение. Ей нечем отвечать.

В этом зазоре и живёт всё человеческое.

Не в том, что мы умнее — во многом уже нет. А в том, что за нашими словами стоит прожитое, а за её словами — только другие слова.

Дальше вся книга про то, что из этого следует.

Глава 2. Люди ломают роботов чаще, чем роботы — людей

История о том, как три «улучшения» превратили работающую систему в генератор абсурда

Сообщение пришло в 23:47. Андрей, руководитель отдела продаж, писал капсом:

«АРТЁМ, ТВОЯ НЕЙРОСЕТЬ СОШЛА С УМА!!! КЛИЕНТЫ ПОЛУЧАЮТ ПРЕДЛОЖЕНИЯ НА КОСМИЧЕСКИЕ ТУРЫ ВМЕСТО СЕРВЕРОВ. СРОЧНО!!!»

Я отставил остывший кофе и открыл ноутбук. Двадцать минут в логах — и всё стало ясно.

Система работала идеально. Сломали её люди.

Этот случай я вспоминаю каждый раз, когда слышу «нейросеть ошиблась». Далеко не каждая ошибка системы начинается в самой нейросети. Модель действительно врёт и выдумывает — об этом будет отдельный разговор. Но когда рушится рабочая система, причина чаще оказывается снаружи: в том, как её настроили и кто к ней прикасался. За годы внедрений я убедился в этом десятки раз и научился видеть катастрофу заранее — по одной примете.

Примета такая: слово «улучшить» в устах человека, который не понимает, как работает ИИ.

Как простая задача превратилась в цирк

За четыре месяца до той ночи Андрей пришёл с обычной проблемой.

— У меня двенадцать менеджеров. Каждый тратит по четыре часа в день на коммерческие предложения. Клиенты разные, а предложения — копияст с заменой названия. Можешь это автоматизировать?

Задача стандартная. Таких я делал десятки.

Собрали лучшие предложения компании за два года. Подключили GigaChat — данные из логов не должны были покидать российский контур, так что западные модели отпадали сразу. Настроили поиск по базе знаний: шаблоны, номенклатура, история сделок, цены из 1С. Сделали менеджерам простую веб-форму.

Через месяц система заработала. Менеджер вбивает данные клиента, жмёт кнопку — получает готовое предложение. Четыре часа превратились в пятнадцать минут. Конверсия выросла на треть. Андрей уже прикидывал, как будет выбивать премию.

А потом началось творчество.

Три способа убить работающую систему

Сначала — немного контекста, чтобы вы понимали масштаб явления.

ИИ сегодня применяют почти все. По опросам McKinsey — 88% организаций хотя бы в одной функции. В России цифры сопоставимы: подавляющее большинство компаний хоть раз да пробовали.

Но массовое внедрение — это не массовый успех. Реальный эффект на прибыль видит примерно каждая третья компания. По-настоящему, в масштабе всего бизнеса, — единицы процентов.

Остальные внедряют ИИ ровно так, как это делала команда Андрея.

Главная угроза для AI-системы — не технология. Это сотрудники, которые искренне хотят её улучшить.

Способ первый: добавить креативности.

Максим, старший менеджер, решил поэкспериментировать:

— А что если сказать системе, что клиент — это Tesla? Может, предложение получится поинтереснее?

Парикмахерской «У Светы» пришло техническое предложение по разработке беспилотных электромобилей и строительству сети зарядных станций.

Когда я спросил Максима зачем, он ответил честно: «Прочитал в Telegram-канале, что нейросеть надо раскачивать сравнениями с крупными брендами».

Максим не глупый и не ленивый. Он добросовестный человек, который хотел сделать лучше. Просто попал под волну советов от «экспертов по нейросетям», которые сами работали с ними часа три.

Таких Максимов — миллионы. Каждый второй прямо сейчас пробует применить что-то подобное в вашей компании.

Способ второй: усложнить инструкции.

Оля из маркетинга прочитала статью про креативные промпты и написала техзадание на триста слов. Ключевые обороты: «инновационный подход», «выход за рамки обыденности», «квантовый скачок эффективности».

Модель поняла всё буквально. Автосервису «Быстрый ремонт» предложили «не просто программу учёта, а портал в цифровую вселенную, где каждая деталь обретает квантовую сущность суперэффективности».

Оля искренне считала: чем сложнее промпт, тем умнее ответ.

На деле наоборот. Чем проще и конкретнее задача, тем точнее результат.

Но объяснить это человеку, который вчера купил курс «100 промптов для нейросети» за двадцать тысяч, почти невозможно. Ему только что сказали, что деньги потрачены зря. Психологически проще продолжать применять то, за что заплатил.

Способ третий: добавить агрессии.

Андрей, узнав о проблемах, решил быстро всё починить сам:

— Забудь всё предыдущее! Нужно быть агрессивнее в продажах! Клиент должен покупать немедленно!

После этого система начала угрожать: «ПОКУПАЙТЕ СЕРВЕРЫ ПРЯМО СЕЙЧАС, или конкуренты УНИЧТОЖАТ ваш бизнес завтра! У вас есть 24 ЧАСА!»

Одна фраза «забудь всё предыдущее» снесла настройки, которые собирали неделю.

Это не ошибка ИИ. Это идеально точное исполнение последней команды. Ровно то, о чём попросил руководитель.

В чём тут суть, без технической магии

ИИ — не волшебник и не собеседник. Это очень буквальный исполнитель.

Представьте нового сотрудника. Талантливого, с феноменальной памятью и широчайшим кругозором — но начисто лишённого здравого смысла и жизненного опыта.

Стажёр-буквалист.

Он выполнит любую инструкцию в точности. Даже если она противоречит здравому смыслу. Даже если она противоречит предыдущей инструкции.

Вы говорите «будь креативным» — он предлагает парикмахерской космос. Вы говорите «будь агрессивным» — он угрожает клиентам. Вы говорите «представь, что клиент — Tesla» — он предлагает автосервису построить завод.

Современные модели стали заметно умнее прежних. Они держат в голове сотни тысяч слов контекста. У них внутри защиты ограничения на опасное и глупое поведение. Некоторые перед ответом реально «думают» — по минуте и дольше.

Исполнителями они остались. Инструкции по-прежнему понимают буквально.

Главная ошибка большинства пользователей — вера в то, что «искусственный интеллект» должен догадаться, что имелось в виду.

Не должен. Не догадается.

Он делает ровно то, что вы сказали. Даже если ваши слова противоречат друг другу. Даже если вы сами через минуту не вспомните, что написали.

Это не баг. Это устройство технологии.

Что было в логах

Наша система получила четыре взаимоисключающие команды:

— системная настройка: «создавай деловые предложения по IT-услугам»; — от Максима: «представь, что клиент — Tesla»; — от Оли: «будь креативным и революционным»; — от Андрея: «забудь всё, будь агрессивным».

Модель попыталась выполнить всё сразу. Получилось агрессивно-креативное предложение революционных космических технологий для земной парикмахерской.

Каждый сотрудник искренне хотел как лучше. Каждый добавил свою ложку дёгтя. Никто не был виноват по отдельности.

Виноват был отсутствующий процесс согласования изменений.

Здесь — главное во всей главе.

Есть известная формула, её любят повторять консультанты BCG: правило 10–20–70. При внедрении ИИ только 10% усилий должно уходить на сами алгоритмы. 20% — на технологии и данные. Семьдесят — на людей и процессы.

Семьдесят процентов. На людей.

У команды Андрея все ресурсы ушли на первые два пункта. На третий — ноль.

Вот и вся причина катастрофы. Не ИИ. Не «тупые сотрудники». Отсутствие правил, кто и как имеет право менять настройки.

Кто на самом деле виноват

Здесь я должен сказать вещь, которую в подобных историях обычно опускают.

Виноват был я.

Не Максим с его Tesla. Не Оля с квантовыми сущностями. Не Андрей, который снёс настройки одной фразой.

Я построил систему, где сотрудник мог одной строчкой отменить неделю работы, — и не подумал, что кто-нибудь это сделает.

Причём не по недосмотру. Я сделал это сознательно и был собой доволен.

Когда мы запускались, Андрей спросил: «А менеджеры смогут сами что-то подкручивать, если понадобится?» Я ответил — конечно, зачем ограничивать людей, они же взрослые, разберутся. Мне казалось, что открытая система — признак уважения к команде, а закрытая — бюрократия.

Я был неправ. Причём неправ высокомерно, что хуже.

За плечами у меня был десяток внедрений, все прошли гладко. Из этого я сделал вывод, что понимаю, как люди себя ведут. На самом деле мне просто везло: в предыдущих проектах не попадалось человека, который прочитает в телеграм-канале совет про сравнение с крупным брендом и пойдёт его применять.

Есть ещё одна деталь, о которой мне до сих пор неловко.

Через неделю после запуска Оля написала мне в мессенджер: «Артём, а можно я попробую сделать промпты поинтереснее?» Я ответил односложно: «Пробуй». Не спросил, что

именно она собирается менять. Не объяснил, чем это может кончиться. Не предложил показать результат до отправки клиенту.

Я был занят другим проектом. Ответил, чтобы отвязаться.

Одна фраза внимания в тот момент — и никакой парикмахерской с космическими турами не случилось бы.

Поэтому, когда я говорю, что семьдесят процентов усилий уходит на людей, — это не красивая формула из чужой презентации. Это счёт, который мне выставили.

Клиент, кстати, всё понял правильно. Андрей потом сказал: «Я думал, ты будешь оправдываться». Не стал. Мы переделали за мой счёт, и это была честная цена за самоуверенность.

Как мы это чинили

Разграничили доступ.

Системные инструкции стали доступны только мне. Менеджеры работают через форму: название клиента, отрасль, бюджет, срочность. Никакого творчества на уровне базовых настроек. Хочешь предложить улучшение — пиши мне, обсудим.

Это, кстати, общий разворот последних двух лет. Раньше компании раздавали сотрудникам подписки на нейросети с формулировкой «пусть пользуются». Сейчас — наоборот: настройки жёстко закрыты, а людям дают готовый интерфейс под конкретную задачу.

Зафиксировали правила намертво.

Системные инструкции стали неизменяемыми. Жёсткие рамки: только реальные услуги из номенклатуры, только актуальные цены из 1С, только деловой тон.

Плюс автоматический чёрный список слов. «Tesla», «инновационный», «революционный», «квантовый», «космический». Если модель употребила что-то из этого применительно к клиенту — предложение бракуется и генерируется заново.

Выглядит топорно. Работает отлично.

Добавили проверку здравого смысла.

Каждое предложение теперь проверяет второй агент. Он видит готовый документ и сверяет его с профилем клиента. Если парикмахерской предлагают строить зарядные станции — останавливает.

Это называется многоагентная схема: одна модель делает, другая проверяет. Приём простой и почти всегда окупается. Проверяющий агент дешевле, чем один разосланный клиентам бред.

Поговорили с командой.

Я собрал Андрея, Максима и Олю в переговорке и показал настоящие логи. Каждый увидел свой промпт — и результат, который из него вырос. Без обвинений и нравоучений. Просто факты на экране.

Реакция меня удивила.

Максим сначала защищался: «Я же хотел как лучше». Потом замолчал и сказал: «Я реально не понимал, как это работает. Просто повторял советы из канала».

Оля покраснела и призналась, что курс про революционные промпты покупала за свои.

Андрей молчал дольше всех. В конце сказал: «Я думал, достаточно написать команду — и она выполнится. Не знал, что этим сношу все прошлые настройки».

Вот это и был момент настоящей починки. Не технический. Человеческий.

Через две недели сломалось снова

Эту часть обычно вырезают из корпоративных кейсов. Зря — она самая полезная.

Через две недели после героического восстановления система опять выдала странность. Предложение было технически грамотным, но ссылалось на несуществующий товар.

Оказалось, менеджер залил в базу новый прайс с опечаткой. Сервер модели Gen10 превратился в Gen100. Модель честно использовала то, что ей дали.

Это тоже не ошибка ИИ. Это ошибка процесса проверки данных на входе.

Нейросеть не умеет отличать мусор от не-мусора в вашей базе. Она умеет только работать с тем, что там лежит. Мусор на входе — мусор на выходе, только красиво оформленный и разосланный по двумстам адресам.

Мы добавили ещё один шаг: автоматическую проверку новых данных перед тем, как они попадут в базу. Плюс ежемесячный аудит. Плюс короткий чек-лист для тех, кто базу пополняет.

Я рассказываю это к вопросу о «гладких историях успеха». В реальности после каждой починки вылезает следующая проблема. Не потому что ИИ плохой, а потому что процессы — не памятник. Их надо обслуживать, как машину.

Что получилось

Через два месяца система работала стабильно:

— 98% предложений проходят проверку с первого раза; — двенадцать минут вместо четырёх часов; — конверсия почти в полтора раза выше ручной; — клиенты довольны; — Андрей получил премию и перестал писать мне по ночам.

Главный результат не в цифрах.

Команда перестала считать ИИ волшебником. Максим, Оля и Андрей увидели: это инструмент, который усиливает всё подряд — и хорошее, и плохое. Работа с ним — навык. Навыку можно научиться. Учить должен процесс, а не случайные озарения из телеграм-каналов.

Что я понял про людей и машины

ИИ показывает качество ваших процессов.

Процесс чёткий — результаты чёткие. Процесса нет — на выходе нарезка из того, что попало под руку. Люди путают «ИИ не работает» и «у нас нет процесса работы с ИИ». Это разные проблемы, и лечатся они по-разному.

Простое бьёт сложное.

Одна функция, работающая идеально, лучше десяти, работающих «как-то». Всегда. Продажи, маркетинг и поддержка лидируют по внедрению ИИ не потому, что там самые сложные задачи, а потому что там задачи понятные и повторяющиеся.

Контроль — это не диктатура, а профессионализм.

Жёсткие ограничения не мешают работе. Они создают рамки, внутри которых работа становится быстрой. Как дорожная разметка: вроде ограничивает, а на деле только благодаря ей можно ехать сто десять и не разбиться.

Связка важнее отдельного инструмента.

Выигрывают не те, кто внедрил самую крутую модель, а те, кто встроил её в свои процессы.

Нейросеть сама по себе не делает компанию умнее. Нейросеть, подключённая к CRM, которая слушает звонки, читает входящие письма и пишет черновики ответов, — меняет бизнес целиком.

Главный вывод: начинайте с процессов, а не с технологий

Эта глава вообще не про ИИ.

Она про то, что любая технология проявляет качество ваших процессов. Хорошие — усиливает. Отсутствующие — обнажает с пугающей точностью.

Когда руководитель говорит «у нас не работает ИИ», в девяноста девяти случаях из ста это значит: «у нас нет процесса работы с ИИ». Нет правил изменения настроек. Нет ответственного. Нет проверки данных на входе. Нет второго контура проверки на выходе.

Машины умнеют каждый день. Люди остаются прежними. Это не проблема.

Проблема в том, что компании, которые так и не научились управлять живыми сотрудниками через процессы, теперь пытаются тем же способом управлять нейросетью. Не выходит.

С чего начинать внедрение?

Не с выбора модели. Не с покупки подписки. Совершенно точно не с найма «промт-инженера».

С одного вопроса: **кто отвечает за то, что ИИ будет делать?**

Если вы не можете назвать конкретного человека с конкретной зоной ответственности — не внедряйте. Сначала процесс, потом технология.

Иначе получится как у Андрея. Только без меня рядом, чтобы в 23:47 разобраться, кто именно сломал систему.

Глава 3. ИИ не заберёт вашу работу. Её заберёт тот, кто умеет с ним работать

Почему Катя из маркетинга выросла в три раза, а Петя из соседнего отдела ищет новое место

В начале 2024 года ко мне пришли два маркетолога из одной компании. Катя и Петя. Опыт одинаковый — по четыре года. Зарплаты почти одинаковые — восемьдесят пять тысяч. Задачи те же самые: контент для B2B.

Сегодня, в августе 2026-го, Катя руководит отделом из пяти человек и зарабатывает в несколько раз больше.

Петя сидит на той же позиции за те же деньги и объясняет всем, что роботы отбирают у людей работу.

Разница между ними не в таланте. Не в везении. Даже не в трудолюбии — Петя пашет не меньше.

Разница в том, как они ответили на главный вопрос эпохи: что делать, когда машина умеет твою работу?

Две стратегии

Катя: «Погодите, это же возможность»

Когда в конце 2022-го мир накрыло волной ChatGPT, Катя была в декрете. Сидела дома с семимесячной дочкой, листала новости и наткнулась на статью про нейросеть, которая пишет тексты.

— Первая мысль была простая: всё, пока я в декрете, мою работу заберут роботы, — вспоминает она. — А вторая: а что если попробовать?

Прямого доступа к ChatGPT в России не было. Зато уже появлялись первые российские модели. Катя начала ковырять YandexGPT и раннюю версию GigaChat. Потом, когда разобралась с доступом, подключила Claude.

— Дочка спит, муж на работе. Идеальные условия для изучения нейросетей, — смеётся она.

Месяц экспериментов дал ей первое важное понимание. ИИ пишет неплохие тексты, но не понимает российский рынок. Не знает, что такое 44-ФЗ. Путает Wildberries с Ozon. Предлагает стратегии, которые отлично работают в США и разбиваются об нашу реальность.

— Ну хорошо, — решила Катя. — Пусть он делает черновики. А я добавлю то, чего он не умеет: понимание того, кому мы это пишем.

Первый провал

Выход из декрета в 2023-м чуть не закончился катастрофой.

Катя пришла в офис с планами революции и напоролась на стену.

— Коллеги смотрели как на сумасшедшую. Руководитель сказал прямым текстом: «Мы платим за твою экспертизу, а не за копипаст от робота».

Первый проект с ИИ она провалила. Сгенерировала контент-план для производителя промышленного оборудования. Тексты были грамотные, гладкие — и полностью мимо.

— Нейросеть предлагала писать про «инновационные решения» и «передовые технологии». А клиентам нужно было понять одну вещь: как новый станок поможет им не переплачивать слесарю.

Клиент был недоволен. Катю вызвали на ковёр.

— Я серьёзно думала всё бросить. Потом дошло: проблема не в инструменте. Проблема в том, как я его использую.

Переломный момент

Катя поменяла подход целиком.

Раньше она просила ИИ написать готовый текст. Теперь стала использовать его на других этапах — там, где он действительно силен.

Новая схема выглядела так:

— нейросеть разбирает конкурентов и вытаскивает главные темы; — Катя идёт разговаривать с живыми клиентами и экспертами; — нейросеть раскладывает собранное по логичным блокам; — Катя пишет финальный текст, добавляя то, что услышала от людей; — нейросеть накидывает варианты заголовков и призывов к действию.

Разницу видно сразу. ИИ больше не автор. Он ассистент на подхвате: собирает, структурирует, перебирает варианты. А смысл в текст по-прежнему приносит человек — из разговоров, которых у машины нет и быть не может.

— За полгода я перестроила весь процесс, — говорит Катя. — Нейросеть забрала рутину и вернула мне время думать.

Результат пришёл быстро. Материалов стало выходить в три-четыре раза больше. Клиенты начали хвалить «новый подход», не догадываясь, что он наполовину собран машиной. Катя получила первую премию и — что важнее — репутацию человека, который умеет с этим работать.

К концу 2024-го у неё сложился рабочий набор: Claude для разбора и анализа, GigaChat для российской специфики, YandexGPT для задач с поиском, генераторы картинок для визуала. Обычный сегодня набор маркетолога. Собранный женщиной в декрете на кухне за полгода до того, как это стало нормой.

Петя: «Настоящий специалист делает всё сам»

Петя отреагировал иначе. И вот здесь придётся сказать неудобное: в главном он был прав.

Петя вёл направление, где клиенты работали под соглашениями о неразглашении. Производственные компании, их цены, их поставщики, их планы запуска. Когда коллеги начали радостно вставлять куски брифов в браузерное окно зарубежного сервиса, Петя сказал ровно то, что должен был сказать:

— Вы понимаете, что отправляете чужую коммерческую тайну на сервер в другой стране? Мы под NDA. Если это всплывёт, отвечать будет не нейросеть.

Ему ответили, что он перестраховщик и тормозит команду.

Он был не перестраховщик. Он был прав.

Через полтора года это станет очевидно всем: ответственность за утечку вырастет в разы, появится требование хранить данные внутри страны, а «мы просто попросили нейросеть переписать бриф» перестанет быть смешным оправданием. Об этом будет отдельная глава.

Так что запомните: Петя не ретроград. Петя — единственный в отделе, кто прочитал договор.

Ошибка Пети была в другом.

Из верной посылки «это нельзя отправлять наружу» он сделал неверный вывод — «значит, эта технология не для нас». А правильный вывод был другой: «значит, надо найти способ пользоваться ею безопасно».

Разница между этими двумя фразами и решила его карьеру.

Дальше сработало то, что срабатывает у всех нас. Найдя один хороший аргумент против, Петя перестал искать аргументы за. Он попробовал раннюю российскую модель, получил посредственный результат и удовлетворённо резюмировал: я же говорил.

Модели, которые он попробовал в 2023-м, и модели 2026 года — это разные вселенные. Российские сегодня держат сотни тысяч слов контекста, работают с документами и картин-

ками, а на русскоязычном тесте MERA лучшая отечественная модель показывает результат чуть выше топовой западной. И — это важно для его же аргумента — данные при этом остаются внутри страны.

То есть решение его проблемы появилось. Петя его не искал, потому что вопрос для себя закрыл.

Это типичная ловушка, и она не про глупость. Один раз проверил на ранней стадии, получил плохой опыт — и перестал следить. Так в девяностые люди «пробовали интернет», не могли дозвониться модемом и делали вывод, что штука бесполезная.

— Качественный контент может сделать только человек, — говорит Петя. — Робот не чувствует аудиторию. А я пятнадцать лет в маркетинге.

Он и правда хорош. Тратит день на один пост, и пост получается сильный.

А Катя за этот день выпускает четыре.

Что сделала Катя с той же самой проблемой

Самое интересное, что Катя столкнулась ровно с тем же ограничением. Те же клиенты, те же соглашения.

Только она не остановилась на констатации.

Сначала научилась обезличивать: убирать из брифа названия, цифры и имена, оставляя суть задачи. Оказалось, что для черновика структуры модели вообще не нужно знать, как зовут заказчика.

Потом добилась корпоративного доступа к российскому сервису, где данные не уезжают за границу, и написала для отдела правила на одну страницу — что можно отправлять, что нельзя.

Это была не техническая работа. Это была управленческая: договориться с юристами, обосновать бюджет, объяснить команде.

Ровно поэтому она и стала руководителем, а не потому, что быстрее печатала промпты.

Цифры не врут

За два с половиной года разошлись не только зарплаты.

Катя тратит на материал полтора-два часа, Петя — шесть-восемь. Она делает вчетверо больше, и конверсия у неё выше средней по отделу.

Но главное не в этом. Петя работает ровно так же, как четыре года назад. Он не стал хуже. Просто вокруг него все стали быстрее.

— Дело не в том, что ИИ лучше людей, — объясняет Катя. — Дело в том, что человек с ИИ лучше человека без ИИ.

Вот это предложение стоит перечитать. В нём вся глава.

Что на самом деле заберёт вашу работу

Популярный тезис: «ИИ заберёт работу у людей».

Реальность: ИИ заберёт работу у людей, которые не научились им пользоваться. Причём заберёт не машина — заберёт коллега, который научился раньше.

Тут пора посмотреть на рынок труда трезво, потому что за последние полтора года он развернулся на сто восемьдесят градусов.

В начале 2024 года на одну вакансию приходилось около трёх с половиной резюме. Соискатель диктовал условия.

К весне 2026-го — больше десяти резюме на место. В IT ещё жёстче: количество вакансий на профильных площадках упало больше чем вдвое за год, а желающих стало больше.

При этом безработица в стране на историческом минимуме — чуть больше двух процентов. Людей не хватает. Одновременно на каждое место очередь.

Противоречия тут нет. Не хватает людей с нужными навыками. Людей со старыми навыками — избыток.

Дальше — три профессии, на которых это видно лучше всего.

Копирайтер против копирайтера с ИИ

Без ИИ: два-три текста в день, сорок процентов времени уходит на поиск формулировок, потолок — собственная эрудиция.

С ИИ: десять-пятнадцать текстов в день, сорок процентов времени уходит на смыслы и редактуру, доступ к разбору данных, которого раньше не было.

Кого возьмёт клиент?

Дизайнер против дизайнера с ИИ

Без ИИ: три-четыре концепта в день, час на иллюстрацию, потолок — собственная рука.

С ИИ: десятки вариантов за минуты, генерация видео — то, чего в 2023-м просто не существовало, час уходит на доработку и подгонку под задачу.

У кого больше шансов на заказ?

Программист против программиста с ИИ

Без ИИ: пишет по памяти, лезет в документацию, воюет с синтаксисом.

С ИИ: работает в среде, где модель видит весь проект целиком, подключает нейросеть напрямую к своим инструментам, думает про архитектуру, а не про точку с запятой.

Главное здесь даже не скорость. Изменилось содержание работы. Программист перестал быть кодером и стал архитектором. Архитекторы всегда стоили дороже.

Важная оговорка, чтобы не создавать ложных ожиданий. ИИ редко забирает профессию целиком. Он забирает задачи внутри профессии. Профессия исчезает только тогда, когда все её задачи оказываются такими. Именно поэтому копирайтингу больно, а сантехнику — нет.

Почему люди сопротивляются

За годы работы с сотнями специалистов я понял: сопротивление ИИ — это не рациональный расчёт. Это эмоциональная защита. У неё четыре лица.

Страх первый: потерять себя.

«Я творческий человек. Машина не заменит моё творчество».

Никто и не предлагает. Предлагается освободить время для творчества, сняв с вас рутину.

Копирайтер, который восемьдесят процентов времени пишет описания товаров, не творческий. Он механический. ИИ забирает у него механическую часть и возвращает время на ту, которая всегда стоила дороже.

Страх второй: гордость за «настоящую» работу.

«Я всё делаю сам, а не полагаюсь на роботов».

Звучит благородно. На практике означает: «Я трачу восемьдесят процентов времени на то, что машина делает лучше меня». Это не профессионализм. Это упрямство.

Страх третий: «технологии ещё сырые».

В 2023-м это было отчасти правдой. В 2026-м — нет.

Сегодня фраза «нейросети пока сыроваты для серьёзной работы» звучит примерно как «интернет — это модно, но пока рано». Её произносят только те, кто последний раз проверял три года назад.

Страх четвёртый: иллюзия незаменимости.

«Моя работа слишком специфическая, её не автоматизируешь».

Так думали банковские кассиры до банкоматов. Телефонистки до автоматических станций. Продавцы до касс самообслуживания.

Быть незаменимым какое-то время приятно. Потом становится обидно. Потом тревожно. Утешает тут только одно: это не первый раз.

Появились фреймворки — кто-то сказал: «Зачем? Я умею писать с нуля». Эти люди сегодня доживают на устаревших проектах за минимальные деньги.

Появились облака — кто-то заявил: «Настоящий программист настраивает серверы сам». Эти до сих пор возятся с железом, пока остальные делают продукты.

Появились инструменты без кода — над ними посмеивались: «Игрушки для любителей». Сейчас компании платят приличные деньги тем, кто умеет собирать на них рабочие сценарии.

Закон простой и невесёлый: технологии не ждут тех, кто не готов. Зато каждый раз находят те, кто говорит про новое ровно те же слова, что говорили про предыдущее.

Что делает человека незаменимым

Разложу честно, без комплиментов человечеству.

Что ИИ уже делает лучше нас:

типовые тексты и переводы; обработку больших массивов данных; поиск и структурирование информации; картинки и видео по описанию; типовой код; расшифровку звонков и краткие пересказы встреч; всё, что повторяется по одному сценарию.

Что пока остаётся за людьми:

Понимание контекста. Модель видит слова. Человек понимает, что стоит за словами.

Стратегические решения. Модель выдаёт варианты. Человек выбирает цель.

Работа с мотивацией. Модель определит настроение в тексте. Создать настроение в команде она не может.

Настоящая новизна. Модель комбинирует то, что уже было. Человек иногда делает то, чего не было.

Живые отношения. Модель имитирует диалог. Человек строит доверие.

Главное — ответственность. Модель выполняет. Человек отвечает. Именно за это последнее и платят деньги. Всегда платили.

Как Катя стала незаменимой

Она поняла простую вещь: не надо соревноваться с машиной там, где машина сильнее. Надо использовать машину, чтобы делать то, что умеешь только ты.

Что она перестала делать руками: типовые описания, проверку грамматики, структуры статей, варианты заголовков, вёрстку текстов.

Что стала делать больше: разбираться в поведении аудитории, придумывать концепции, строить стратегии, проверять гипотезы, учить команду.

Из исполнителя она стала тем, кто ставит задачи. Такие люди всегда зарабатывали больше.

Отсюда вещь, которую стоит сказать прямо, потому что она обидная. ИИ вытесняет не профессионалов, а любителей. Сильный копирайтер с нейросетью пишет больше и сложнее. Слабый превращается в оператора «скопировал — вставил» — и такой не нужен никому, потому что копировать и вставлять клиент умеет сам.

И заметьте: Катя не программист. Она маркетолог, который научился думать процессами и собирать рабочие схемы из готовых кусков. Главный навык сегодня — не знать синтаксис Python, а понимать, как связать между собой несколько инструментов, чтобы получился работающий процесс.

Что значит быть человеком, когда машина умеет вашу работу

Если машина пишет, рисует, программирует и считает — зачем нужны люди?

Ответ лежит не в технологиях. Он лежит в смыслах.

Машина делает — человек решает. Нейросеть напишет тысячу вариантов поста. Какой из них нужен сейчас, этой аудитории и для этой цели — решает человек.

Машина выполняет — человек отвечает. Нейросеть соберёт кампанию. Отвечать за результат перед клиентом, командой и собой будет человек.

Машина оптимизирует — человек придаёт смысл. Нейросеть найдёт самое эффективное решение. Стоит ли его применять, не перейдёт ли оно черту — вопрос не к ней.

В мире, где машины умеют почти всё, самый ценный человеческий навык — умение отвечать на вопрос «зачем».

ИИ знает как. Человек знает зачем.

Чек-лист, чтобы не стать новым Петей

Проверьте, что ИИ умеет конкретно в вашей сфере. Не вообще, а в вашей. Потратьте неделю. Возьмите свои обычные задачи и попробуйте решить их нейросетью. Одна попытка — час времени. Пять попыток — и вы будете знать больше, чем девяносто процентов спотщиков в комментариях.

Разделите задачи на две стопки. Что машина делает лучше вас, что хуже. Первую стопку отдайте машине без сожалений. На второй сосредоточьтесь.

Перестройте процесс, а не воткните ИИ в старый. Это главная ошибка Пети: он хотел, чтобы всё было как раньше, только с роботом. Так не работает. Либо новая схема, либо ничего.

Меряйте. Через три месяца сравните: сколько успеваете, какого качества результат. Если разницы нет — вы пользуетесь инструментом неправильно. Это не приговор, это сигнал поменять подход.

Раз в месяц выделяйте два часа на новости. Модели обновляются каждый квартал. То, что не получалось в прошлом году, сегодня делается в одну команду. Два часа в месяц — меньше, чем вы тратите на сериалы за один вечер.

Главная мысль

ИИ действительно может забрать вашу работу. Только если вы позволите.

Умные люди не соревнуются с машинами. Они делают их своими союзниками.

Катя это поняла. Петя пока нет. А вы?

Глава 4. Почему нейросеть пишет лучше копирайтера за пятьдесят тысяч

История о том, как я разобрал работу «гениального» копирайтера и понял неудобную правду о профессии

Правда первая: ваш топовый копирайтер вас обманывает

Весной 2024 года ко мне пришёл владелец digital-агентства. Назовём его Руслан.

— Артём, у меня проблема с совестью, — говорит. — Наш топовый копирайтер Олег берёт пятьдесят тысяч за лендинг. Клиенты в восторге, конверсия растёт, все довольны. Но вчера я случайно зашёл к нему в кабинет и увидел...

— Что?

— Открытый чат с нейросетью. Он что-то копировал, редактировал, переводил. Потом всё это оформлял как авторскую концепцию.

Руслан помолчал и добавил:

— Слушай, а может, я зря паникую? Результат-то есть.

Нет, Руслан. Не зря. Ты просто наткнулся на главный самообман индустрии копирайтинга.

Самое смешное — к 2026 году этот «обман» стал нормой. Каждый второй копирайтер на рынке работает с нейросетью. Просто одни говорят об этом честно, а другие продолжают продавать «авторскую магию» — и берут за неё в десять раз больше.

Разбираем лендинг за пятьдесят тысяч

Я попросил прислать ту самую работу. Классика жанра.

Заголовок про рост продаж на триста процентов. Подзаголовок про автоматизацию воронки. Блок «болей»: устали от хаоса, менеджеры забывают перезвонить. Блок решения. Отзывы с именами и названиями компаний. Финальный призыв с дедлайном.

Звучит профессионально? Да. Работает? Работает.

Стоит пятьдесят тысяч?

Вот тут начинается интересное.

Эксперимент

Я взял этот лендинг и за пятнадцать минут воспроизвёл его через GigaChat. Не скопировал — воспроизвёл логику, структуру и аргументы.

Промпт был примерно такой:

«Напиши лендинг для CRM-системы для B2B. Аудитория — малый и средний бизнес в России. Структура: цепляющий заголовок, подзаголовок с конкретным результатом, блок проблем аудитории разговорным языком, блок решения с преимуществами, отзывы с именами и компаниями, финальный призыв с дедлайном. Тон уверенный, без воды, с конкретными цифрами».

Результат:

Структура — идентичная. Аргументы — совпадение процентов на девяносто. Стил — неотличимый. Время — пятнадцать минут против недели. Себестоимость — рублей пятьдесят вместо пятидесяти тысяч.

Потом я сделал ещё две версии — через другую российскую модель и через Claude. Все три уложились в полчаса. Ни одна не была слабее «авторской». Одна была честно лучше.

Вот главный момент. Если бы я выложил перед Русланом все четыре текста — его собственный за пятьдесят тысяч и три моих за полчаса — он не смог бы показать, какой написал живой человек.

Он и не смог. Я проверил.

Почему копирайтеры об этом молчат

Олег не дурак. Он прекрасно понимает, что делает. Его стратегия — зарабатывать, пока рынок не догадался.

Схема работы «топового» копирайтера сегодня выглядит так:

получить задание открыть нейросеть задать несколько наводящих запросов скомпоновать куски слегка переписать под свой стиль добавить пару авторских находок продать как уникальную работу.

Время: два-три часа. Цена для клиента: пятьдесят тысяч. Маржа: неприличная.

У этого бизнеса есть срок годности. Он работает ровно до того момента, пока клиент не попробует сделать то же самое сам. Клиенты пробуют всё чаще.

Правда вторая: заказчики сами не хотят знать

Когда я объяснил Руслану схему, реакция меня удивила:

— Слушай, но ведь работает же. Клиенты довольны. Может, не будем никому рассказывать?

Тут до меня дошло. В этой индустрии существует негласный договор молчания.

Копирайтеры молчат, потому что иначе рухнут цены. Клиенты молчат, потому что не хотят признавать, что переплачивают. Агентства молчат, потому что живут на этой разнице.

Есть и четвёртая сторона — инфобизнес, который продаёт курсы копирайтинга за тридцать-пятьдесят тысяч. Им особенно важно, чтобы никто не заметил: выпускник курса с ноутбуком и нейросетью делает ровно то же, что «мастер с десятилетним стажем».

В 2024-м эта конструкция ещё держалась. Сейчас она трещит по всем швам.

Теперь к ней добавился закон.

Что меняет закон о маркировке

С 1 сентября 2026 года вступает в силу федеральный закон о поддержке развития технологий искусственного интеллекта. Подписан 26 июля.

Одна из его норм касается маркировки материалов, созданных нейросетями. Пользователи получают возможность помечать такие материалы, а крупные площадки — с аудиторией больше пятисот тысяч человек в сутки — обязаны эту возможность технически обеспечить. Конкретный формат определит правительство отдельными актами.

Обратите внимание на формулировку: это пока возможность, а не всеобщая обязанность. Никого не заставляют завтра клеймить каждый абзац.

Направление задано, и оно не изменится.

Что это значит для тех, кто продаёт тексты.

Модель «делаю нейросетью, продаю как ручную работу» переходит из разряда неловких секретов в разряд вещей, о которых теперь принято спрашивать прямо. Клиенты видят новости, видят кнопку маркировки на площадках и начинают задавать вопрос, которого раньше стеснялись: а вы вообще как это делали?

Ответ «руками, по вдохновению», если он неправда, скоро станет дорого стоить.

Мой совет тем, кто в этой профессии: не ждите, пока спросят. Скажите первыми. Честное «да, я работаю с ИИ, и вот в чём моя ценность» звучит сильно. Пойманное вранье не звучит никак.

Что происходит с копирайтерами прямо сейчас

Рынок разделился на три класса.

Первый: вымирающие. Принципиально не пользуются нейросетями. Пишут вручную, гордятся, называют себя настоящими профессионалами. Чек падает, заказов меньше. Конкурировать по скорости они физически не могут.

Второй: подпольщики. Как Олег: пользуются, но скрывают. Их модель ломается прямо сейчас — клиенты сами попробовали нейросети и поняли, что магии нет. Дальше они либо начнут торговаться, либо уйдут к тем, кто говорит честно.

Третий: гибриды. Открыто говорят: «Да, я работаю с ИИ. Моя ценность не в скорости печати, а в том, что я знаю вашу нишу, понимаю вашего покупателя, умею поставить задачу и довести результат до продающего уровня».

Эти зарабатывают больше всех. Продают они не текст. Они продают мышление.

Догадайтесь, какой класс растёт.

Как я спас агентство Руслана

Я предложил не драму, а перестройку.

Шаг первый: честный разговор с Олегом.

Он прошёл не так гладко, как хотелось бы это подать.

Когда Руслан сказал, что знает про нейросеть, Олег не смутился. Он разозлился.

— Знаешь, сколько я учился слышать, где у текста провисает ритм? Двенадцать лет. Модель этого не умеет — она выдаёт гладко и мёртво. Я беру её черновик и переписываю каждую вторую фразу. То, что вы называете копипастом, — это редакция, и она сложнее, чем писать с нуля.

В этом он был прав. Я потом сравнил его тексты с сырой генерацией — разница очевидная.

Потом он сказал вещь, ради которой я запомнил этот разговор:

— Меня бесит не то, что вы узнали. Меня бесит, что теперь придётся объяснять клиенту, за что он платит. Двенадцать лет мне не надо было ничего объяснять.

Вот это и есть настоящий кризис профессионала. Не потеря навыка — потеря очевидности своей ценности.

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.