



РАЗГОВОР ОБ ИИ

КАК ПРЕВРАТИТЬ МОДНУЮ ТЕХНОЛОГИЮ
В ПОЛЬЗУ ДЛЯ КЛИЕНТОВ, СОТРУДНИКОВ
И ПРИБЫЛИ

ИТОН БЛЕЙК

Итон Блейк Разговор об ИИ. Как превратить модную технологию в пользу для клиентов, сотрудников и прибыли

*<https://litres.ru/74253713>
SelfPub; 2026*

Аннотация

Эта книга поможет Вам перестать гадать по отчётам и начать видеть, где реклама продаёт, где скидка съедает прибыль, где бренд держится на доверии, а где просто делает вид. Вы узнаете, как использовать AI: задавать точные вопросы, проверять гипотезы, читать язык клиентов, находить слабые места продукта и не отдавать алгоритму право думать за Вас. Вы сможете меньше тратить на пустые кампании, быстрее замечать ошибки, честнее смотреть на данные и собирать маркетинг, который работает не в презентации, а в реальной покупке.

Содержание

Введение	4
Трезвый разговор об ИИ	10
Это началось раньше, чем кажется	34
Не магия, а привычка считать	51
Конец ознакомительного фрагмента.	57

Итон Блейк Разговор об ИИ. Как превратить модную технологию в пользу для клиентов, сотрудников и прибыли

Введение

Последние два года я почти не встречал руководителей, которые не думали бы об искусственном интеллекте. Одни уже купили несколько подписок, раздали их сотрудникам и ждут чуда. Другие провели пилот, получили аккуратную презентацию с зелеными галочками и теперь не знают, что делать дальше. Третьи вообще не начали, но на каждом совещании чувствуют, что от них ждут решения. ИИ оказался в странном положении: о нем говорят все, а ясности от этого не становится больше.

В этом нет ничего удивительного. Когда технология становится модной, вокруг нее быстро вырастает рынок обеща-

ний. Продавцы рассказывают, что она сократит расходы, увеличит выручку, ускорит сотрудников, улучшит клиентский опыт и заодно спасет компанию от конкурентов. Скептики отвечают, что все это пузырь, что модели выдумывают факты, что сотрудники разучатся думать, а расходы на внедрение съедят любой выигрыш. Оба лагеря говорят громко. Руководителю от этого легче не становится.

Эта книга написана для человека, которому нужно принять решение, а не победить в споре о будущем человечества. Вам не нужно становиться инженером машинного обучения. Вам не нужно знать названия всех моделей, читать каждую новость о новом чат-боте и спорить о том, где проходит граница между статистикой и разумом. Но вам придется понимать достаточно, чтобы не путать демонстрацию с результатом, пилот с внедрением, удобную игрушку с рабочим инструментом.

ИИ уже вошел в бизнес, но вошел неровно. Где-то он помогает сотрудникам писать письма, искать информацию, собирать черновики документов. Где-то он управляет ценами, прогнозирует спрос, ищет дефекты на производстве или помогает врачу увидеть то, что человек мог пропустить. Где-то он просто стоит в виде дорогой лицензии, которой никто толком не пользуется. И вот это последнее встречается чаще, чем принято признавать на конференциях.

Главная ошибка начинается с вопроса: где нам применить ИИ? Вопрос звучит разумно, но он слишком рано ставит ор-

ганизацию в режим поиска инструмента. Лучше начать иначе: какую проблему мы хотим решить, почему она стоит денег, кому станет лучше, что изменится в работе людей и как мы поймем, что получилось? Только после этого имеет смысл спрашивать, нужен ли здесь ИИ. Иногда нужен. Иногда хватит нормального поиска по документам, переписанного процесса или честного разговора между отделами.

Мне нравится трезвый подход. Не потому, что ИИ слабая технология. Наоборот, он силен именно там, где от него ждут конкретной работы. Он хорошо ищет закономерности, обрабатывает большие объемы информации, помогает с черновиками, классифицирует, сравнивает, предлагает варианты. Но он плохо переносит туман в постановке задачи. Если руководитель сам не понимает, чего хочет, ИИ не превратит эту неопределенность в стратегию. Он красиво ее переформулирует. Иногда этого достаточно, чтобы всех успокоить. Для бизнеса этого мало.

В книге будет много разговоров о пользе, расходах, рисках, данных, людях, клиентах и управлении. Технических деталей будет ровно столько, сколько нужно руководителю. Не больше. Человек, отвечающий за подразделение или компанию, должен понимать, почему модель ошибается, почему ей нужны данные, почему хороший пилот может умереть при масштабировании, почему сотрудники саботируют полезные инструменты, и почему безопасное внедрение почти всегда медленнее красивой демонстрации.

ИИ не является задачей ИТ-отдела в одиночку. Это уже не та история, где технологи что-то покупают, настраивают и потом показывают бизнесу. Если ИИ меняет способ работы, значит, его должны обсуждать те, кто отвечает за людей, клиентов, риски, финансы, право и репутацию. Иначе получится знакомая картина: модель есть, процесс не изменился; лицензии куплены, привычки остались прежними; отчет о внедрении написан, пользы никто не видит.

Я буду говорить о руководителе, который хочет стать достаточно грамотным в ИИ. Не фанатом и не противником. Грамотность здесь означает четыре простые вещи. Первое: понимать основу технологии настолько, чтобы отличать реальную возможность от рекламного шума. Второе: уметь сформулировать причину внедрения. Третье: видеть организационную сторону дела, а не только экран с чат-ботом. Четвертое: помнить о последствиях, включая этику, право, безопасность и доверие клиентов.

Эти четыре вещи звучат спокойно. На практике они требуют дисциплины. Например, сказать совету директоров, что компания пока не готова к сложному ИИ-проекту, гораздо труднее, чем одобрить модный пилот. Признать, что данные разбросаны по системам, старые, неполны и плохо описаны, неприятно. Объяснить сотрудникам, что ИИ не заменит их завтра утром, но изменит ожидания к их работе, тоже непросто. Зато именно такие разговоры отделяют внедрение от имитации.

Эта книга не обещает универсального рецепта. Организации разные. Банк, больница, сеть кафе, производственный завод, университет и небольшая юридическая фирма будут применять ИИ по-разному. Но вопросы у них похожи. Что уже можно делать? Сколько это стоит? Где риски? Почему одни сотрудники получают пользу, а другие тратят больше времени? Как выбрать первые проекты? Когда не надо использовать ИИ? Как не превратить все это в очередную волну корпоративного энтузиазма, которая через год оставит после себя только папку с презентациями?

Я не предлагаю бояться ИИ. Бояться обычно начинают там, где нет ясности. Я предлагаю относиться к нему как к сильной, полезной и местами упрямой технологии. Она может дать компании преимущество. Может сэкономить время. Может улучшить сервис. Может открыть новый продукт. Но она же может уверенно ошибиться, усилить старую несправедливость, усложнить работу людей и создать репутационную проблему из мелочи, которую никто не заметил на этапе пилота.

Хорошая новость в том, что руководителю не нужно ждать идеального момента. Его не будет. Технология будет меняться, цены будут меняться, регулирование будет меняться, сотрудники будут приносить новые инструменты быстрее, чем компания успеет написать политику использования. Поэтому задача не в том, чтобы один раз выбрать правильный ответ. Задача в том, чтобы выстроить способ-

ность организации учиться, проверять, отказываться от плохих идей и масштабировать хорошие.

Начнем с самого простого вопроса, который почему-то часто пропускают. Новый ли ИИ на самом деле? Ответ важен. Если считать, что ИИ появился вчера вместе с очередным чат-ботом, легко переоценить его магию и недооценить накопленный опыт провалов. Если помнить, что эта область развивается десятилетиями, становится понятнее другое: шум новый, но многие управленческие уроки старые. Технология изменилась. Люди и организации изменились меньше, чем нам хотелось бы.

Трезвый разговор об ИИ

Я впервые увидел, как руководитель всерьез растерялся перед ИИ, не на технологической конференции, а в обычной переговорной. На столе стояли чашки с остывшим кофе, кто-то пытался подключить экран, финансовый директор листал распечатку бюджета. Вопрос был простой: покупать ли корпоративный доступ к новому генеративному инструменту для всей компании. Через десять минут стало ясно, что обсуждают не лицензию. Обсуждают страх отстать.

Маркетинг хотел быстрее делать тексты и презентации. Юристы спрашивали, куда будут уходить данные. ИТ-директор говорил о безопасности. HR видел шанс обучить сотрудников новым навыкам. Финансовый директор хотел цифры. Генеральный директор молчал дольше всех, а потом сказал: я понимаю, что это важно, но не понимаю, что именно мы покупаем. Это была самая честная фраза в комнате.

С ИИ такое происходит постоянно. Он приходит в организацию не как один понятный продукт, а как набор возможностей, угроз, слухов и чужих историй успеха. Один конкурент уже заявил, что внедрил ИИ во все процессы. Другой сократил службу поддержки. Третий открыл лабораторию искусственного интеллекта, хотя вчера еще не мог свести клиентские данные из трех систем. На этом фоне спокойное решение кажется почти слабостью. Но как раз спокойствие здесь

и нужно.

Для начала полезно убрать ощущение, что мы столкнулись с чем-то полностью новым. Да, современные системы, особенно генеративные, выглядят удивительно. Они пишут текст, создают изображения, объясняют код, переводят, резюмируют встречи, отвечают на вопросы в человеческом тоне. Но сама идея машин, которые выполняют задачи, требующие человеческого умения, появилась не вчера. У нее длинная история, и эта история важна не ради эрудиции, а ради управленческой памяти.

Еще в середине прошлого века ученые задавали вопрос, может ли машина мыслить. Потом появились первые программы, имитирующие разговор. Они были простыми, но люди уже тогда легко приписывали им понимание. Затем пришли экспертные системы. Инженеры записывали правила, по которым машина могла делать выводы, например в медицинской диагностике или технической поддержке. Некоторое время казалось, что достаточно собрать знания экспертов в огромную базу правил, и разумная машина готова.

Потом выяснилось, что реальный мир слишком плохо поддается таким правилам. Симптомы пересекаются, оборудование ведет себя странно, клиент говорит не так, как написано в инструкции, а опытный специалист часто решает задачу на основании тонких признаков, которые сам не может объяснить. У энтузиазма наступило похмелье. Деньги уходили, ожидания не оправдывались, и наступали периоды

разочарования. В истории ИИ их обычно называют зимами. Название красивое, но смысл простой: обещали больше, чем смогли сделать.

Эти зимы полезно помнить каждому руководителю. Не потому, что нынешняя волна обязательно закончится так же. Нет. Сегодня у ИИ есть то, чего не хватало раньше: огромные данные, дешевые по историческим меркам вычисления, облака, зрелая инженерия, миллионы пользователей и сильный коммерческий спрос. Но причина многих провалов осталась прежней. Люди видят впечатляющую демонстрацию и сразу переносят ее на сложную организацию, где есть устаревшие системы, сопротивление сотрудников, плохие данные и клиенты, которые не обязаны восхищаться вашей инновацией.

Представьте ресторан, который решил внедрить ИИ для обслуживания гостей. В ролике все выглядит отлично: посетитель пишет в чат, система бронирует столик, учитывает аллергию, предлагает блюдо, передает заказ на кухню. А потом начинается жизнь. У постоянного гостя две фамилии в базе. Меню меняется каждый день, но сайт обновляют по пятницам. Официанты записывают пожелания в личные мессенджеры. Кухня заменяет ингредиенты без уведомления. Клиент просит не просто столик, а тот самый у окна, где он сидел в прошлом феврале. ИИ не провалился. Провалилась фантазия, что поверх хаоса можно положить умный интерфейс и получить порядок.

Поэтому первый урок такой: ИИ усиливает зрелость орга-

низации, но плохо заменяет ее. Если процессы понятны, данные доступны, ответственность распределена, а люди знают, зачем им инструмент, технология может дать заметный эффект. Если всего этого нет, ИИ часто превращается в зеркало. Он показывает старые проблемы быстрее и ярче, чем обычная автоматизация.

Теперь надо договориться, что мы вообще называем ИИ. В бизнесе этим словом часто называют все подряд: от простого правила в CRM до большой языковой модели. Такое расширение удобно для продавцов, но мешает управлению. Для руководителя достаточно различать две большие семьи. Первая работает на правилах. Вторая учится на данных. В реальных системах они часто смешиваются, но разделение помогает думать трезво.

Правила понятны. Если клиент не оплатил счет семь дней, отправить напоминание. Если сумма операции выше обычной и страна входа изменилась, отправить на проверку. Если датчик показывает температуру выше порога, остановить линию. Это не всегда называют ИИ, но сложные системы правил действительно могут вести себя достаточно умно. Они просматривают тысячи условий, выбирают подходящие, иногда оценивают вероятность и предлагают решение.

Плюс правил в том, что их можно объяснить. Минус в том, что их надо написать. А написать можно только то, что вы понимаете. Если эксперт не может сформулировать, почему он считает эту заявку подозрительной, правило будет грубым.

Если ситуаций слишком много, база правил разрастается и начинает спорить сама с собой. В какой-то момент поддерживать ее становится тяжелее, чем решать задачу вручную.

Системы, которые учатся на данных, работают иначе. Им не объясняют каждую возможность. Им показывают много примеров. Вот письма клиентов и правильные ответы. Вот фотографии деталей и отметки, где есть дефект. Вот история покупок и факт, ушел клиент или остался. Модель ищет закономерности и потом применяет их к новой ситуации. Когда данных много и они хороши, такой подход бывает очень сильным.

Но слово 'хороши' здесь не украшение. Данные должны отражать задачу. Если вы учите модель на старых решениях, в которых уже были ошибки и предвзятость, она может аккуратно воспроизвести эти ошибки. Если данных мало, она будет угадывать. Если данные собраны в разных форматах, с пропусками и странными кодами, то сначала придется заниматься скучной работой. Чистить, согласовывать, описывать, проверять права на использование. Это не похоже на футуристическую презентацию. Это больше похоже на ремонт склада, где десять лет складывали коробки без подписей.

Генеративный ИИ сделал разговор сложнее, потому что он выглядит слишком человеческим. Обычная модель прогнозирования спроса не пытается быть обаятельной. Она выдает числа. А чат-бот отвечает вежливо, шутит, признает

ошибки, объясняет сложное простыми словами. Руководитель видит текст и невольно думает: раз он так говорит, значит, он понимает. Это опасная привычка. Текст может быть гладким, а основание слабым.

Большая языковая модель, если говорить грубо, учится предсказывать продолжение текста. Она видела огромные массивы документов и по ним научилась подбирать вероятные слова, фразы, ответы, стили. Это потрясающее инженерное достижение. Но оно не отменяет простой факт: уверенный ответ не равен истинному ответу. Модель может ошибиться спокойно и убедительно. Для черновика письма это терпимо. Для юридической позиции, медицинского совета, кредитного решения или публичного заявления компании это уже другая история.

Один из самых неприятных управленческих вопросов звучит так: где ошибка ИИ будет стоить нам дороже, чем ошибка человека? Не только в деньгах. В доверии, в репутации, в отношениях с регулятором, в безопасности. Люди тоже ошибаются, конечно. Устают, торопятся, забывают, защищают любимую гипотезу. Но общество часто иначе воспринимает машинную ошибку. Когда сотрудник службы поддержки дал неверную скидку, это конфликт с сотрудником. Когда чат-бот компании пообещал клиенту то, чего компания не собиралась выполнять, это уже разговор о контроле над собственными системами.

Отсюда не следует, что ИИ надо держать под замком. Сле-

дует другое: у каждого применения должен быть уровень риска. В одном случае можно разрешить сотрудникам использовать инструмент свободно, потому что последствия малы. В другом нужен человек, который проверяет результат. В третьем ИИ вообще не должен принимать решение, а может только готовить информацию. В четвертом лучше пока не трогать задачу, потому что данных нет, ответственности нет, а желание показать прогресс слишком велико.

Я часто слышу от руководителей вопрос: с чего начать? Самый плохой ответ - начать с самого красивого. Красивые проекты любят презентации. Виртуальный помощник для клиентов, умная аналитика совета директоров, автоматический стратегический советник, цифровой сотрудник на все случаи жизни. Все это может быть полезно, но для первого шага часто слишком сложно. Начинать лучше там, где задача понятна, данные доступны, риск ограничен, а результат можно измерить без гадания.

Например, у компании есть сотни внутренних документов, и сотрудники тратят время, чтобы найти правила командировок, закупок или оформления отпусков. Здесь можно сделать помощника по базе знаний. Он не должен увольнять юристов или менять стратегию. Он должен быстрее находить ответы и ссылаться на источник. Такой проект учит организацию многому: как готовить документы, как ограничивать доступ, как проверять ответы, как обучать сотрудников задавать вопросы, как собирать обратную связь.

Другой пример: отдел продаж пишет типовые письма клиентам после встреч. ИИ может готовить черновик на основе заметок менеджера и карточки клиента. Но черновик должен оставаться черновиком. Менеджер отвечает за тон, обещания и факты. Польза здесь проста: меньше пустого времени на первый вариант текста. Риск управляемый, если запретить отправку без проверки и не подсовывать модели лишние персональные данные.

Есть и более сложные случаи. Прогнозирование спроса, динамическое ценообразование, выявление мошенничества, диагностика дефектов, персональные рекомендации. Там ИИ может дать серьезные деньги. Но там же растут требования. Нужны специалисты, экспериментальный дизайн, контроль качества данных, мониторинг модели после запуска, понятные правила остановки. Модель, которая хорошо работала весной, может начать ошибаться осенью, потому что изменился рынок, поведение клиентов или ассортимент. ИИ не поставили один раз и забыли. Его эксплуатируют.

Слово 'эксплуатация' звучит менее вдохновляюще, чем 'трансформация', зато ближе к реальности. Любая рабочая система требует владельца. Кто смотрит на показатели? Кто принимает решение о переобучении модели? Кто отвечает, если клиент пожаловался? Кто имеет право отключить инструмент? Кто проверяет, что данные используются законно? Пока эти вопросы не имеют ответов, проект остается экспериментом, даже если он уже встроен в процесс.

Отдельно стоит сказать о производительности. Это самое популярное обещание ИИ. Сотрудники будут писать быстрее, искать быстрее, анализировать быстрее. Иногда так и происходит. Особенно если задача повторяемая, критерий качества понятен, а ИИ помогает человеку, не заменяя его с первого дня. Новички часто получают большую пользу, потому что инструмент подсказывает им практики опытных коллег. Но у производительности есть неприятная сторона: выигрыш в одном месте может создать работу в другом.

Допустим, ИИ стал писать больше маркетинговых вариантов. Теперь юридический отдел проверяет больше материалов. Или модель ускорила подготовку отчетов, но руководители тратят время на выяснение, какие цифры взяты из системы, а какие красиво домыслены. Или программист быстрее создает код, а тестирующий дольше ищет ошибки. Это не аргумент против ИИ. Это аргумент против ленивой арифметики, где экономию считают там, где удобно, а новые расходы прячут в чужом календаре.

Если вы хотите оценивать продуктивность, сначала уточните уровень. Индивидуальный сотрудник может реально сэкономить час в день. Отдел может не сэкономить ничего, если процесс не меняется. Организация может даже замедлиться, если каждый начинает использовать разные инструменты, хранить данные в разных местах и спорить о качестве результатов. Национальная экономика тем более не меняется от того, что у всех появились красивые ассистенты. Меж-

ду личной пользой и большим эффектом лежит управление. Теперь о деньгах. ИИ часто продают как путь к снижению затрат. В отдельных задачах это правда. Автоматическая обработка обращений, предиктивное обслуживание оборудования, оптимизация маршрутов, поддержка операторов, анализ документов - все это может снизить расходы. Но внедрение тоже стоит денег. Лицензии, облачные вычисления, интеграция, подготовка данных, обучение, безопасность, юристы, консультанты, внутренние команды, изменение процессов. Иногда главная стоимость спрятана не в счете поставщика, а во времени людей, которые должны привести организацию в состояние, пригодное для ИИ.

У генеративного ИИ есть еще одна особенность: он кажется дешевым, пока им пользуется один человек. Подписка стоит как несколько обедов, и руководителю кажется, что масштабирование будет таким же простым. Но корпоративное использование быстро поднимает другие вопросы. Какие данные можно вводить? Что будет храниться у поставщика? Нужно ли подключать внутренние базы? Сколько будут стоить запросы через API при большом потоке клиентов? Кто поддерживает приложение? Как ограничить расходы, если нагрузка выросла? Как выйти из зависимости от одного поставщика?

Хороший финансовый директор в такой ситуации не портит праздник. Он задает правильные вопросы. Какой бизнес-показатель изменится? Где базовая линия? Сколько сто-

ит текущий процесс? Какие расходы одноразовые, какие постоянные? Что будет, если точность модели окажется ниже ожиданий? Сколько стоит ручная проверка? Какой план остановки? Если ответов нет, бюджет на ИИ превращается в ставку на настроение рынка.

Польза для клиента требует отдельного разговора. Многие компании внедряют ИИ так, будто клиент должен быть благодарен уже за сам факт автоматизации. Но клиенту все равно, насколько современный у вас стек. Он хочет решить свою задачу. Быстрее получить ответ. Не повторять одно и то же трем каналам поддержки. Не ждать сорок минут из-за вопроса, который можно было решить письмом. Не разговаривать с ботом, который делает вид, что понял, а потом отправляет в тот же раздел сайта.

ИИ может дать клиенту ценность, если начинать с его проблемы. Не с того, что мы можем автоматизировать, а с того, где человек теряет время, деньги, уверенность или ощущение контроля. В страховании это может быть понятное объяснение полиса и быстрый разбор случая. В образовании - персональный темп и обратная связь. В медицине - напоминания, сортировка симптомов, поддержка врача, а не замена доверия. В рознице - рекомендации, которые не выглядят как навязчивая распродажа остатков.

Иногда лучшая клиентская ценность вообще не видна клиенту напрямую. ИИ помогает точнее планировать запасы, и нужный товар есть на полке. Он прогнозирует полом-

ку оборудования, и сервис не падает. Он ищет мошенничество, и честные клиенты меньше страдают от лишних проверок. Но и здесь нельзя забывать о восприятии. Люди обычно не против удобства. Они против ощущения, что их данные используют без спроса, а решение принимает непрозрачная машина, с которой невозможно спорить.

Поэтому вопрос 'можем ли мы?' всегда должен идти рядом с вопросом 'стоит ли нам?'. Закон задает минимальную границу. Репутация живет выше этой границы. Можно формально иметь право использовать данные клиента для персонализации, но сделать это так навязчиво, что клиент почувствует слежку. Можно автоматизировать отбор кандидатов и не заметить, что модель повторяет старые перекосы. Можно поставить распознавание лиц ради безопасности и получить общественный конфликт, потому что людей не спросили и плохо объяснили цель.

Руководителю не нужно самому писать этические кодексы на пятьдесят страниц. Но он должен добиться, чтобы в компании был простой механизм: кто оценивает риск, какие применения запрещены, где нужен человек в контуре, как клиент может оспорить решение, что делаем при инциденте, кто говорит с прессой, какие данные нельзя использовать даже ради красивого результата. Без такого механизма этика остается плакатом на сайте.

Есть еще один слой, о котором многие вспоминают поздно: сотрудники. ИИ меняет не только процессы, но и чув-

ство профессиональной ценности. Опытный человек может воспринять инструмент как помощника, а может как угрозу. Новичок может научиться быстрее, а может перестать разбираться в основе работы, потому что всегда есть подсказка. Менеджер может использовать ИИ, чтобы снять рутину, а может начать требовать больше задач за то же время, не понимая, что проверка результатов тоже работа.

Сотрудников нельзя просто поставить перед фактом. Им нужно объяснить, где ИИ помогает, где запрещен, как проверять результат, что будет оцениваться, какие навыки станут важнее. Очень полезно честно сказать: да, часть задач исчезнет; да, появятся новые; да, нам придется учиться; нет, мы не знаем всего заранее. Люди обычно лучше переносят неприятную ясность, чем бодрые лозунги.

Внутри компании быстро появятся три группы. Первые уже используют ИИ и никого не ждут. Они найдут инструменты сами, иногда нарушая правила просто потому, что правил нет. Вторые осторожно пробуют и ждут разрешения. Третьи сопротивляются, иногда разумно, иногда из страха. Задача руководителя не в том, чтобы всех загнать в одинаковый энтузиазм. Задача в том, чтобы энергию первых сделать безопасной, вторым дать обучение, а возражения третьих разобрать без насмешки. Среди скептиков часто сидят люди, которые первыми заметят будущую проблему.

Когда организация выбирает первые проекты, полезно разложить их по простым уровням. Самый легкий уровень -

личная помощь сотруднику: черновик текста, резюме встречи, перевод, поиск идей. Второй - помощник внутри процесса, например бот по внутренним политикам или система подготовки ответов клиентам. Третий - специализированная модель, которая влияет на деньги, риск или качество, например прогноз спроса или выявление брака. Четвертый - сложные автономные решения, где ошибка может быть дорогой и где нужна серьезная инженерная зрелость.

Новички часто прыгают сразу на третий или четвертый уровень, потому что там больше славы. Но если организация не умеет управлять первым и вторым, сложный проект быстро покажет слабые места. Это как учиться водить на грузовике в гололед. Возможно, получится. Но лучше сначала понять, где педали, кто отвечает за маршрут и почему тормозной путь не обсуждают после аварии.

Есть ситуации, когда ИИ лучше не применять. Если задача редкая и плохо описанная, данных мало, а ошибка дорогая, стоит остановиться. Если руководитель хочет ИИ только потому, что совет директоров спросил о нем на прошлой встрече, стоит остановиться. Если клиенту станет хуже, но внутренний отчет покажет экономию, стоит остановиться. Если никто не может объяснить, кто владелец результата, стоит остановиться. Остановка не означает отказ навсегда. Это нормальная часть управления технологией.

Трезвый руководитель отличается не тем, что говорит 'нет' чаще других. Он умеет говорить 'пока нет' и 'вот при

каких условиях да'. Например: мы вернемся к этой идее, когда приведем данные в порядок; когда определим метрику качества; когда найдем ответственного; когда проведем тест с реальными пользователями; когда юристы и служба безопасности согласуют границы; когда поймем, как отключить систему без остановки бизнеса.

Теперь можно сформулировать практический минимум для начала. Выберите несколько задач, где боль понятна. Опишите текущий процесс: кто делает, сколько времени, где ошибки, сколько стоит задержка. Проверьте данные. Оцените риск. Назначьте владельца. Сделайте небольшой тест, но не называйте его успехом только потому, что он понравился участникам. Сравните с базовой линией. Посмотрите на скрытую работу, которую создал инструмент. Спросите сотрудников и клиентов, стало ли лучше. Потом решайте, масштабировать или закрывать.

Закрывать плохие проекты надо быстро и без стыда. В ИИ это особенно важно, потому что вокруг технологии много самолюбия. Никто не хочет признать, что модный пилот не дал пользы. Но зрелая организация не обязана защищать каждую идею. Она обязана учиться. Иногда лучший результат пилота - доказательство, что задачу надо решать иначе. Это тоже экономия, если вывод сделан рано.

С масштабированием свои трудности. Рабочий прототип часто держится на энтузиазме небольшой команды. В большой организации появляются права доступа, обучение, под-

держка, интеграция, регламенты, исключения, локальные особенности, закупки, аудит. То, что работало на двадцати пользователях, начинает скрипеть на двух тысячах. Поэтому масштабирование надо проектировать с самого начала. Не полностью, но достаточно, чтобы не строить демонстрацию, которую потом придется выбросить.

Еще один совет: не отдавайте язык ИИ только технологам или продавцам. Слова формируют ожидания. Если вы называете систему 'умным сотрудником', люди будут ждать самостоятельности и понимания. Если называете ее 'помощником для черновигов', ожидания станут точнее. Если говорите 'автоматизация решения', риски одни. Если говорите 'поддержка решения человека', риски другие. Внедрение начинается не с кода, а с честного описания того, что система делает и чего не делает.

Руководителю полезно самому пользоваться ИИ. Не поручать все ассистенту, не ждать отчета от цифровой команды, а открыть инструмент и попробовать. Попросить резюмировать длинный документ. Сравнить варианты письма. Задать вопрос по теме, в которой вы разбираетесь, и увидеть, где ответ хорош, а где он уверенно ошибается. Личный опыт быстро лечит две крайности: мистическое восхищение и ленивое презрение.

Но личный опыт не должен подменять корпоративное решение. То, что помогло вам подготовиться к встрече, не означает, что систему можно выпускать к клиентам. То, что

модель красиво ответила на десять вопросов, не означает, что она выдержит десять тысяч разных клиентов, плохие формулировки, злонамеренные запросы, устаревшие документы и юридические ограничения. Между 'работает у меня' и 'работает в бизнесе' лежит большая дистанция.

ИИИ сейчас переживает жаркое лето. Каждый день появляются новые инструменты, рейтинги, заявления, прогнозы. Возможно, часть нынешнего восторга спадет. Возможно, какие-то компании разочаруются, бюджеты сократятся, громкие проекты закроются. Это не будет означать, что ИИИ исчез. После каждой волны шума остается то, что действительно полезно. Так было с предыдущими технологиями, так будет и здесь.

Лучшее, что может сделать руководитель в такой период, - не играть роль пророка. Не обещать, что ИИИ заменит половину работы. Не уверять, что это всего лишь игрушка. Не строить стратегию на чужих пресс-релизах. Надо строить способность задавать скучные вопросы и доводить полезные решения до реальной работы. В бизнесе это часто важнее дальновидности. Дальновидность без исполнения быстро превращается в красивую речь.

Есть еще практическая деталь о пилотах, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых

слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о данных, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не

на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о ответственности, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о обучении, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о клиентском опыте, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром?

Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о стоимости, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не по-

нимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о риске, которую редко произносят в начале проекта. На старте всем хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

Есть еще практическая деталь о масштабировании, которую редко произносят в начале проекта. На старте всем

хочется говорить о будущем состоянии: как станет быстро, удобно и современно. Но проект живет в настоящем, где люди заняты, системы не готовы, а старые договоренности сильнее новых слайдов. Поэтому перед каждым решением полезно спросить: что должно измениться в понедельник утром? Кто откроет другой экран, кто перестанет делать старое действие, кто проверит результат, кто объяснит клиенту новый порядок? Если ответа нет, идея пока не стала рабочей.

Этот вопрос кажется мелким, но именно на мелочах ломаются большие внедрения. Не на слове 'искусственный', не на слове 'интеллект', а на том, что сотрудник не знает, можно ли вставить в систему письмо клиента; руководитель не понимает, какую метрику смотреть; служба безопасности подключается за день до запуска; обучение проводят один раз, когда половина команды в отпуске. И потом все удивляются, почему технология, которая на демонстрации выглядела убедительно, в жизни дает скромный результат.

В конце этой первой главы я хочу оставить не вывод, а рабочую установку. Относитесь к ИИ как к сильному младшему коллеге, который много читал, быстро пишет, не устает, иногда блестяще помогает, но может придумать уверенную ерунду и не понимает вашего бизнеса так, как его понимают люди внутри. Такого коллегу не надо ни увольнять с порога, ни ставить генеральным директором. Ему дают задачу, контекст, ограничения, проверку и ответственность взрослого человека рядом.

Если организация научится делать это спокойно, она получит пользу. Не сразу во всем. Не без ошибок. Но получит. А если она начнет с лозунгов, закупок и страха отстать, ИИ просто ускорит старую привычку делать вид, что движение равно прогрессу. В следующей части этой книги мы будем разбирать, из чего складывается такая спокойная работа: как понимать технологию, как выбирать причины для внедрения, как считать пользу и почему иногда самый разумный проект начинается не с модели, а с уборки в собственных данных.

Это началось раньше, чем кажется

Я часто слышу от руководителей одну и ту же фразу: «Ну теперь-то всё изменилось. Теперь появился искусственный интеллект».

В этой фразе есть правда. И есть ошибка.

Правда в том, что последние годы действительно изменили ощущение от технологии. Раньше искусственный интеллект жил где-то за стеной. В лабораториях, в отделах аналитики, в банках, в поисковых системах, в рекомендациях интернет-магазинов, в навигации, в антифроде. Он работал, но не здоровался. Не спорил. Не писал письмо поставщику. Не сочинял план совещания. Не объяснял бухгалтеру, почему формула в таблице не сходится. Не отвечал человеческим голосом на человеческий вопрос.

А потом люди открыли чат-бота, написали туда обычную фразу и получили обычный ответ. Не команду. Не код. Не набор пунктов из справки. Ответ.

Вот почему многим показалось, что искусственный интеллект родился вчера.

Ошибка в том, что он не родился вчера. Он просто вышел из подсобки в зал.

Это полезно понимать, потому что отношение к новой технологии сильно зависит от того, какой возраст мы ей приписываем. Если нам кажется, что перед нами свежая игруш-

ка, мы ведём себя одним образом. Кто-то бросается попробовать всё подряд. Кто-то ждёт, пока игрушка сломается. Кто-то нервничает: слишком быстро, слишком непонятно, слишком шумно. Но если мы видим, что за этой игрушкой стоят десятилетия попыток, ошибок, денег, разочарований и инженерного упрямства, разговор становится спокойнее.

Искусственный интеллект не прилетел к бизнесу на блестящей тарелке. Он шёл к нему долго. Иногда бодро, иногда на костылях.

Мне нравится начинать этот разговор не с серверов, не с нейросетей и не с графиков инвестиций, а с простой сцены. Представьте человека, который впервые говорит с машиной не как с калькулятором, а как с собеседником. Он не нажимает набор кнопок. Не вводит команду на странном языке. Он спрашивает: «Что мне делать?» И машина отвечает так, что на секунду хочется забыть, что это машина.

У многих такая сцена случилась в 2022 или 2023 году. У кого-то в офисе, когда коллега показал чат-бота. У кого-то дома, когда ребёнок попросил нейросеть объяснить дробь. У кого-то ночью, перед дедлайном, когда нужно было быстро привести в порядок текст, который уже не было сил перечитывать.

Но сама мечта говорить с машиной обычным языком намного старше.

Фантасты давно понимали, что настоящий прорыв будет не тогда, когда компьютер научится считать быстрее чело-

века. Считать он давно умеет быстрее. Прорыв будет тогда, когда с ним можно будет говорить без переводчика. Когда между человеком и машиной исчезнет неприятная прослойка из команд, меню, инструкций и ошибок ввода.

В старых фильмах и сериалах корабельный компьютер слушал капитана, уточнял детали, давал справку, запускал системы. Сейчас это выглядит почти смешно: люди летят через галактику, но всё равно формулируют запрос короткими приказами, будто боятся, что машина не поймёт лишнее слово. В этом есть своя честность. Авторы угадывали желание, но ещё не верили в свободный разговор.

И долго были правы.

Машины десятилетиями плохо понимали человеческую речь. Они могли выполнять команды, если эти команды были заранее предусмотрены. Могли отвечать по шаблону. Могли изображать беседу. Но между «изображать беседу» и «быть полезным собеседником» лежит большая дистанция. В бизнесе эта дистанция особенно заметна. Игрушечный диалог может впечатлить на демонстрации. Через неделю он начинает раздражать, если не помогает решать настоящие задачи.

Поэтому первый урок истории искусственного интеллекта для руководителя звучит просто: не путайте впечатление от демо с зрелостью технологии.

Демо всегда опережает реальность. Так было почти всегда.

Ещё в середине прошлого века один из главных вопросов звучал дерзко: может ли машина думать? Вопрос был не технический, а почти философский. Если машина отвечает так, что мы не можем отличить её от человека, что тогда считать мышлением? Нам обязательно нужно заглянуть ей внутрь, или достаточно поведения? Если сотрудник даёт хороший анализ рынка, мы же не вскрываем ему голову, чтобы проверить, как он пришёл к выводу. Но с машиной нам хочется именно этого. Нам нужно знать, что там не фокус.

Эта напряжённость никуда не исчезла. Сегодня руководитель смотрит на результат работы модели и думает: «Вроде хорошо. Но можно ли этому доверять?» И это правильный вопрос. Только он не новый. Он сопровождает искусственный интеллект с самого начала.

Ранние программы, которые имитировали беседу, были устроены довольно просто. Они подхватывали слова человека, переставляли их, задавали встречные вопросы. Пользователь писал: «Мне трудно разговаривать с начальником». Машина отвечала чем-нибудь вроде: «Почему вам трудно разговаривать с начальником?» Ничего глубокого. Никакого понимания в человеческом смысле. Но некоторым людям этого хватало, чтобы почувствовать контакт.

Это важный и немного неприятный факт. Человеку не всегда нужна настоящая разумность, чтобы начать относиться к системе как к разумной. Иногда достаточно правильно-го ритма ответа.

Для бизнеса это не мелочь. Если клиент разговаривает с ботом, он может приписать боту больше понимания, чем там есть. Если сотрудник получает уверенный ответ, он может решить, что система «знает». Если совет директоров видит красивую презентацию, созданную с помощью ИИ, он может подумать, что за ней стоит такая же красивая аналитика. А там может стоять только ловкая сборка слов.

История искусственного интеллекта учит не цинизму, а трезвости. Машины давно умеют производить впечатление. Полезность надо проверять отдельно.

Сам термин «искусственный интеллект» появился не как маркетинговый слоган в презентации стартапа, а как название исследовательского направления. Учёные пытались понять, можно ли сделать машины, которые решают задачи, требующие рассуждения, поиска, распознавания, планирования. Тогда казалось, что если разбить мышление на правила, всё получится довольно быстро.

Это вообще типичная ошибка умных людей: они недооценивают сложность того, что сами делают легко.

Человек смотрит на комнату и понимает, где стол, где дверь, где чашка, где злой кот, которого лучше не трогать. Ребёнок учится этому без учебника. А машина долго спотыкалась о такие задачи. Зато она могла быстро перебирать варианты в игре, решать формальные головоломки, доказывать теоремы в узких областях. Получался странный перекося. То, что человеку кажется умным, машине иногда давалось про-

ще. То, что человеку кажется бытовой ерундой, машине давалось мучительно трудно.

В офисной жизни этот перекося тоже есть. Система может быстро подготовить черновик сложного письма, но неправильно понять простую внутреннюю договорённость. Может найти закономерности в тысячах обращений клиентов, но не уловить, что один конкретный клиент сейчас злится не из-за тарифа, а из-за того, что ему три раза обещали перезвонить и не перезвонили. Может предложить десять вариантов стратегии, но не знать, что два ключевых отдела не разговаривают друг с другом после прошлогоднего конфликта.

Интеллект машины не похож на человеческий. Это не хуже и не лучше. Это надо помнить.

В первые десятилетия искусственного интеллекта было много ожиданий. Казалось, ещё немного, ещё несколько лет, ещё немного финансирования, и машины начнут переводить языки, ставить диагнозы, вести диалог, управлять сложными системами. Часть задач действительно двигалась. Появлялись программы для игр, ранние роботы, системы перевода, математические решатели.

Но потом выяснялось, что аккуратный пример в лаборатории не выдерживает грязного мира.

В лаборатории предложения чистые. В жизни люди говорят с ошибками, намёками, раздражением и сокращениями. В лаборатории данные подготовлены. В жизни данные лежат в десяти системах, часть устарела, часть введена вручную,

часть вообще хранится в голове у сотрудницы, которая работает в компании с 1998 года и одна знает, почему этот клиент называется в базе двумя разными именами. В лаборатории задача сформулирована. В жизни руководитель сам не всегда понимает, какую задачу хочет решить.

Вот почему в истории ИИ были не только подъёмы, но и резкие похолодания.

Сначала все обещают скорую революцию. Потом выясняется, что революция не пришла по расписанию. Деньги заканчиваются. Отчёты становятся мрачнее. Исследователи осторожнее. Журналисты переключаются на другую тему. Наступает зима.

Эти зимы полезны. Неприятны, но полезны.

Когда вокруг меньше шума, остаются те, кто действительно строит работающие вещи. Не обязательно эффектные. Иногда скучные. Системы планирования. Распознавание образов. Поиск. Оптимизация маршрутов. Кредитный скоринг. Обнаружение мошенничества. Рекомендации. Управление запасами. Всё это не всегда называли громко искусственным интеллектом, но именно там технология выросла.

Бизнес часто не замечал ИИ не потому, что его не было. Он был встроен в процессы так, что перестал выглядеть как чудо.

Это хороший признак зрелой технологии. Лампочка не кажется магией. Лифт не кажется роботом. Поисковая стро-

ка не кажется искусственным интеллектом, хотя за ней давно стоят сложные модели. Когда технология становится полезной, она часто становится невидимой.

С генеративным ИИ произошло обратное. Он снова сделал технологию видимой. Причём слишком видимой. Люди начали показывать друг другу картинки, стихи, письма, ответы на экзамены, планы отпусков, куски кода. Впечатление было сильным, потому что система работала в человеческой зоне: язык, образ, объяснение, просьба.

Раньше ИИ помогал где-то внутри машины. Теперь он сел рядом за стол.

Для руководителя это меняет многое. Когда ИИ спрятан в антифрод-системе, его обсуждают специалисты. Когда ИИ появляется в руках каждого сотрудника, это уже вопрос управления. Кто может пользоваться? Для чего? С какими данными? Кто отвечает за ошибку? Как меняется работа новичков? Как меняется работа экспертов? Что считать хорошим результатом? Что делать, если сотрудники тихо используют внешние инструменты, потому что так быстрее, а политика компании ещё не написана?

И тут история снова помогает.

Каждый большой подъём ИИ сопровождался одним и тем же соблазном: принять быстрый прогресс за прямую линию в будущее. Сегодня модель отвечает лучше, чем год назад. Значит, через год она будет делать всё. Сегодня она пишет код. Значит, программисты исчезнут. Сегодня она подбира-

ет аргументы. Значит, стратегия станет автоматической. Сегодня она проходит тесты. Значит, можно доверить ей решения.

Так думать удобно. Особенно если нужно продать проект, выбить бюджет или напугать конкурентов. Но история технологий редко движется прямой линией. Она идёт рывками, откатами, обходными тропами. Где-то прогресс оказывается быстрее ожиданий. Где-то упирается в стену, которую никто не заметил в презентации.

В искусственном интеллекте таких стен много.

Одна стена называется данными. Модель может быть мощной, но если ваши данные грязные, разрозненные или юридически непригодные для использования, магии не будет. Другая стена называется процессами. Если процесс сам по себе бессмысленный, ИИ ускорит бессмыслицу. Третья стена называется ответственностью. Машина может предложить решение, но кто подпишется под ним? Четвёртая стена называется доверием. Сотрудники могут саботировать даже хороший инструмент, если чувствуют, что его внедряют против них. Клиенты могут отказаться от удобного сервиса, если видят в нём слежку или равнодушие.

История ИИ не говорит: «Не верьте технологии». Она говорит: «Не верьте гладким траекториям».

Есть ещё один урок, менее очевидный. Старые подходы не исчезают полностью. Они возвращаются в новых сочетаниях.

Когда-то многие системы строились на правилах. Эксперты формулировали: если у пациента такие симптомы, возможен такой диагноз; если в заявке такие признаки, нужно проверить риск; если клиент сделал то-то, предложить ему это. Потом пришло машинное обучение и сказало: не надо вручную писать все правила, дайте данные, модель найдёт закономерности. Это был сильный сдвиг. Но правила не умерли.

В бизнесе правила нужны. Законодательство, политика безопасности, кредитные лимиты, регламенты, профессиональные стандарты, этические запреты. Нельзя просто сказать модели: «Разберись как-нибудь». Иногда нужно жёстко задать границы. Не предлагать клиенту продукт, если он ему противопоказан. Не использовать персональные данные без основания. Не принимать автоматическое решение там, где требуется человек. Не выдавать уверенный ответ, если уверенности нет.

Будущее, скорее всего, не будет чисто «правильным» или чисто «данным». Оно будет смешанным. Данные помогут видеть закономерности, правила будут задавать рамки, люди будут отвечать за смысл и последствия. На бумаге это звучит аккуратно. На практике потребует много скучной управленческой работы.

А ИИ любит, когда за него сделали скучную работу.

Я иногда прошу руководителей представить не блестящую лабораторию, а обычную среду своей организации. Почтовые цепочки на двадцать семь писем. Файлы с названиями

«финал_новый_точно_последний». Клиентские записи, заполненные по-разному в разных регионах. Отделы, которые используют одни и те же слова в разных смыслах. Старые системы, к которым никто не хочет прикасаться, потому что они «как-то работают». Новые инициативы, которые накладываются на старые и создают ещё один слой отчётности.

И вот в эту среду приходит искусственный интеллект.

Он не приходит в пустую комнату. Он приходит в вашу комнату. Со всеми её проводами, привычками, конфликтами, страхами, обходными путями и табличками на стене.

Поэтому вопрос «новый ли ИИ?» менее важен, чем вопрос «готовы ли мы не повторять старые ошибки?»

Старые ошибки хорошо известны.

Первая: ждать слишком много слишком быстро. Это ведёт к разочарованию. Руководитель объявляет большую программу, сотрудники видят несколько сырых инструментов, через полгода энтузиазм превращается в усталость. Потом появляется фраза: «Мы уже пробовали ИИ, не сработало». Хотя на самом деле пробовали не ИИ, а плохо выбранный пилот без ясной цели.

Вторая: недооценивать технологию после первых неудач. Система ошиблась в смешном месте, и все решили, что она бесполезна. Но смешная ошибка в одном типе задач не означает бесполезность в другом. Человек тоже может плохо петь и хорошо вести переговоры. Инструменты надо оценивать по задаче, а не по общему впечатлению.

Третья: отдавать тему только техническим специалистам. Технические специалисты необходимы, но они не могут в одиночку решить, где ИИ принесёт пользу бизнесу, какой риск допустим, как изменится работа людей, что будет считаться успехом. Если руководитель прячется за фразой «это пусть айтишники разбираются», он уже проигрывает часть управления.

Четвёртая: превращать ИИ в символ современности. Компания покупает модный инструмент, потому что так делают другие. Появляется внутренний портал, несколько обучающих сессий, красивые слайды. Вопрос «зачем?» остаётся без ответа. Через год инструмент либо тихо умирает, либо живёт как игрушка для энтузиастов.

Пятая: забывать, что каждая волна ИИ приносит не только возможности, но и новую ответственность. Чем естественнее машина разговаривает, тем легче человеку забыть, что она может ошибаться. Чем дешевле создавать контент, тем выше риск мусора. Чем проще автоматизировать решение, тем важнее понять, кого это решение затронет.

В истории ИИ есть много поводов для оптимизма. Машины научились делать то, что раньше казалось далёкой фантастикой. Они помогают искать лекарства, управлять сложными сетями, переводить языки, распознавать изображения, писать код, анализировать документы, поддерживать клиентов. Они уже стали частью экономики, даже если мы не всегда называем их по имени.

Но в этой же истории есть много поводов для осторожности. Обещания часто опережали результат. Уверенность часто оказывалась преждевременной. Системы, хорошо работающие в узком примере, ломались при столкновении с реальностью. Деньги тратились быстрее, чем появлялась польза. Люди сопротивлялись не из глупости, а потому что им предлагали не помощь, а очередной плохо объяснённый контроль.

Мне кажется, хороший руководитель должен держать в голове обе стороны.

Без оптимизма вы будете слишком медленны. Без осторожности слишком дороги.

Сейчас многие организации находятся в странном положении. Они уже не могут игнорировать ИИ, но ещё не понимают, как встроить его в нормальную работу. Это похоже на момент, когда в офисах появился интернет. Сначала он был отдельной штукой: кто-то умел пользоваться, кто-то боялся, кто-то печатал адреса сайтов на визитках как знак прогресса. Потом интернет стал инфраструктурой. Никто не говорит: «У нас интернет-стратегия для отправки писем». Он просто есть в каждом процессе.

С ИИ может произойти нечто похожее. Не сразу. И не само.

Пока технология шумит, её трудно рассмотреть. Одни кричат, что она заменит всех. Другие говорят, что пузырь лопнет. Третьи продают курсы, платформы, консультации и

чудеса. В такой обстановке полезно сделать шаг назад и посмотреть на длинную историю. В ней видно, что ИИ переживал восторги и зимы, провалы и спокойные инженерные победы, модные названия и тихое внедрение в повседневные системы.

Это снижает градус.

Если искусственный интеллект не новорождённый, то не надо носиться вокруг него с криками. Если он не всемогущий, то не надо падать перед ним на колени. Если он не игрушка, то не надо отдавать его на самотёк. Если он не магия, то можно управлять им как технологией, хоть и сложной.

Для начала руководителю стоит задать себе несколько простых вопросов.

Какие задачи в нашей организации уже решаются с помощью ИИ, даже если мы так их не называем? Где у нас есть автоматические рекомендации, скоринг, прогнозы, распознавание, классификация, поиск аномалий? Кто отвечает за эти системы? Как мы проверяем их работу? Какие решения они поддерживают? Где человек остаётся в контуре, а где только делает вид?

Потом другой набор вопросов. Что изменилось именно с появлением генеративных моделей? Какие сотрудники уже используют их в работе? Для каких задач? С какими данными? Что они делают лучше? Где тратят больше времени из-за проверки результата? Где появляются риски, о которых раньше не думали?

И ещё один вопрос, самый неприятный: где мы сейчас верим красивому ответу больше, чем проверенному процессу?

История ИИ показывает, что красивая форма ответа обладает силой. Человеку трудно сопротивляться уверенной речи. Особенно если она экономит время. Особенно если он устал. Особенно если начальство ждёт результата вчера. Поэтому зрелая организация должна учиться не только пользоваться ИИ, но и сомневаться в нём без истерики.

Сомнение не тормозит внедрение. Оно делает внедрение взрослым.

Взрослый подход начинается с признания: ИИ не новый гость, а давний сосед, который вдруг стал разговорчивым. Раньше он помогал незаметно. Теперь он просит место за столом. Его стоит выслушать. Иногда он скажет полезное. Иногда глупость. Иногда опасную глупость уверенным тоном. Иногда увидит закономерность, которую люди пропустили. Иногда сэкономит часы. Иногда создаст новую работу.

Наша задача не угадать, станет ли очередная волна ИИ вечным летом или снова принесёт холод. Наша задача проще и труднее: строить решения так, чтобы они приносили пользу при любой погоде.

Если завтра ажиотаж спадёт, останутся хорошие данные, понятные процессы, обученные люди, ясные правила и несколько работающих применений. Это уже неплохо. Если развитие ускорится, те же самые вещи позволяют двигаться

быстрее других. А если организация ограничится восторгом, закупкой лицензий и общими словами, она будет зависеть от погоды целиком.

История искусственного интеллекта не даёт готовой стратегии. Но она даёт хорошую прививку от двух болезней: паники и самоуверенности.

Паника говорит: «Всё изменилось, мы не успеем». Самоуверенность говорит: «Мы сейчас внедрим платформу, и всё заработает». Обе говорят слишком громко.

Лучше начать тише. Посмотреть, что ИИ уже умеет. Вспомнить, как долго он к этому шёл. Признать, что у него были зимы не случайно. Выбрать несколько задач, где польза проверяема. Не обещать чудес. Не прятаться за осторожностью. Не путать старые мечты с сегодняшними возможностями. И не забывать, что самые прочные изменения обычно выглядят скучно до тех пор, пока не становятся незамеченными.

Когда технология взрослеет, она перестаёт требовать аплодисментов. Она начинает работать.

Искусственный интеллект ещё не везде дошёл до этой стадии. Где-то он уже обычный инструмент. Где-то всё ещё ярмарочный фокус. Где-то дорогой эксперимент. Где-то риск, который плохо понимают. Но если смотреть на него не как на внезапное чудо, а как на длинную историю, становится легче выбрать правильный тон.

Не восторг. Не страх. Рабочее любопытство.

С него и стоит начинать.

Не магия, а привычка считать

Один из самых вредных способов говорить об искусственном интеллекте выглядит невинно. Человек поднимает глаза от экрана и говорит: «Это же почти как живой ум».

Я понимаю, почему так хочется сказать. Иногда ответ машины действительно похож на мысль. Она подбирает слова, объясняет, спорит, шутит, признаёт ошибку, переписывает фразу, предлагает другой вариант. После первых удачных опытов легко забыть, что перед нами не существо, а система. Не коллега, который много читал и мало спал. Не консультант с опытом. Не новый сотрудник, которого нужно познакомить с командой.

Перед нами математика, статистика, данные, правила, вычисления и очень много инженерной работы.

Сказано сухо. Зато честно.

Проблема в том, что слово «интеллект» само толкает нас в сторону лишних ожиданий. Мы слышим его и сразу думаем о человеке. О способности понимать, хотеть, сомневаться, помнить детство, обижаться, менять мнение после тяжёлого разговора. Но искусственный интеллект не обязан быть похожим на человеческий ум, чтобы быть полезным. И не обязан быть бесполезным только потому, что не похож.

Это первая вещь, которую руководителю стоит принять без раздражения: ИИ не надо любить как чудо и не надо раз-

облачать как мошенника. С ним надо разобраться как с инструментом, который иногда ведёт себя неожиданно.

Слово «инструмент» тоже не совсем точное. Молоток не предлагает вам новый способ построить дом. Таблица не спорит с финансовым директором. Навигация не пишет письмо клиенту. А современные ИИ-системы уже могут участвовать в работе так, что граница между инструментом и помощником становится размытой. Но размытой не значит исчезнувшей.

Когда сотрудник говорит: «ИИ решил», в компании должен найтись кто-то, кто спросит: «Что именно он сделал?» Предложил текст? Выбрал вариант? Посчитал вероятность? Нашёл похожие случаи? Составил прогноз? Сгенерировал картинку? Сократил документ? Принял решение без чловека?

Все эти действия разные. Риски разные. Требования к качеству разные. Ответственность разная.

Поэтому разговор об ИИ лучше начинать не с больших слов, а с простого вопроса: какую задачу он выполняет?

Когда мы говорим «искусственный интеллект», мы складываем в одну корзину множество разных методов. Там есть старые экспертные системы, которые работают по правилам. Есть машинное обучение, которое ищет закономерности в данных. Есть нейронные сети. Есть системы распознавания речи и изображений. Есть рекомендательные алгоритмы. Есть большие языковые модели. Есть генераторы изоб-

ражений, музыки, видео, кода. Есть автономные агенты, которые не просто отвечают, а пытаются выполнить цепочку действий.

Для технического специалиста различия между ними важны. Для руководителя тоже, но в другом смысле. Ему не обязательно знать устройство каждого алгоритма. Но ему нужно понимать, какой тип мышления заложен в системе. Если выразиться грубо, есть два больших семейства: системы, которые следуют правилам, и системы, которые учатся на данных.

Этого деления уже достаточно, чтобы перестать воспринимать ИИ как туман.

Представьте небольшую сеть кофеен. Владелец хочет предугадывать, что постоянный клиент закажет утром, днём или вечером. Можно пойти старым путём и написать правила. Если клиент приходит до девяти утра в будний день, предложить американо. Если в субботу после одиннадцати и с друзьями, предложить капучино. Если после четырёх и выглядит усталым, предложить двойной эспрессо. Если на улице жара, показать холодный напиток.

Такую систему легко объяснить. Она почти бытовая. Вот правило, вот условие, вот ответ. Руководитель может открыть список правил и понять, почему клиенту показали именно это предложение. Юрист может проверить, нет ли там запрещённых признаков. Маркетолог может спорить с логикой. Операционный директор может спросить, кто об-

новляет правила и как часто.

В этом сила подхода на правилах. Он прозрачен.

Но у него быстро появляется слабое место: жизнь не помещается в список правил.

Клиент пришёл в восемь сорок, но сегодня воскресенье. Он с друзьями, но не хочет сладкого. Он обычно пьёт кофе, но неделю назад врач посоветовал ему снизить кофеин. Он нажал в приложении на латте, потому что ошибся пальцем. Он покупает напиток не себе, а коллеге. Он устал, но не из-за работы, а потому что всю ночь летел. Какое правило срабатывает? Кто будет записывать все исключения? Сколько правил понадобится, чтобы не выглядеть глупо?

В маленьком примере это забавно. В медицине, страховании, банках, логистике, кадровом отборе и промышленности уже не забавно.

Системы на правилах долго были основой многих полезных решений. Они и сейчас не исчезли. Если нужно жёстко соблюдать закон, регламент или корпоративную политику, правила незаменимы. Нельзя просить модель «как-нибудь» соблюдать требования. Ей нужно задать границы. Нельзя выдавать кредит выше лимита. Нельзя отправлять персональные данные туда, куда нельзя. Нельзя рекомендовать клиенту продукт, если он запрещён условиями.

Но если вся логика держится только на правилах, компания быстро упирается в трудоёмкость. Правила надо придумать, проверить, согласовать, обновлять. Они конфликтуют.

Они устаревают. Они закрывают вчерашние случаи и пропускают завтрашние.

Тут появляется другой подход: не описывать поведение вручную, а дать системе данные и позволить ей найти закономерности.

Вернёмся к кофейне. Вместо того чтобы писать правила, можно собрать историю заказов. Когда клиент приходил, что покупал, какая была погода, был ли он один, какие акции действовали, сколько времени прошло с прошлого визита. Система смотрит на тысячи таких случаев и начинает предсказывать следующий заказ. Она не знает клиента как бариста, который видит его каждое утро. Но она видит повторы. Если похожие обстоятельства часто заканчиваются капучино, она предложит капучино.

В этом сила подхода на данных. Он масштабируется.

Человеку трудно просмотреть миллион историй заказов. Машине не трудно. Человеку трудно заметить слабую закономерность между временем визита, погодой, скидкой и поведением клиента через три дня. Машина может заметить. Человеку сложно удержать в голове тысячи переменных. Машина для этого и создана.

Но у подхода на данных есть своя неприятная правда: система учится не на мире, а на ваших данных о мире.

Если данные бедные, система будет бедной. Если данные кривые, система будет кривой. Если в данных есть перекосы, система научится перекосам. Если вы учили модель на

клиентах одного региона, а применяете в другом, не удивляйтесь странностям. Если исторические решения компании были несправедливыми, модель может аккуратно воспроизвести эту несправедливость и придать ей вид объективного расчёта.

Фраза «данные говорят» звучит убедительно, но данные сами не говорят. Их кто-то собрал. Кто-то решил, что измерять. Кто-то оставил пропуски. Кто-то ввёл значения вручную. Кто-то назвал один и тот же объект разными словами. Кто-то удалил неудобные записи. Кто-то построил систему так, что ошибки копились годами и никого не трогали, пока на них не начали обучать модель.

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.