

— С Е Р Г Е Й А Н Т О Н О В —

ВНЕДРИТЬ ИИ.

**ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ
ТАМ, ГДЕ ОН ДЕЙСТВИТЕЛЬНО
ПРИНОСИТ РЕЗУЛЬТАТ.**



Сергей Антонов
Внедрить ИИ. Искусственный
интеллект там, где
он действительно
приносит результат

<https://litres.ru/74253573>

SelfPub; 2026

Аннотация

После этой книги Вы будете смотреть на AI не как на моду, а как на инструмент, который помогает работать быстрее, спокойнее и умнее.

Вы точно знаете, где у вас уходят деньги? Обычно нет. Они прячутся в лишних отчётах, плохих правилах, запасах «на всякий случай», красивых показателях и привычке терпеть потери. Здесь Вы начнёте видеть работу без тумана: поймёте, с кем себя сравнивать, где настоящая граница возможностей, почему средний результат часто обманывает и как отличить экономию от саморазрушения. Вы перестанете верить голым KPI, узнаете, какие цифры Вас дурят, сможете найти слабое место в процессе, не устраивая охоту на виноватых. Эта книга поможет Вам сэкономить деньги, время и силы там, где они сейчас тихо сгорают каждый день.

Содержание

Введение	4
Сначала понять, потом автоматизировать	10
Как выбирать рамку, а не модное слово	29
Когда код перестаёт быть только кодом	46
Конец ознакомительного фрагмента.	54

Сергей Антонов Внедрить ИИ. Искусственный интеллект там, где он действительно приносит результат

Введение

Искусственный интеллект вошел в деловую речь слишком быстро. Вчера его обсуждали инженеры, исследователи и несколько упрямых людей из продуктовых команд. Сегодня о нем спрашивает владелец небольшой фабрики, директор клиники, руководитель отдела продаж, основатель интернет-магазина, бухгалтер, юрист и человек, который еще месяц назад считал, что нейросети нужны только для картинок. Это не мода в обычном смысле. Мода приходит и уходит, а здесь меняется сам способ работы: как мы ищем информацию, пишем, считаем, проверяем гипотезы, разговариваем с клиентами, обучаем сотрудников и принимаем решения.

При этом вокруг ИИ уже успели вырасти две крайности.

Одни говорят о нем как о чуде, которое само закроет все слабые места компании. Другие стараются не смотреть в эту сторону, потому что за новыми словами слышат угрозу: людей заменят, опыт обесценится, рынок станет чужим. Обе реакции понятны. Но обе плохо помогают управлять бизнесом. Чудо не внедряется по плану, а страх не дает проверить, где технология действительно полезна. Поэтому начинать приходится не с восторга и не с обороны. Начинать приходится с трезвого разговора.

Эта книга написана для руководителя, предпринимателя и специалиста, которому надо не просто знать, что такое искусственный интеллект, а понять, как с ним работать в обычной компании. Не в лаборатории, не на сцене конференции, не в рекламной презентации поставщика, а там, где есть сроки, счета, клиенты, усталые сотрудники, старые таблицы, неполные данные и начальник, которому нужно объяснить, зачем все это стоит денег. Технологии часто описывают так, будто бизнес живет в стерильной комнате. На практике он живет в помещении, где одновременно звонит телефон, ломается интеграция, клиент просит скидку, менеджер забыл обновить CRM, а маркетинг уже пообещал то, чего продукт еще не умеет.

Именно в такой среде ИИ либо приносит пользу, либо превращается в игрушку. Полезная система не обязана выглядеть эффектно. Иногда она просто быстрее распределяет заявки. Иногда заранее замечает, что станок скоро остано-

вится. Иногда помогает врачу увидеть риск в наборе снимков. Иногда пишет черновик письма, который человек потом спокойно правит. Иногда собирает из разрозненных заметок нормальный отчет. Ценность здесь не в том, что машина имитирует ум. Ценность в том, что люди получают больше времени на работу, где действительно нужен человеческий опыт.

Но для этого руководителю надо избавиться от одной опасной привычки: думать об ИИ как об отдельной программе, которую можно купить и поставить рядом с остальными. Искусственный интеллект почти всегда цепляется за уже существующие процессы. Он требует данных, понятных целей, ответственных людей, проверки качества, доступа к системам и ясных правил. Если компания не понимает, как у нее устроен процесс продаж, ИИ не сделает его зрелым. Если данные грязные, модель не станет мудрой. Если люди боятся задавать вопросы, пилотный проект будет выглядеть успешным только в отчете.

Поэтому первая задача не техническая. Надо понять, какую работу компания хочет улучшить. Не абстрактно повысить эффективность, а конкретно: сократить время ответа клиенту, уменьшить число ошибок в документах, точнее прогнозировать спрос, быстрее готовить коммерческие предложения, выявлять поломки до аварии, помогать сотрудникам учиться на внутренних знаниях компании. Чем яснее задача, тем легче выбрать инструмент. Чем мутнее за-

дача, тем больше шансов купить красивую обертку и потом полгода объяснять, почему она не окупилась.

В этой книге я буду говорить об ИИ как о практическом инструменте деловой трансформации. Не как о магии и не как о новой религии производительности. Мы разберем базовые понятия, типы систем, сценарии применения, этапы внедрения, культуру, роль команд, управление рисками, работу с генеративными моделями, связь ИИ с новыми технологиями и то, почему этические вопросы нельзя оставлять юристам на последнюю неделю проекта. Каждый раздел будет возвращаться к одному простому вопросу: что должен сделать человек в компании, чтобы технология не осталась красивой игрушкой?

У ИИ есть неприятное свойство: он увеличивает масштаб того, что уже есть. Хороший процесс он ускоряет. Плохой процесс он делает плохим быстрее. Сильную команду он разгружает. Слабую команду он запутывает. Честную аналитику он усиливает. Самообман он превращает в аккуратные графики. Поэтому внедрение начинается с управленческой честности. Нужно признать, где компания не знает своих данных, где решения принимаются на слух, где сотрудники держат ключевые знания в голове, где клиентский опыт строится на героизме отдельных людей, а не на системе.

Эта честность нужна не для самобичевания. Она нужна для выбора маршрута. Малому бизнесу не обязательно сразу строить сложную платформу. Крупной компании нель-

зя ограничиться чат-ботом для презентации совету директоров. Производству мало генератора текстов, если главная боль в простоях оборудования. Клиника не должна бросаться в модную модель, если не решены доступы, безопасность и ответственность за решение. В каждом случае ИИ начинается с разных вопросов, но почти всегда упирается в одно: умеет ли организация учиться быстрее, чем меняется ее рынок.

Я буду часто говорить о генеративном ИИ, потому что именно он сделал тему массовой. Текст, изображения, звук, код, диалоговые интерфейсы - все это стало ближе к обычному пользователю. Но генеративные модели не отменяют другие виды ИИ. Прогнозирование спроса, компьютерное зрение, анализ рисков, оптимизация маршрутов, рекомендательные системы, обработка документов, контроль качества и обнаружение мошенничества во многих компаниях дают более прямой экономический эффект, чем самый красивый чат. Ошибка начинается там, где все многообразие ИИ сводят к одному окну с запросом.

Введение нужно закончить простой мыслью. Внедрение ИИ не начинается с вопроса: какую модель выбрать. Оно начинается с вопроса: какую способность должна получить компания. Быстрее видеть изменения? Лучше помнить собственные знания? Точнее обслуживать клиентов? Меньше тратить людей на повторяемую работу? Быстрее проверять идеи? Когда ответ найден, технологии становятся не шумом,

а набором инструментов. Тогда можно двигаться дальше.

Сначала понять, потом автоматизировать

Современный бизнес привык жить рядом с технологиями, но искусственный интеллект отличается от большинства прежних инструментов. Электронная почта ускорила переписку. CRM собрала клиентов в одну базу. Облачные сервисы упростили доступ к файлам. ИИ делает другое: он вмешивается в самую работу с неопределенностью. Он помогает увидеть закономерность там, где человек видит шум, предложить черновик решения там, где раньше была пустая страница, заметить отклонение до того, как оно станет аварией. Поэтому к нему нельзя относиться как к очередной кнопке в интерфейсе.

Чтобы не запутаться, стоит начать с простого определения. Искусственный интеллект - это область, где создают системы, способные выполнять задачи, которые раньше требовали человеческого мышления: распознавать речь и изображения, классифицировать данные, учиться на примерах, строить прогнозы, находить связи, генерировать текст или код, помогать в принятии решений. В этом определении нет мистики. Машина не становится человеком. Она обрабатывает данные и действует по правилам, которые заданы архитектурой, обучением и ограничениями. Но результат иногда

выглядит так, будто перед нами помощник с опытом.

Главная ошибка новичка - смешивать все виды ИИ в один комок. Когда руководитель говорит: нам нужен искусственный интеллект, он часто имеет в виду что угодно: чат-бота, автоматизацию отчетов, прогноз продаж, поиск по документам, генератор изображений, ассистента для сотрудников или систему контроля брака на линии. Это разные задачи. Для них нужны разные данные, разные подходы, разные критерии качества и разные риски. Нельзя выбрать лекарство, пока не назван диагноз. Нельзя выбрать ИИ, пока не описана работа, которую он должен улучшить.

Обычно говорят о двух больших полюсах. Первый - узкий искусственный интеллект. Он решает конкретную задачу: распознает номер автомобиля, рекомендует товар, сортирует заявки, переводит текст, ищет мошеннические операции, прогнозирует поломку. Он может быть очень сильным в своей области, но не понимает мир целиком. Второй полюс - общий искусственный интеллект, система, которая могла бы решать широкий круг задач на уровне человека или выше. О нем много спорят, строят прогнозы, пугаются и мечтают. Для обычного бизнеса сегодня практическую ценность почти всегда дает узкий ИИ, даже если внешне он кажется очень умным.

Узкий ИИ полезен именно потому, что ему можно дать границы. Компания не нуждается в машине, которая рассуждает обо всем на свете, если ей надо сократить время обра-

ботки претензий. Ей нужна система, которая понимает документы, вытаскивает нужные факты, предлагает ответ и передает сложные случаи человеку. В этом есть управленческая зрелость: не просить технологию заменить здравый смысл, а поручать ей работу, где правила, данные и результат можно проверить.

Внутри ИИ важное место занимает машинное обучение. Его смысл прост: вместо того чтобы прописывать каждое правило вручную, мы даем системе примеры, и она учится находить закономерности. Если у нас есть тысячи писем, отмеченных как спам и не спам, модель может научиться различать новые письма. Если есть история продаж, погода, праздники, цены и рекламные акции, модель может помогать прогнозировать спрос. Если есть снимки с отметками врачей, система может учиться видеть признаки заболевания. Машинное обучение не отменяет человека. Оно берет на себя часть наблюдения, где данных слишком много для обычного взгляда.

Самый понятный вариант - обучение с учителем. У модели есть примеры и правильные ответы. Ей показывают входные данные и говорят, каким должен быть результат. Это похоже на обучение сотрудника на разобранных случаях: вот заявка, вот ее категория; вот клиент, вот вероятность ухода; вот фотография детали, вот отметка о дефекте. Такой подход хорош там, где компания уже накопила историю и может объяснить, что считала правильным решением. Но если ис-

тория хаотична, модель честно выучит хаос.

Есть обучение без учителя. Здесь правильных ответов заранее нет. Система ищет группы, похожие объекты, необычные отклонения, скрытые связи. В бизнесе это полезно, когда нужно сегментировать клиентов, найти нетипичное поведение, разобрать массив документов или увидеть, что часть операций живет по особому сценарию. Человек мог бы заметить это после долгого анализа, но у него редко есть время и терпение просматривать тысячи строк. Модель не устает. Зато она не знает, какая находка имеет смысл для бизнеса. Объяснять смысл все равно придется людям.

Есть смешанные подходы, где часть данных размечена, а часть нет. Это важная практическая история, потому что у компаний почти всегда мало идеальных данных. Документы лежат в разных папках, клиенты заполнены с ошибками, фотографии подписаны не везде, старые сделки описаны в свободной форме. Полностью разметить все дорого и долго. Поэтому хорошие проекты часто строятся постепенно: немного ручной разметки, проверка качества, расширение датасета, новая версия модели, снова проверка. Это не похоже на эффектный запуск. Это больше похоже на ремонт дороги, по которой уже едут машины.

Отдельно стоит назвать обучение с подкреплением. Здесь система учится через действия и последствия. Она пробует, получает награду или штраф, меняет поведение и ищет стратегию, которая дает лучший результат. Такой подход изве-

стен по играм и робототехнике, но в бизнесе он тоже встречается: оптимизация цен, логистики, рекомендаций, распределения ресурсов. Однако с ним надо быть аккуратным. Если неверно задать награду, система начнет побеждать в той игре, которую вы сами плохо описали. Например, если оценивать поддержку только по скорости закрытия обращений, она может быстрее закрывать тикеты, но хуже помогать клиентам.

Последние годы внимание забрал генеративный ИИ. Его особенность в том, что он не только классифицирует или прогнозирует, а создает новый материал: текст, изображение, звук, видео, код, структуру документа, план, варианты ответа. Он строит результат на основе закономерностей, найденных в огромных массивах данных. Для бизнеса это заметный сдвиг. Раньше автоматизация чаще касалась повторяемых операций. Теперь она подошла к задачам, которые считались почти полностью человеческими: написать черновик, придумать варианты, объяснить сложную тему, собрать резюме встречи, адаптировать текст под аудиторию.

Но генеративный ИИ нельзя пускать в компанию без правил. Он может ошибаться уверенно. Он может придумывать источники, сглаживать важные различия, повторять предвзятость обучающих данных, раскрывать лишнее, если сотрудники бездумно вставляют в запросы конфиденциальную информацию. Поэтому хороший пилот начинается не с восторженного списка возможностей, а с вопроса: где ошибка

модели допустима, а где нет. Черновик рекламного текста можно поправить. Неверный медицинский совет, финансовая рекомендация или юридический вывод без проверки могут стоить дорого.

В деловой среде ИИ дает эффект в нескольких больших направлениях. Первое - прогнозирование. Компании всегда пытаются угадать будущее: сколько товара понадобится, какие клиенты уйдут, где будет перегрузка, когда оборудование выйдет из строя, какой регион даст рост. Раньше это часто делалось на опыте и простых таблицах. Опыт важен, но он плохо масштабируется и зависит от конкретного человека. Модель может взять больше факторов и быстрее пересчитать прогноз. Она не знает бизнеса лучше директора, но может показать директору то, что он не успел заметить.

Второе направление - автоматизация повторяемой работы. Здесь особенно много тихой пользы. Сотрудники тратят часы на перенос данных, сверку документов, подготовку отчетов, поиск информации в переписке, стандартные ответы клиентам. Часть этой работы требует не таланта, а внимания. ИИ и связанные с ним инструменты могут брать на себя черновую обработку: извлечь данные из счета, сопоставить заявку с договором, подготовить краткое содержание, заполнить шаблон, направить обращение нужной команде. Человек остается там, где надо оценить контекст, договориться, принять ответственность.

Третье направление - диалоговые интерфейсы. Чат-боты

существовали давно, но многие из них раздражали клиентов, потому что работали как телефонное меню в текстовом виде. Новые модели позволяют строить более гибкие ответы, понимать намерение, учитывать историю, объяснять шаги. Это полезно для поддержки, продаж, обучения сотрудников, внутреннего поиска знаний. Но и здесь простая правда: плохая база знаний превращает умного бота в уверенного болтуна. Если документы устарели, процессы не описаны, права доступа не настроены, бот будет быстро выдавать сомнительный ответ красивыми словами.

Четвертое направление - компьютерное зрение. Для производства, медицины, ритейла, транспорта и безопасности это не менее важно, чем текстовые модели. Система может находить дефекты на деталях, анализировать снимки, считать объекты, проверять выкладку товара, распознавать документы, отслеживать опасные ситуации. Здесь ценность часто очень конкретна: меньше брака, быстрее контроль, меньше аварий, точнее диагностика. Но внедрение требует камеры, правильного освещения, разметки, испытаний в реальных условиях. Презентационная точность на чистых примерах мало значит, если цех пыльный, свет меняется, а объект повернут не так, как в учебном наборе.

Пятое направление - работа с языком. Обработка естественного языка помогает разбирать обращения клиентов, анализировать отзывы, находить темы в больших массивах сообщений, переводить, суммировать, извлекать факты из

договоров, готовить ответы. Для компаний с большим количеством текста это уже не роскошь. Юридические отделы, HR, продажи, закупки, поддержка, маркетинг, обучение - все они живут в документах и переписке. Если сократить время поиска и черновой подготовки хотя бы на четверть, эффект быстро становится заметным.

Шестое направление - творчество и создание материалов. Здесь легко уйти в игрушки, но не стоит недооценивать практическую пользу. Маркетинговая команда может быстрее проверять варианты сообщений. Дизайнер может получать наброски для обсуждения. Продавец может адаптировать презентацию под отрасль клиента. Обучающий отдел может готовить примеры и упражнения. Разработчик может получить черновик кода или тестов. Вопрос не в том, заметит ли ИИ творческого человека. Вопрос в том, сколько пустого листа и рутинной подготовки он снимет с команды.

В здравоохранении ИИ уже помогает в диагностике, анализе изображений, поиске лекарственных молекул, планировании лечения, работе с пациентскими обращениями и медицинскими исследованиями. Но именно здесь особенно видны границы. Медицинская система не может просто сказать: модель решила. Нужны доказательства, проверка, ответственность, защита данных и понимание того, как врач использует подсказку. ИИ может ускорить поиск и усилить внимание специалиста. Но доверие строится не на красивом интерфейсе, а на проверяемом качестве и ясной роли чело-

века.

В ритейле и электронной коммерции ИИ давно стал частью повседневности. Рекомендации, динамические витрины, прогноз спроса, персонализация, описание товаров, генерация изображений, анализ отзывов, управление запасами - все это напрямую влияет на деньги. Покупатель не думает о модели, когда видит подходящий товар. Он просто быстрее находит нужное. Для бизнеса это означает рост конверсии, снижение лишних запасов и более точное общение с клиентом. Но персонализация легко становится навязчивой, если компания забывает о мере и прозрачности.

В производстве и цепочках поставок ИИ часто звучит менее эффектно, зато работает очень предметно. Прогноз спроса помогает планировать закупки. Анализ датчиков предупреждает о поломке. Оптимизация маршрутов снижает расходы на топливо и время доставки. Оценка поставщиков показывает риски. Контроль качества на линии сокращает брак. Здесь ИИ редко бывает отдельной красивой витриной. Он вшит в операции. Его успех измеряется не аплодисментами, а простоями, списаниями, скоростью, точностью и себестоимостью.

В сфере путешествий и гостеприимства ИИ меняет не только поддержку, но и сам способ подбора предложения. Клиент хочет не список отелей, а подходящий сценарий поездки. Семья с ребенком, командировка на два дня, путешествие ради еды, отпуск у моря, поездка с ограниченным бюд-

жетом - это разные задачи. Система может собрать предпочтения, цены, отзывы, доступность и выдать более полезные варианты. Но туристический бизнес живет на доверии. Если ИИ советует уверенно, но неточно, раздражение клиента будет сильнее, чем от обычной ошибки фильтра.

Финансовые компании используют ИИ для оценки рисков, обнаружения мошенничества, анализа документов, обслуживания клиентов, прогнозирования и персональных предложений. Здесь цена ошибки тоже высока. Если модель необоснованно отказывает человеку в услуге, проблема становится не только технической, но и этической. Если алгоритм ловит подозрительные операции, но не объясняет причины, сотрудники не могут нормально расследовать случаи. Поэтому в финансах особенно важны объяснимость, контроль, аудит и разделение полномочий между машиной и человеком.

Образование получает от ИИ возможность персонализировать обучение, быстро готовить материалы, помогать студентам с объяснениями, проверять задания, переводить и адаптировать содержание. Но здесь есть ловушка. Если заменить обучение выдачей готовых ответов, навыки у ученика не растут. Хорошая система не делает работу вместо человека, а помогает ему пройти путь: задает вопрос, предлагает пример, объясняет ошибку, дает упражнение сложнее или проще. В бизнес-обучении это особенно важно: сотруднику нужен не красивый текст, а способность лучше выпол-

нять работу завтра утром.

Когда руководитель впервые собирает карту возможностей ИИ, полезно разделить задачи на четыре группы. Первая - где нужно быстрее видеть информацию. Вторая - где нужно быстрее создавать черновики. Третья - где нужно точнее прогнозировать. Четвертая - где нужно автоматически выполнять повторяемые действия. Такая карта сразу снижает шум. Вместо разговоров про искусственный интеллект вообще появляются конкретные участки работы: входящие заявки, склад, ремонт, обучение, договоры, отчеты, клиентские письма, проверка качества, планирование спроса.

Дальше начинается неприятная, но нужная часть: оценка готовности. Есть ли данные? Кто владелец процесса? Как измерить результат? Что будет считаться ошибкой? Кто проверяет ответы? Где хранится информация? Какие ограничения по безопасности? Какие сотрудники будут пользоваться системой? Что изменится в их работе? Если на эти вопросы нет ответов, проект рано отдавать разработчикам. Они могут сделать прототип, но не смогут сами придумать управленческий смысл.

Путь внедрения обычно проходит несколько стадий. Сначала осведомленность и исследование. Компания изучает возможности, смотрит примеры, пробует инструменты, обсуждает риски. На этом этапе опасны две вещи: поверхностный восторг и бесконечное чтение без действия. Нужно быстро перейти к нескольким реальным задачам, где можно

проверить пользу. Не к десяти направлениям сразу, а к двум-трем, которые понятны бизнесу и имеют измеримый результат.

Потом идет эксперимент и прототип. Здесь команда собирает минимальное решение, проверяет данные, смотрит качество, слушает пользователей. Это этап, где обычно вскрываются старые проблемы. Оказывается, что данные разбросаны. Что сотрудники по-разному называют одно и то же. Что в CRM половина полей заполнена как попало. Что юридический отдел не согласовал использование внешнего сервиса. Это не провал. Это диагностика. Хороший прототип полезен даже тогда, когда показывает: сначала надо привести в порядок основу.

Следующая стадия - операционное внедрение и масштабирование. Тут ИИ перестает быть демонстрацией и входит в рабочий день. Появляются интеграции, права доступа, инструкции, поддержка, метрики, ответственность, обновления. У многих компаний именно здесь ломается энтузиазм. Пилот сделал сильный сотрудник, а пользоваться должны десятки людей с разным уровнем подготовки. В презентации все было просто, а в реальной системе есть исключения, жалобы, нестандартные случаи. Поэтому масштабирование - это не копирование прототипа. Это перестройка процесса вокруг нового инструмента.

Зрелое использование ИИ начинается тогда, когда компания перестает относиться к нему как к проекту с да-

той окончания. Модели надо наблюдать, обновлять, проверять на смещение, сравнивать с новыми данными, отключать неудачные сценарии, добавлять ограничения, обучать людей. Рынок меняется, клиенты меняются, данные меняются. То, что работало год назад, может начать ошибаться. Без мониторинга ИИ превращается в старую автоматизацию с умным названием.

Самый трудный вопрос - доверие. Люди не обязаны верить системе только потому, что ее назвали интеллектуальной. Сотруднику надо понимать, где подсказка помогает, а где он обязан проверить. Клиенту надо знать, когда с ним говорит бот. Руководителю нужны метрики качества, а не рассказы о потенциале. Юристу нужны правила обработки данных. Специалисту по безопасности нужны ограничения доступа. Если доверие строится на вере, оно быстро рушится. Если на проверке, ролях и прозрачности, оно выдерживает ошибки и улучшается вместе с системой.

Есть еще один момент, о котором часто говорят поздно: ИИ меняет работу людей. Даже небольшой инструмент может изменить привычный порядок. Кто раньше писал отчеты вручную, теперь проверяет черновик. Кто раньше искал информацию, теперь формулирует запрос и оценивает ответ. Кто раньше отвечал на типовые вопросы, теперь разбирает сложные случаи. Это не всегда воспринимается как облегчение. Иногда человек чувствует, что его опыт обесценили. Поэтому руководителю надо объяснять не только пользу для

бизнеса, но и новую роль сотрудников.

Обучение команды нельзя сводить к инструкции с кнопками. Людям нужно научиться задавать задачи модели, проверять результат, видеть ошибки, защищать данные, понимать границы автоматизации. Хороший сотрудник в эпоху ИИ - не тот, кто слепо верит ответу машины, а тот, кто умеет использовать ее как сильного, но не безошибочного помощника. Это новая грамотность. Она включает техническую привычку, деловое мышление и здоровое недоверие.

Этические вопросы не висят где-то сбоку. Они входят в экономику проекта. Если модель дискриминирует часть клиентов, компания теряет доверие и получает риски. Если система непрозрачна, сотрудники не могут защищать решения. Если данные использованы без согласия, проблема выходит за пределы ИТ. Если автоматизация ухудшает качество обслуживания ради скорости, клиенты уходят. Ответственный ИИ - не украшение для годового отчета. Это способ не разрушить то, что бизнес строил годами.

Руководитель, который хочет начать правильно, может сделать простой шаг: выбрать одну важную, но ограниченную задачу. Например, ускорить обработку входящих запросов от клиентов. Затем описать текущий процесс: откуда приходит запрос, кто его читает, как определяется категория, где берется информация для ответа, сколько времени уходит, какие ошибки самые частые. После этого можно понять, что именно должен делать ИИ: классифицировать, ис-

кать похожие случаи, предлагать черновик, проверять полноту ответа или передавать сложные обращения старшему специалисту.

Такой подход кажется медленным только на бумаге. На практике он экономит месяцы. Компания не спорит о терминах, а смотрит на рабочий участок. Не покупает платформу из страха отстать, а проверяет конкретную пользу. Не заставляет всех срочно становиться экспертами, а собирает маленькую команду из владельца процесса, специалиста по данным, пользователя, IT и человека, который отвечает за риски. У этой команды появляется шанс сделать не красивый эксперимент, а основу для следующего шага.

Именно поэтому первая глава не должна закончиться призывом срочно внедрять ИИ везде. Спешка здесь вредна так же, как и промедление. Бизнесу нужна дисциплина: понять тип задачи, выбрать подходящий вид ИИ, проверить данные, сделать прототип, измерить результат, встроить в процесс, обучить людей, наблюдать качество. Звучит не так громко, как обещания революции. Зато именно так технология становится частью компании, а не темой для презентации.

Если убрать шум, искусственный интеллект дает бизнесу три сильные способности. Он помогает видеть больше, чем человек успевает увидеть сам. Он помогает делать черновую работу быстрее, чем команда делает ее вручную. Он помогает проверять варианты и решения чаще, чем позволял старый ритм. Но он не освобождает руководителя от выбора.

Наоборот, делает выбор заметнее. Там, где раньше можно было спрятаться за занятостью и привычкой, теперь приходится сказать: вот задача, вот данные, вот риск, вот человек, который отвечает за результат.

В этом и состоит настоящая сила ИИ в современном бизнесе. Не в том, что он заменяет разум. А в том, что заставляет компанию лучше организовать собственный разум: знания, процессы, решения, ответственность, обучение. Технология может быть сложной, но первый шаг прост. Нужно честно посмотреть на работу компании и выбрать место, где машина способна помочь людям делать дело лучше. Не громче. Не моднее. Лучше.

Есть простой тест на зрелость ИИ-проекта. Попросите команду объяснить его без слов: модель, инновация, трансформация, экосистема. Если после удаления этих слов смысл исчез, значит проекта пока нет. Есть желание казаться современными. Настоящий проект можно объяснить обычным языком: сегодня оператор тратит семь минут на заявку, мы хотим сократить до двух; сегодня менеджер ошибается в двадцати счетах из тысячи, мы хотим снизить до пяти; сегодня ремонт происходит после поломки, мы хотим видеть риск за неделю. В таких формулировках уже есть жизнь.

Второй тест касается данных. Спросите, кто отвечает за качество данных, которые должна использовать система. Если все смотрят друг на друга, ИИ лучше не трогать. Модель не знает, что ваша таблица велась временно, что поле регион

заполнялось по настроению, что один и тот же клиент записан тремя способами, что часть сделок перенесли вручную после сбоя. Она принимает это как реальность. Потом бизнес удивляется странным выводам. Но странность часто появилась не в модели, а в привычках компании.

Третий тест - место человека в процессе. Если команда говорит, что ИИ будет сам все решать, надо остановиться и спросить: кто проверяет, кто отвечает, кто исправляет, кто видит спорные случаи. Полная автономия возможна только там, где последствия ошибки понятны и ограничены. Во многих деловых задачах лучше работает схема помощника: система делает черновую работу, подсказывает, ранжирует, предупреждает, а человек принимает решение. Это не слабость. Это нормальная инженерная осторожность.

Еще одна практическая вещь - цена бездействия. Иногда компания боится вложиться в ИИ и ничего не делает, но не считает, сколько уже теряет. Сколько часов сотрудники тратят на ручные отчеты? Сколько клиентов уходят из-за медленного ответа? Сколько денег заморожено в лишнем запасе? Сколько ошибок проходит через документы? Сколько знаний исчезает, когда опытный сотрудник увольняется? Если эти потери посчитать, разговор становится спокойнее. ИИ перестает быть затратой на модную тему и становится способом закрыть понятную дыру.

Но и обратная сторона важна. Не каждая проблема заслуживает ИИ. Иногда достаточно нормального регламента,

простой автоматизации, чистой формы ввода, обучения сотрудников или пересмотра мотивации. Если склад ошибается из-за того, что товар лежит без маркировки, модель не спасет. Если продажи не заносят данные в CRM, прогноз будет гаданием. Если поддержка отвечает плохо, потому что у нее нет полномочий решить вопрос клиента, чат-бот лишь быстрее покажет бессилие компании. ИИ силен, когда его ставят на подготовленную почву.

Поэтому первое управленческое правило звучит почти скучно: не автоматизируйте неразобранный беспорядок. Сначала найдите повторяемую работу, понятные входы, измеримый результат и человека, который готов отвечать за процесс. Потом проверяйте технологию. Такой порядок кажется менее эффективным, зато он дает шанс дойти до внедрения, а не остановиться на демо. В бизнесе ценится не тот ИИ, который впечатлил на встрече, а тот, который через три месяца продолжает экономить время и снижать ошибки.

Наконец, нужно помнить о языке, на котором компания говорит об ИИ. Если говорить только языком угрозы, люди будут прятаться и саботировать. Если говорить только языком восторга, они перестанут сообщать о проблемах, потому что не захотят портить праздник. Нужен рабочий язык: что меняем, зачем, как проверяем, где границы, кого обучаем, что делаем при ошибке. В таком языке меньше блеска, но больше доверия. А без доверия даже сильная технология остается чужой.

Первая глава нужна, чтобы поставить основание. Искусственный интеллект не обязан быть загадкой для руководителя. Не нужно знать все детали архитектуры, чтобы принимать здравые решения. Но нужно понимать типы задач, природу данных, роль человека, цену ошибки и этапы внедрения. Тогда разговор с технической командой становится предметным. Тогда поставщик не может легко продать туман. Тогда сотрудники видят, что от них хотят не поклонения машине, а новой дисциплины работы.

С этого места можно двигаться к рамкам внедрения. Любая компания, которая хочет использовать ИИ серьезно, рано или поздно приходит к вопросу управления: кто выбирает сценарии, как оценивается готовность, как строится пилот, как масштабируется решение, как меняется культура, как удерживаются этические границы. Но прежде чем говорить о рамках, надо было разобраться с силой самой технологии. Она велика. Но она становится полезной только тогда, когда попадает в руки людей, которые понимают свой бизнес лучше, чем чужие презентации.

Как выбирать рамку, а не модное слово

После первого разговора об искусственном интеллекте обычно возникает соблазн сразу перейти к инструментам. Кто-то просит список сервисов. Кто-то хочет нанять специалиста. Кто-то открывает чат-бот, загружает туда пару документов и через неделю сообщает команде, что компания теперь использует ИИ. На бумаге это выглядит бодро. В реальной работе быстро выясняется неприятная вещь: технология есть, а способа обращаться с ней нет.

Я видел это не один раз. Руководитель приходит на встречу с хорошей мыслью: автоматизировать поддержку клиентов, ускорить подготовку коммерческих предложений, сделать прогноз спроса, убрать ручную возню с таблицами. Мысль здравая. Деньги на эксперимент есть. Люди тоже есть. Но через месяц команда спорит уже не о пользе, а о том, кто вообще отвечает за результат. Продажи говорят, что модель не понимает клиента. Аналитики отвечают, что им дали грязные данные. Юристы спрашивают, кто разрешил загружать внутренние письма во внешний сервис. Разработчики просят нормальное описание задачи. Владелец проекта злится, потому что ожидал быстрый эффект, а получил клубок старых проблем.

Вот почему разговор о внедрении ИИ почти всегда должен начинаться не с модели, а с рамки работы. Не с красивого названия, не с презентации на сорок слайдов, а с простого вопроса: каким способом мы будем принимать решения, проверять гипотезы, защищать данные, обучать людей и признавать ошибки? Если ответа нет, проект живёт на энтузиазме одного человека. Энтузиазм полезен на старте, но на нём трудно построить устойчивую систему.

В этой главе речь пойдёт о таких рамках. В источнике рядом стоят классические подходы Agile и Waterfall, лидерская рамка ARIA и практичный AI Use Case Canvas. Их можно воспринимать как набор готовых схем. Но пользы больше, если смотреть на них проще: это разные способы не потерять голову, когда компания пытается встроить ИИ в живую работу.

Начнём с главного. Внедрение ИИ отличается от обычного внедрения софта тем, что заранее редко известно, какой результат даст модель. В классическом проекте можно описать экран, кнопку, отчёт, интеграцию, роль пользователя. В ИИ-проекте часто приходится сначала проверить, существует ли нужный сигнал в данных. Может оказаться, что клиентское поведение прогнозируется хуже, чем ожидали. Может оказаться, что данные есть, но они собраны в разных системах и противоречат друг другу. Может оказаться, что бизнес хочет точность девяносто пять процентов, а экономический смысл появляется уже при семидесяти пяти. А может быть и

обратная история: модель показывает приличный результат на тесте, но люди в отделе не доверяют ей и продолжают работать по старой привычке.

Поэтому рамка нужна не для красоты. Она нужна, чтобы компания не путала прототип с внедрением, демонстрацию с пользой, прогноз с решением, автоматизацию с ответственностью.

Agile часто выбирают там, где много неизвестного. Его смысл не в том, чтобы каждый день стоять у доски и произносить правильные слова. Смысл в коротком цикле: выбрали гипотезу, сделали малую версию, проверили на данных и пользователях, получили обратную связь, изменили направление. Для ИИ это естественная среда. Модель редко становится полезной с первой попытки. Её приходится кормить данными, сравнивать подходы, менять метрики, обсуждать с теми, кто будет пользоваться результатом.

Представим сеть клиник, которая хочет предсказывать неявку пациентов на приём. Идея понятна: если заранее знать, кто с высокой вероятностью не придёт, можно отправить напоминание, предложить другое время или иначе распределить загрузку врачей. В Agile-подходе команда не станет сразу строить огромную систему. Она возьмёт один город, одну категорию приёмов, несколько месяцев исторических записей и проверит, есть ли там устойчивые признаки. Затем покажет результат администраторам. Не дирекции в виде абстрактного графика, а людям, которые каждый день

звонят пациентам и видят реальные причины неявок. Они скажут: тут вы правы, а тут модель путает людей после переноса записи с теми, кто просто забыл. После этого команда исправит набор признаков и снова проверит.

Такой подход выглядит медленнее, чем обещания консультанта, но на практике часто быстрее. Потому что он рано вскрывает слабые места. Не через год, когда бюджет потрачен, а через две недели, когда ещё можно спокойно повернуть.

У Agile есть и слабая сторона. Он легко превращается в бесконечное движение без решения. Команда может улучшать прототип, добавлять графики, спорить о качестве модели и никогда не дойти до промышленного режима. Особенно если нет владельца, который связывает эксперимент с деньгами, рисками и поведением пользователей. ИИ любит итерации, но бизнес не может жить в вечном черновике. В какой-то момент нужно сказать: эта версия достаточно хороша для пилота, эти риски приняты, эти ограничения понятны, за этот результат отвечает конкретная группа людей.

Waterfall, или водопадный подход, часто ругают за жёсткость. Иногда заслуженно. Если требования меняются каждую неделю, если данные ещё не изучены, если пользователи сами не знают, что им нужно, длинный линейный проект быстро начинает скрипеть. Но у Waterfall есть своя область, где он не просто уместен, а необходим. Это проекты с жёсткими требованиями к безопасности, документации,

аудиту, интеграции со старыми системами. Банки, страхование, медицина, государственный сектор, крупное производство. Там нельзя просто сказать: мы попробовали, вроде работает, выкатываем.

Допустим, банк внедряет ИИ для помощи в оценке кредитного риска. Даже если модель не принимает финальное решение, а только даёт рекомендацию, вокруг неё возникает много вопросов. Какие данные использовались? Можно ли объяснить результат? Не дискриминирует ли модель отдельные группы клиентов? Кто имеет право менять параметры? Как хранится история решений? Как откатиться назад, если обнаружена ошибка? Здесь линейность Waterfall может быть полезна. Сначала требования. Потом архитектура. Потом работа с данными. Потом разработка. Потом тестирование. Потом проверка безопасности и юридическая оценка. Потом запуск с понятными правилами контроля.

Это не означает, что Waterfall лучше. Это означает, что в некоторых условиях цена хаоса слишком высока. Там, где ошибка модели может привести к судебным претензиям, потере лицензии или прямому вреду человеку, спешка выглядит не как смелость, а как управленческая безответственность.

На практике многие компании приходят к смешанному варианту. Исследование и прототипирование идут гибко. Промышленный запуск идёт строже. Пока команда выясняет, есть ли смысл в модели, ей нужны короткие циклы. Ко-

гда модель входит в реальные процессы, нужны документы, права доступа, журналирование, контроль качества, обучение пользователей, процедуры отключения. Это нормальная взрослая логика. Сначала ищем рабочий способ. Потом строим вокруг него надёжную систему.

Проблема многих руководителей в том, что они выбирают методологию как знак принадлежности. Agile звучит современно, Waterfall звучит старомодно. Но бизнесу не платят за современное звучание. Ему платят за работающий результат. Если проект исследовательский, нужен воздух для проб. Если проект регулируемый, нужна дисциплина. Если проект касается клиента, нужна обратная связь. Если проект касается персональных данных, нужна защита. Выбор рамки должен вытекать из риска и неопределённости, а не из вкуса руководителя.

Теперь перейдём к лидерской стороне. Технологический проект с ИИ редко остаётся только технологическим. Он почти сразу затрагивает власть, привычки и страхи. Люди спрашивают, заменят ли их. Руководители среднего звена боятся потерять контроль. Специалисты опасаются, что их опыт обесценят. Юристы видят новые угрозы. Финансы хотят понятный эффект. Собственник хочет скорость. Всё это нельзя решить одной моделью.

И тут появляется логика ARIA: Assemble, Reflect, Innovate, Adapt. Если перевести на обычный язык, лидер сначала собирает контекст, затем обдумывает стратегию, затем

запускает изменения, затем учится на последствиях и корректирует курс. Вроде бы просто. Но именно эти простые шаги чаще всего пропускают.

Собрать контекст значит не только спросить у технической команды, что можно автоматизировать. Нужно понять, где у бизнеса боль, какие данные доступны, где лежат ограничения, кто выиграет, кто проиграет, кто будет сопротивляться. В хорошем проекте руководитель не стесняется задавать приземлённые вопросы. Где сейчас человек тратит время? Что он копирует вручную? Какие решения принимаются по привычке? Где ошибка стоит дорого? Какие данные мы собираем только потому, что так давно заведено? На какие процессы сотрудники жалуются уже третий год?

Один основатель интернет-магазина хотел внедрить ИИ для персональных рекомендаций. На встрече команда долго говорила о моделях, похожих товарах, истории покупок и векторных базах. Потом операционный директор тихо сказал: у нас карточки товаров заполнены как попало. В одном месте материал указан как хлопок, в другом как хлопковая ткань, в третьем вообще пусто. Фотографии тоже разного качества. Возвраты часто связаны не с рекомендациями, а с неверным ожиданием клиента. Проект резко изменился. Сначала навели порядок в товарных данных и описаниях. Только после этого рекомендации начали иметь смысл.

Reflect, или осмысление, требует паузы. Это труднее, чем кажется. В мире ИИ всем хочется действовать быстро. Но

если руководитель не отделяет деловую цель от технологического увлечения, команда начинает решать не ту задачу. Вопрос не в том, можем ли мы встроить модель. Вопрос в том, что изменится в работе компании, если модель окажется полезной. Уменьшится время ответа клиенту? Снизится стоимость обработки заявки? Вырастет точность прогноза? Освободятся специалисты для более сложных задач? Появится новый продукт?

Если цель не названа, метрики начинают жить отдельно от бизнеса. Модель показывает высокую точность, но отдел продаж не видит пользы. Чат-бот закрывает много обращений, но недовольные клиенты чаще уходят к оператору в раздражении. Система скоринга ускоряет обработку, но менеджеры перестают замечать необычные случаи. Осмысление нужно именно для того, чтобы не перепутать технический успех с деловым.

Innovate в этой логике не означает устраивать лабораторию чудес. Это про аккуратное внедрение нового способа работы. Иногда инновация выглядит скучно: убрать три ручных шага из процесса, дать сотруднику подсказку перед звонком, автоматически подготовить черновик письма, подсветить риск в договоре. Руководителям часто хочется большого жеста. Но самые полезные ИИ-изменения нередко начинаются с малых вещей, которые каждый день экономят людям по двадцать минут и уменьшают число ошибок.

Adapt, последний элемент ARIA, особенно важен. ИИ-

проект нельзя запустить и забыть. Данные меняются. Клиенты меняются. Рынок меняется. Сотрудники учатся обходить систему. Модель, которая вчера была точной, через полгода может начать ошибаться. Поэтому адаптация не должна быть героическим ремонтом после провала. Она должна быть встроенной привычкой: проверять качество, слушать пользователей, пересматривать правила, обновлять обучение, смотреть на последствия.

Лидерская рамка нужна ещё и потому, что ИИ усиливает слабые места управления. Если в компании нет культуры ответственности, ИИ даст новые способы перекладывать вину. Если нет прозрачности в данных, ИИ сделает непрозрачность дороже. Если руководитель привык принимать решения по ощущению, он будет использовать модель как украшение к уже принятому решению. Если команда боится говорить о проблемах, она промолчит и о рисках модели. Технология не делает организацию зрелой сама по себе. Она только быстрее показывает, где незрелость мешает.

AI Use Case Canvas полезен другой стороной. Он заставляет упаковать идею в понятный рабочий контур. Не просто «внедрим ИИ в маркетинг», а: какую проблему решаем, кто пользователь, какие данные нужны, какой текущий способ уже существует, почему он недостаточен, как мы измерим успех, какие риски принимаем, кто участвует, какие изменения в процессе потребуются.

Это особенно ценно для предпринимателей и владельцев

бизнеса. У них часто нет времени на большие методологические дискуссии. Им нужно быстро понять, стоит ли идея денег и внимания. Canvas помогает сделать неприятную проверку до того, как начался проект. А точно ли нужна модель? Может быть, достаточно нормального правила в CRM? Может быть, проблема не в прогнозе, а в том, что менеджеры не заполняют поля? Может быть, клиентам не нужна персонализация, им нужна честная информация о сроках доставки?

Одна из сильных мыслей Canvas: сначала оптимизируй существующее, потом автоматизируй. Это звучит буднично, но спасает от множества глупых затрат. Автоматизация плохого процесса делает плохой процесс быстрее. ИИ в таком случае не лечит, а разносит проблему шире. Если отдел поддержки отвечает клиентам путано, модель, обученная на этих ответах, будет производить более быстрый поток путаницы. Если данные о запасах неточны, прогноз спроса будет выглядеть научно, но закупки всё равно ошибутся. Если в компании нет ясного владельца процесса, ИИ не назначит его за вас.

Canvas также возвращает разговор к данным. Для машинного обучения данные не являются приложением к проекту. Это материал, из которого проект сделан. Нужно смотреть не только на объём, но и на качество, происхождение, права использования, полноту, свежесть, смещения. Руководитель может не знать всех технических деталей, но он обязан задавать вопросы. Откуда эти данные? Кто их вводит? Почему мы им доверяем? Что в них отсутствует? Кого они не пред-

ставляют? Что произойдёт, если модель ошибётся именно на редкой группе случаев?

Здесь появляется тема выбора типа решения. Не всякая задача требует генеративной модели. Иногда нужна классификация: определить категорию обращения, тип риска, вероятность оттока. Иногда нужна регрессия: предсказать сумму, срок, спрос, нагрузку. Иногда нужен поиск похожих случаев. Иногда нужен простой набор правил, который быстрее, дешевле и понятнее. Мода на большие языковые модели заставила многих забыть, что ИИ не сводится к чату. Разумный Canvas не влюбляется в инструмент. Он держит внимание на задаче.

После данных идёт готовность организации. В источнике описаны стадии: осведомлённость, исследование, операционный уровень, трансформационный уровень. Я бы сказал проще. Сначала компания слышит об ИИ и пытается понять, где он может помочь. Потом пробует малые эксперименты. Потом встраивает удачные решения в работу. И только потом начинает менять продукты, структуру и способ конкуренции. Ошибка возникает, когда руководитель хочет сразу на последнюю ступеньку. Сегодня у нас хаос в данных, завтра мы станем AI-first компанией. Так не работает. Можно сказать это на сцене. В операционке придётся идти ступенями.

На стадии осведомлённости главная задача не в том, чтобы всех вдохновить. Главная задача в том, чтобы убрать ту-

ман. Люди должны понимать базовые возможности и ограничения ИИ. Не в виде академического курса, а применительно к своей работе. Бухгалтеру не нужны лекции о трансформерах. Ему нужно понимать, какие документы можно проверять быстрее, где нельзя отдавать машине финальное решение, какие данные нельзя загружать во внешний сервис. Маркетологу нужно видеть разницу между черновиком текста и утверждённым сообщением бренда. Руководителю продаж нужно понимать, что прогноз вероятности сделки не освобождает менеджера от разговора с клиентом.

На стадии исследования нужны безопасные песочницы. Команды должны пробовать, но не тащить в эксперимент всё подряд. Нужны правила: какие данные разрешены, какие инструменты можно использовать, как фиксируются результаты, кто оценивает пользу. Без этого исследование превращается в набор личных игрушек. Один сотрудник делает промпты, другой строит таблицу, третий подключает сервис, четвёртый случайно отправляет туда то, что нельзя отправлять. Потом руководство узнаёт о «теновом ИИ» и начинает запрещать всё подряд. Запрет редко помогает. Лучше заранее дать людям понятный коридор.

Операционный уровень начинается там, где появляются владельцы, регламенты, обучение и метрики. Если чат-бот отвечает клиентам, кто проверяет качество ответов? Если модель помогает оценивать заявки, кто разбирает спорные случаи? Если система генерирует описания товаров, кто от-

вечает за юридические ошибки и тон бренда? Если ИИ помогает разработчикам писать код, как проверяется безопасность? Без таких вопросов внедрение остаётся витриной.

Трансформационный уровень наступает, когда ИИ перестаёт быть отдельным проектом и становится частью способа думать о бизнесе. Компания уже не спрашивает: куда бы нам добавить модель? Она спрашивает иначе: какие процессы мы можем построить заново, если часть анализа, генерации, поиска и контроля стала дешевле и быстрее? Здесь появляются новые продукты, новые роли, новые цепочки создания ценности. Но до этой стадии нельзя допрыгнуть лозунгом. Её приходится заработать скучной работой: данными, безопасностью, обучением, доверием и дисциплиной.

Отдельно нужно сказать об этике и безопасности. Эти слова часто звучат как обязательный раздел в конце презентации. На практике они должны стоять в начале. Не потому что юристы любят мешать инновациям. А потому что плохое решение в ИИ может масштабироваться быстрее, чем обычная ошибка сотрудника. Один неверный шаблон письма неприятен. Тысячи неверных писем уже вредят репутации. Один предвзятый менеджер плох. Предвзятая модель, встроенная в процесс, может тихо воспроизводить несправедливость каждый день.

Этическая проверка не должна быть театром. Она должна отвечать на конкретные вопросы. Какие решения мы отдаём машине, а какие оставляем человеку? Может ли кли-

ент понять, что с ним взаимодействует ИИ? Есть ли способ оспорить результат? Проверяли ли мы модель на смещение? Что мы делаем, если система уверенно ошибается? Кто имеет право остановить её работу? Как мы объясняем сотрудникам границы использования?

Безопасность тоже не сводится к паролям. В ИИ-проектах она включает доступ к данным, шифрование, хранение запросов, работу с поставщиками, проверку вывода модели, защиту от утечек, контроль прав. Малый бизнес часто думает: это вопросы корпораций, нам бы быстрее продавать. Но утечка клиентской базы одинаково неприятна и корпорации, и небольшой компании. Разница только в том, что у небольшой компании меньше запаса прочности.

Ещё одна больная тема: обучение людей. Часто его понимают как разовую презентацию. Пригласили эксперта, показали инструменты, все поаплодировали, разошлись. Через неделю половина сотрудников вернулась к старым способам, другая половина начала использовать ИИ как попало. Обучение должно быть привязано к роли и процессу. Что именно меняется в работе менеджера? Что должен делать аналитик? Как руководитель проверяет результат? Какие ошибки типичны? Где граница между помощью ИИ и профессиональной халатностью?

В хорошей программе обучения есть не только инструкции, но и практика. Сотрудники приносят свои реальные задачи, пробуют новые способы, получают обратную связь,

видят примеры удачных и плохих решений. И ещё важно: людям нужно разрешить говорить, что инструмент неудобен, путает, мешает или создаёт лишнюю работу. Если обратная связь наказуема, компания получит красивый отчёт об успешном внедрении и тихий саботаж в ежедневной работе.

Выбор рамки зависит от нескольких простых признаков. Если требования неясны и ценность надо ещё доказать, начинайте гибко. Если область регулируемая и ошибка дорого стоит, добавляйте строгие этапы проверки. Если проблема управленческая, а не только техническая, используйте лидерскую логику вроде ARIA. Если идея расплывчата и нужно понять, стоит ли за неё браться, берите Canvas и заставляйте себя отвечать на неприятные вопросы. Если проект большой, почти наверняка понадобится смесь подходов.

При этом ни одна рамка не спасает от плохой честности. Можно заполнить Canvas и всё равно подогнать ответы под заранее выбранный инструмент. Можно провести Agile-спринты и ни разу не поговорить с реальными пользователями. Можно написать водопадный план и закрыть глаза на то, что данные не годятся. Можно произнести ARIA и продолжать управлять страхом. Рамка помогает только тем, кто готов видеть факты.

Мне нравится простой тест для любого ИИ-проекта. Попросите команду за десять минут объяснить пять вещей: какую боль решаем, почему без ИИ плохо, какие данные нуж-

ны, кто будет пользоваться результатом, что мы будем делать при ошибке. Если ответы расплываются, проект ещё не готов. Не надо его хоронить. Но надо вернуться на шаг назад. Возможно, нужна не разработка, а интервью с пользователями. Не модель, а чистка данных. Не закупка платформы, а решение о владельце процесса. Не большая трансформация, а маленький честный пилот.

Руководителю в эпоху ИИ приходится держать сразу две скорости. Первая скорость быстрая: пробовать, учиться, смотреть на новые возможности, не бояться малых экспериментов. Вторая скорость медленная: выстраивать доверие, безопасность, ответственность, качество данных, обучение людей. Если есть только медленная скорость, компания будет бесконечно обсуждать риски и отстанет. Если есть только быстрая, она наломает дров и потом потратит больше сил на исправления.

Поэтому зрелая стратегия внедрения ИИ не похожа на рывок. Она похожа на серию осознанных шагов. Сначала понять, где боль. Потом проверить, есть ли данные и смысл. Потом выбрать подходящую рамку. Потом сделать малый опыт. Потом встроить удачное решение в процесс. Потом обучить людей. Потом измерять результат. Потом исправлять. Скучно? Иногда да. Зато работает.

И здесь мы возвращаемся к началу. ИИ сам по себе не даёт компании преимущества. Преимущество появляется, когда технология соединяется с ясной целью, хорошими дан-

ными, ответственными людьми и подходящей рамкой работы. У одних компаний рамка будет ближе к Agile. У других к Waterfall. У третьих к ARIA. У четвёртых к Canvas. У многих к их сочетанию. Вопрос не в названии. Вопрос в том, помогает ли выбранный способ двигаться без самообмана.

Если после этой главы оставить одну практическую мысль, я бы выбрал такую: не начинайте ИИ-проект с фразы «нам надо внедрить ИИ». Начните с фразы «у нас есть такая проблема, вот как мы сейчас её решаем, вот почему это недостаточно, вот что мы готовы проверить, вот чем будем измерять пользу, вот где проходят границы риска». После такой фразы технология получает шанс стать инструментом, а не игрушкой руководства.

Следующая часть естественно переносит разговор к тем, кто будет строить такие решения руками: разработчикам, архитекторам, инженерам данных и продуктовым командам. Руководитель может выбрать направление и задать правила игры. Но дальше начинается другая работа: жизненный цикл модели, данные, развёртывание, мониторинг, API, безопасность разработки и постоянное обслуживание. Там тоже нужны рамки. Просто они говорят уже не столько языком совета директоров, сколько языком инженерной практики.

Когда код перестаёт быть только кодом

Разработчик, который впервые берётся за искусственный интеллект в рабочем проекте, быстро замечает неприятную вещь. Обычных привычек становится мало. Можно идеально настроить репозиторий, аккуратно разложить задачи, написать чистый сервис, собрать контейнер, выкатить в staging, но вся конструкция всё равно будет шататься. Не из-за плохого кода. Из-за данных, модели, метрик, ожиданий бизнеса и той серой зоны, где никто заранее не знает, что именно получится.

В обычной разработке ошибка часто выглядит понятно. Не проходит тест, падает метод, не сходится контракт API, сервер отвечает не тем статусом. В проектах с машинным обучением всё хитрее. Модель может не падать. Она может работать. Просто работать плохо. Или сегодня хорошо, а через месяц хуже, потому что изменились данные, поведение клиентов, сезон, цены, язык пользователей, ассортимент, правила внутри компании. На экране всё зелёное, пайплайн прошёл, деплой успешный, а результат уже не тот.

Вот почему разговор о внедрении ИИ для разработчиков начинается не с выбора модной библиотеки. Он начинается с дисциплины. С вопроса, кто отвечает за данные, кто про-

веряет качество, кто следит за моделью после запуска, кто понимает, что именно она оптимизирует, и кто первым заметит, что она начала уверенно ошибаться.

Многие команды приходят к этому через боль. Сначала появляется прототип. Он впечатляет на демо. Руководитель улыбается, продуктовая команда строит планы, кто-то уже говорит про масштабирование. Потом прототип сталкивается с настоящими пользователями, грязными данными, неполными полями, странными исключениями, безопасностью, правами доступа, логами, задержками и усталыми инженерами, которым надо поддерживать это не один день, а годами. На этом месте энтузиазм часто сползает с лица.

Я видел этот переход много раз. Сначала все говорят о возможностях. Потом о данных. Потом о процессах. Потом о том, что нужен человек, который хотя бы примерно понимает весь путь от бизнес-вопроса до модели в продакшене. И тут выясняется, что разработка ИИ не живёт отдельно от старой инженерной культуры. Она требует DevOps, но с поправкой на данные. Она требует архитектуры, но с поправкой на неопределённость. Она требует мониторинга, но не только CPU, память и ошибки в логах. Нужно следить за тем, что модель делает с реальностью.

Старый DevOps был ответом на разрыв между разработкой и эксплуатацией. Разработчики писали, администраторы запускали, потом обе стороны спорили, кто виноват. DevOps заставил команды смотреть на продукт как на непрерывный

поток: код, тесты, сборка, доставка, наблюдение, исправление. Для ИИ этого мало, но без этого тоже нельзя. Если команда не умеет нормально выкатывать обычный сервис, она не станет внезапно зрелой только потому, что подключила модель.

В проектах с ИИ к привычному циклу добавляется ещё один живой слой. Данные приходят, меняются, очищаются, размечаются, попадают в обучение. Модель обучается, сравнивается с прежней версией, проверяется на тестовых наборах, выкатывается, получает новые входные данные, даёт прогнозы, ошибается, снова требует проверки. Это уже не просто доставка кода. Это доставка поведения.

Так появляется MLOps. Не как красивое слово для презентации, а как набор привычек, без которых команда однажды просыпается среди десятков моделей, непонятных датасетов, устаревших ноутбуков и скриптов, которые запускал один сотрудник, давно перешедший в другой отдел.

Самый простой признак незрелого ИИ-проекта звучит буднично: никто не может быстро ответить, на каких данных обучалась текущая модель. Не примерно, не «вроде на выгрузке за прошлый квартал», а точно. Какая версия данных, какие фильтры, какие исключения, какие признаки, какая разметка, какие параметры обучения, какая метрика, какой результат на контрольной выборке. Если ответа нет, модель уже стала чёрным ящиком не потому, что нейросеть сложная, а потому что процесс распался.

Разработчику здесь приходится быть занудой. Не токсичным, не бюрократом, а человеком, который не даёт команде потерять следы. Версионирование данных, версионирование модели, воспроизводимые эксперименты, понятные пайплайны, журнал изменений, проверяемые метрики. Звучит скучно. На практике это спасает проект.

Представьте интернет-магазин, который внедряет рекомендательную систему. На демо всё отлично: пользователю показывают товары, похожие на его интересы, средний чек растёт, менеджеры довольны. Потом приходит сезонная распродажа. Поведение покупателей резко меняется. Люди покупают не то, что обычно, а то, на что есть скидка. Модель, обученная на спокойном периоде, начинает выдавать странные рекомендации. Формально она работает. Фактически она мешает продажам.

Если у команды есть MLOps-процесс, она видит сдвиг в данных, падение качества, изменение распределения. Если процесса нет, проблему сначала замечают менеджеры поддержки, потом маркетинг, потом руководитель, потом все собираются на срочную встречу и начинают гадать, что произошло. Разница между этими двумя сценариями и есть зрелость.

Модель нельзя оставить без присмотра. Код тоже нельзя, но код хотя бы чаще ломается шумно. Модель умеет портиться тихо. Она продолжает отвечать, только ответы становятся менее полезными. Это неприятный вид сбоя, потому

что он не всегда вызывает аварию. Он просто медленно съедает доверие.

Разработчик должен заранее договориться с бизнесом о метриках. Не только технических. Точность, полнота, F1, MAE, latency, стоимость запроса, процент отказов, доля ручных исправлений, влияние на выручку, влияние на время обработки заявки. Если метрики не связаны с задачей, команда будет улучшать то, что удобно измерять, а не то, что нужно компании.

Одна из частых ошибок в ИИ-проектах выглядит так: бизнес приносит большую мечту, команда переводит её в машинную задачу слишком быстро, а потом все удивляются, что результат не помогает. Например, руководитель говорит: «Мы хотим предсказывать уход клиентов». Инженеры строят модель оттока. Она выдаёт вероятности. Всё вроде правильно. Но отдел продаж не знает, что делать с клиентом, у которого вероятность ухода 0,73. Позвонить? Дать скидку? Отправить письмо? Передать менеджеру? Если действие не встроено в процесс, прогноз превращается в украшение.

Поэтому техническая работа начинается до технической работы. Нужно разобрать задачу на части. Какое решение будет принято на основе модели? Кто его примет? В какой момент? Что изменится для клиента? Что считается успехом? Сколько стоит ошибка первого рода и сколько стоит ошибка второго рода? Можно ли объяснить результат человеку, который будет за него отвечать?

В финансовом сервисе ложное срабатывание антифрода может заблокировать честного клиента. В медицинском контексте пропущенный сигнал может стоить здоровья. В логистике неточный прогноз спроса может забить склад лишним товаром. Ошибки модели не одинаковы. Их нельзя оценивать одной красивой цифрой.

Здесь разработчик должен говорить с людьми из бизнеса на их языке. Не прятаться за терминами. Если модель классифицирует заявки, надо объяснить, какие заявки она будет отправлять на ручную проверку и почему. Если модель строит прогноз, надо показать, где она уверена, а где лучше не доверять ей без человека. Если модель генерирует текст, надо определить, кто несёт ответственность за финальную публикацию.

ИИ не отменяет API, безопасность и архитектуру. Скорее наоборот, делает их больнее. Модель нужно куда-то встроить. Её нужно вызвать, ограничить, защитить, залогировать, протестировать, обновить, откатить. Если это внешний сервис, появляются вопросы доступности, стоимости, конфиденциальности и зависимости от поставщика. Если модель развёрнута внутри, появляются вопросы инфраструктуры, ускорителей, масштабирования и поддержки.

Команда часто недооценивает стоимость эксплуатации. Прототип на ноутбуке выглядит почти бесплатным. Продакшен уже нет. Нужны вычисления, хранение данных, мониторинг, резервные сценарии, контроль доступа, аудит, доку-

ментация, поддержка инцидентов. Чем серьёзнее сфера, тем меньше права на импровизацию.

Есть ещё одна вещь, о которой не любят говорить в начале: данные редко готовы. Они разбросаны по системам, в них дубли, пропуски, устаревшие поля, разные форматы, непонятные владельцы. Иногда данные есть, но использовать их нельзя из-за регуляторных ограничений. Иногда можно, но никто не знает, кто должен дать разрешение. Иногда разрешение есть, но качество такое, что модель будет учиться на мусоре.

Разработчик, который хочет сделать хороший ИИ-проект, должен уважать эту грязную часть работы. Не считать её низкой. Подготовка данных часто решает больше, чем выбор алгоритма. Можно взять модную архитектуру и получить посредственный результат, потому что признаки собраны плохо. Можно взять простой метод и получить пользу, если данные понятны, задача правильно сформулирована, а процесс внедрения продуман.

В этом месте полезно вспомнить разницу между традиционным жизненным циклом разработки и жизненным циклом ИИ. В обычном проекте требования тоже меняются, но код пишется под относительно понятное поведение. В ИИ-проекте поведение извлекается из данных. Это значит, что данные становятся частью исходного кода, только менее послушной. Их нужно проверять так же серьёзно, как проверяют код.

Хорошая практика начинается с простого каталога. Где лежат источники данных. Кто владелец. Как часто обновляются. Какие поля критичны. Какие ограничения на использование. Какие проверки качества выполняются. Какие признаки попадают в модель. Какие исключаются и почему. Это не роскошь. Это нормальная инженерная гигиена.

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.