

ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ

третье издание



Джимшер
Челидзе
С НЕБА
НА ЗЕМЛЮ

Оглавление

Предисловие.....	8
Введение	8
Что нового в третьей редакции.....	10
Карта книги	12
Карта моих книг: пять книг — один путь	14
Часть 1. Введение в искусственный интеллект	16
Глава 1. Знакомство с ИИ.....	16
Что такое искусственный интеллект?	16
Как обучаются ИИ-модели?	20
Общие недостатки текущих решений на основе ИИ	22
Пять страхов новичка — коротко и честно	26
Резюме главы	27
Глава 2. Виды искусственного интеллекта	28
По видам решаемых задач	28
По «силе» и работе с неопределённостью.....	29
Ограничения на пути к сильному ИИ	33
Ко-пилоты, ИИ-агенты и мультиагентные системы.....	47
Что агенту позволено решать	53
Резюме главы	56
Глава 3. А что может слабый ИИ и общие тренды	56
Слабый ИИ в прикладных задачах	56
Ключевые тренды	58
Три слоя рынка ИИ: где чья игра.....	68
Резюме главы	69
Глава 4. Генеративный ИИ.....	70
Что такое генеративный искусственный интеллект?	70
Проблемы генеративного ИИ	71
Ограничения ИИ, которые приводят к проблемам.....	74
Что же делать?.....	78
Резюме главы	81
Глава 5. Большие языковые модели (LLM)	84
Феномен LLM	84
Что такое RAG?.....	85

Fine-tuning	86
Как работает LLM?	87
Прослеживаемость источника: требование, а не пожелание	90
Резюме главы	90
Глава 6. Промпт-инжиниринг: дисциплина постановки задач для ИИ	91
Промпт как управленческий инструмент	91
Базовая структура запроса	92
Формула качества задачи	95
Четыре шаблона для разных ситуаций	98
Техники повышения качества работы с ИИ	100
Системный промпт и сборка промпта с помощью ИИ	104
Стартовый набор: 10 готовых промптов руководителя	105
Использование вопросов при работе с ИИ: от плана выполнения задачи к разбору ответов	106
Чек-лист и типичные ошибки	113
Резюме главы	115
Глава 7. Регулирование ИИ	117
Почему развитие ИИ вызывает беспокойство?	117
Карта регулирования: инициативы 2023–2026	119
Российский закон об ИИ: что на самом деле приняли	123
Прогноз на будущее	126
Кто пишет правила рынка передовых моделей	129
Резюме главы	134
Глава 8. Безопасность ИИ	135
Недопустимые события, которые надо исключить	135
Сценарии использования ИИ злоумышленниками	136
Из-за чего эти сценарии могут реализоваться?	138
Краш-тесты и требования к ИИ-решениям	139
Безопасность как организационный барьер: экономика периметра	152
Резюме главы	153
Глава 9. Системы поддержки принятия решений, или Цифровые советники	157
Что такое цифровые советники	157
Как устроены СППР и что с ними не так	159
Что делать и куда бежать (горизонт 5–20 лет)	165

Резюме главы	167
Глава 10. Нейроморфные и квантовые ИИ	168
Нейроморфные системы	169
Квантовые системы	172
Резюме главы	176
Часть 2. Искусственный интеллект в системном подходе	178
Глава 11. Стратегия и организационная структура	178
Влияние ИИ	178
Влияние на ИИ	180
Практический пример	182
А что Россия: профиль смекалки	202
Резюме главы	205
Глава 12. Бизнес-процессы	206
Влияние ИИ	207
Влияние на ИИ	213
Резюме главы	214
Глава 13. Управление проектами, изменениями и продуктами	214
Влияние ИИ	216
Влияние на ИИ	219
7 смертных грехов в цифровизации (в том числе ИИ)	221
Резюме главы	232
Глава 14. Коммуникация	233
Влияние ИИ	233
Влияние на ИИ	236
Резюме главы	237
Глава 15. Цифровые технологии, работа с данными и кибербезопасность	238
Цифровые технологии	238
ИТ-системы	245
Работа с данными	253
Резюме главы	253
Глава 16. Практики регулярного менеджмента	254
Управленческое планирование	255
Делегирование полномочий и задач	255
Работа с обратной связью подчинённых	255

Организация совещаний	255
Работа с кадрами и решения относительно людей, оценка эффективности работы сотрудников, развитие кадров	256
Три сценария «ПРМ × ИИ»: с чего начать	256
Резюме главы	257
Глава 17. Бережливое производство и теория ограничений систем (ТОС).....	257
Влияние ИИ	258
Взаимное влияние	260
Резюме главы	272
Глава 18. Влияние ИИ на процессы функциональных направлений.....	273
Цифровой бизнес	273
Цифровой маркетинг	277
Цифровое производство	277
Цифровая аналитика	278
Цифровые продукты.....	279
Новые способы работы.....	279
Инженерия и конструирование: от генерации к верификации	279
Резюме главы	281
Резюме второй части	281
Общие выводы и прогнозы	281
Часть 3. Практические примеры и рекомендации по применению искусственного интеллекта.....	283
Глава 19. ИИ для малого бизнеса и предпринимателей	283
Применение ИИ в решениях 1С.....	283
Маркетплейсы специализированных ИИ-решений.....	286
Цифровой советник для управления проектами	289
Генеративный ИИ для создателей контента	292
Создание ассистентов	294
Бьюти-ассистент	295
Рекомендации для малого и среднего бизнеса	295
Резюме главы	298
Глава 20. Примеры и рекомендации для среднего и крупного бизнеса	299
Генеративный ИИ при заключении сделок	300
Промышленность и инфраструктура	301
Клиентские сервисы	311

Планирование и кадры.....	316
Рекомендации для крупного бизнеса	320
Резюме главы	327
Глава 21. ИИ в работе самого руководителя.....	328
Задачи руководителя: что можно отдать ИИ	328
Мой личный контур: как это устроено.....	329
Три кейса из практики	335
Что даёт союз руководителя с ИИ	342
Граница союза: непохожесть остаётся вам.....	343
Как это превратилось в продукт	344
Резюме главы	346
Глава 22. Как вести проекты по внедрению технологий	347
Жизненный цикл проекта.....	349
Определение целей проекта и возможных эффектов	350
Роли в проекте	351
Три алгоритма — коротко	354
Резюме главы	355
Глава 23. Системный подход к внедрению и дорожная карта	356
О волне хайпа на искусственном интеллекте	356
Системный подход к внедрению ИИ	357
Дорожная карта и связи с другими проектами	357
Резюме главы	358
Глава 24. Что должен знать и уметь руководитель в эпоху ИИ.....	359
Введение.....	359
Модель компетенций	359
Матрица компетенций	368
Компетенции эпохи верификации.....	371
Резюме главы	372
Заключение.....	373
Кто под ударом — и что с этим делать.....	373
Преодоление барьеров внедрения «сильных» моделей.....	375
Преимущества специализированных моделей	376
Направления развития в области специализированного ИИ.....	376
Фундамент — качество данных и защита	377

Практическая реализация.....	378
Постскрипtum 2026: четыре сигнала зрелости рынка	380
Чем я хочу закончить	387
Что дальше	388
Приложения	389
Приложение А. Глоссарий ключевых терминов	389
Приложение Б. Чек-лист: где нужен ИИ, а где классическая автоматизация	391
Приложение В. Первые 7 дней с ИИ: план для новичка	392
Приложение Г. Полный план действий: 24 шага из 24 глав	393

Предисловие

Введение

2023 год стал началом третьего и пока самого жаркого «лета» искусственного интеллекта. Тогда вопросы звучали остро: что делать разработчикам и бизнесу? Стоит ли сокращать людей и полностью доверяться новой технологии — или это всё наносное, и вслед за летом придёт новая ИИ-зима? Как ИИ сочетается с классическими инструментами и практиками менеджмента — управлением проектами, продуктами, бизнес-процессами? Неужели ИИ ломает все правила игры?

К моменту этой, третьей, редакции — в 2026 году — на часть вопросов рынок успел ответить сам. Я обновил книгу с учётом этих ответов: новых данных, исследований, трендов и примеров. Но суть осталась прежней — мы погружаемся в эти вопросы, и в процессе размышлений и анализа различных факторов я постараюсь найти на них ответы.

Скажу сразу: здесь будет минимум технической информации, и даже более того — возможны технические неточности. Но эта книга и не про то, какой тип нейросетей выбрать для той или иной задачи. Я предлагаю вам взгляд аналитика, управленца и предпринимателя на то, что происходит здесь и сейчас, чего ожидать глобально в ближайшем будущем и к чему готовиться.

Кому и чем полезна эта книга?

- **Собственникам и топ-менеджерам компаний.** Они поймут, что такое ИИ, как он работает, куда идут тренды и чего ожидать, — а значит, избегут главной ошибки в цифровизации — искажённых ожиданий. Это позволит им лидировать в этих направлениях, минимизируя затраты, риски и сроки.
- **ИТ-предпринимателям и основателям стартапов.** Они увидят, куда движется отрасль, какие ИТ-продукты и интеграции стоит разрабатывать и с чем придётся столкнуться на практике.

- **Техническим специалистам из ИТ.** Они посмотрят на развитие не только с технической, но и с экономической и управленческой точки зрения. Поймут, почему до 95 % внедрений ИИ не дают измеримого финансового эффекта (MIT, 2025), — возможно, это поможет им в карьерном развитии.
- **Обычным людям.** Они поймут, что их ждёт в будущем и стоит ли бояться, что ИИ их заменит. Спойлер: под угрозой прежде всего оказались специалисты творческих профессий.

Наше путешествие пройдёт через три большие части. Сначала разберём, что такое ИИ: с чего всё начиналось, какие у него проблемы и возможности, куда идут тренды и что нас ждёт впереди. Это теория и долгосрочные перспективы. Затем рассмотрим синергию ИИ с инструментами системного подхода: как они влияют друг на друга и в каких сценариях ИИ уже применяется в ИТ-решениях. В конце сосредоточимся на практике — примерах и рекомендациях.

По моей любимой традиции в книге будут QR-коды (для печатной версии) и активные гиперссылки (для электронной) на полезные статьи и интересные материалы.

Что касается самого ИИ, то в написании книги он применялся для демонстрации его работы — такие фрагменты выделены отдельно, — а также для поиска идей. Уровень технологии пока недостаточен, чтобы использовать сгенерированное ИИ как готовый материал. Отмечу и то, что мой основной язык — русский, поэтому все запросы к ИИ я делаю на русском.

Из практики

Многие выводы, примеры и подходы в этой книге выросли из моей собственной практики — в том числе из создания продукта «Сокол». По ходу глав я ссылаюсь на него как на живой кейс встраивания ИИ в реальные управленческие задачи.

Завершить предисловие я хочу благодарностями людям, которые помогли мне:

- Алисе и сыну Валерию;
- моим родителям;
- моему коучу Евгению Бажову;
- моей команде, в особенности Александру Перемышлину;
- моим партнёрам и клиентам, которые давали мне пищу для размышлений, в особенности Кириллу Неелову;
- моим коллегам по проектам.

Что нового в третьей редакции

Если вы читали первое или второе издание — вот что изменилось и почему стоит прочесть третье.

Две новые главы. «Промпт-инжиниринг: дисциплина постановки задач для ИИ» (Глава 6): шаблоны, техники, готовые промпты руководителя, формула качества задачи — постановка × объём × данные × модель — и раздел «Использование вопросов при работе с ИИ: от плана выполнения задачи к разбору ответов»: пять вопросов, которыми принимают план до начала работ, и девять типов вопросов для проверки ответов ИИ с исследованиями угодливости моделей. И «ИИ в работе самого руководителя» (Глава 21): личный контур, три кейса из моей практики — от разбора рабочего диалога до решения о крупной сделке, — четыре уровня настройки ассистента (от общей инструкции аккаунта и реестра проектов до постановки задачи в чате), интервью вместо шаблона и мета-кейс «статья за два дня». Плюс в Главе 5 — контекстное окно и схема «чат + файл».

Личные приёмы из практики, которых не было в прежних изданиях: как точность термина меняет качество ответа и расход токенов; правило «не останавливаться на первом ответе» — прогоны с разных сторон вместо одного запроса; хитрость «попроси модель разбить задачу на шаги» и утверди план; история о том, как самая мощная модель испортила простую задачу, а я потратил на это час; что делать, когда ассистент «тупеет»; как выглядит день руководителя без привычного инструмента.

Плюс семь врезок «Из практики» и разборы от Bridgewater до отчёта Anthropic по 400 000 рабочих сессий.

Безопасность автономных агентов. Глава 8: три кейса 2026 года — от лабораторных прогонов OpenAI и Anthropic до первого потребительского, где бытовой ИИ-агент взломал сайт фитнес-клуба по дороге к обычной задаче, — и типология «пять механик отказа» с защитой от каждой. Общий вывод главы: границу агента задаёт архитектура, а не сообразительность модели. Отдельно — кейс трёх агентов в одном пространстве (Anthropic, август 2026): конфликт возник без лишних полномочий и без внешнего противника, и это добавляет к границам второй управленческий вопрос — координацию между агентами.

Факты середины 2026 года. Полная актуализация: экономика ИИ — от «вопроса на 600 миллиардов» до триллионных капвложений, новый рейтинг галлюцинаций и карта регулирования 2023–2026, включая российский закон об ИИ.

Прогнозы, которые сбылись. События, подтвердившие прогнозы первых изданий: экспортные ограничения на передовые модели, разворот рынка к специализированным моделям, «Постскрипtum 2026» с четырьмя сигналами зрелости рынка.

Стратегия и Россия. В Главе 11 к разбору китайской связки «ИИ+» и «1397» добавлен раздел «А что Россия: профиль смекалки»: цифры GII 2025, историческая колея развилок — от атомного проекта до ОГАС — и авторская позиция «ИИ закрывает наше слабое — и умножает наше сильное»; две инфографики и выход на полную статью по QR-коду.

Компетенции эпохи ИИ. Глава 24 дополнена разбором пяти дефицитных компетенций специалистов, чью работу ИИ уже перестраивает, выводом о перепроектировании обучения под верификацию вместо генерации и новой дефицитной фигурой — эрудитом с ИИ-инструментом.

Визуальный слой. Больше десяти новых схем и инфографик; среди них — S-кривая зрелости ИИ, три слоя рынка, экономика ИИ в цифрах, таймлайн регулирования, ИТ-ландшафт со слоем ИИ, четыре сигнала рынка и постер принципов.

Навигация и подача. Манифест «7 принципов „С неба на землю“» — вся книга на одной странице; карта серии «пять книг — один путь» и раздел «Что дальше»; врезки «Простыми словами», расширенный глоссарий, чек-лист «где нужен ИИ, а где классическая автоматизация», выводы и действия в каждой главе, «Первые 7 дней с ИИ» и сводный план из 24 шагов.

Если времени мало: начните с Главы 21, Главы 6 и Постскриптума — там сосредоточена бóльшая часть нового.

Карта книги

Для тех, кто занят и не готов читать всю книгу, приведу её краткую карту: на какой вопрос отвечает каждая глава.

Глава	На какой вопрос отвечает
Часть 1. Введение в искусственный интеллект	
Глава 1. Знакомство с ИИ	Что такое ИИ и когда он нужен, а когда дешевле классическая автоматизация
Глава 2. Виды ИИ	Какой бывает ИИ; почему сильный ИИ не завтра; уровни автономности агентов
Глава 3. Слабый ИИ и тренды	Что уже работает и куда идёт рынок (9 ключевых трендов)
Глава 4. Генеративный ИИ	Возможности и реальные ограничения генеративного ИИ (GenAI)
Глава 5. LLM	Как устроены языковые модели, RAG и дообучение (fine-tuning)
Глава 6. Промпт-инжиниринг	Как получать качественные ответы и стандартизировать промпты
Глава 7. Регулирование ИИ	Куда движется регулирование и что закладывать заранее
Глава 8. Безопасность ИИ	AI-специфичные риски и как защищаться
Глава 9. Цифровые советники	Перспективы и ограничения систем поддержки принятия решений (СППР), как их внедрять

Глава 10. Нейроморфные и квантовые ИИ	Горизонт развития «железа» ИИ
Часть 2. ИИ в системном подходе	
Глава 11. Стратегия и оргструктура	Как ИИ меняет стратегию и оргструктуру — и как они определяют успех внедрения
Глава 12. Бизнес-процессы	Что ИИ делает с процессами и почему хаос нельзя автоматизировать
Глава 13. Проекты, изменения, продукты	Как ИИ помогает вести проекты; «7 грехов» цифровизации
Глава 14. Коммуникация	Где ИИ экономит 50–80% времени на обработке информации
Глава 15. Технологии и данные	С чем сочетается ИИ и почему данные — потолок его пользы
Глава 16. Практики регулярного менеджмента	Как встроить ИИ в управленческие ритуалы
Глава 17. Lean и ТОС	Почему «Lean перед ИИ» и как они усиливают друг друга
Глава 18. Функциональные направления	Что ИИ меняет в маркетинге, производстве, аналитике, продуктах
Часть 3. Практика и рекомендации	
Глава 19. ИИ для малого бизнеса	С чего начать малому и среднему бизнесу (МСБ) — без бюджета и разработчиков
Глава 20. Средний и крупный бизнес	Библиотека отраслевых кейсов и принципы масштабирования
Глава 21. ИИ в работе руководителя	Личный контур: дневник, решения, коммуникации — и короткая петля роста
Глава 22. Проекты внедрения	Три типа проектов и управленческая рамка ведения
Глава 23. Системный подход и дорожная карта	Как связать ИИ со стратегией цифровизации и выстроить дорожную карту
Глава 24. Компетенции руководителя	Что должен знать и уметь руководитель в эпоху ИИ
Заключение + Постскрипtum 2026	Четыре сигнала зрелости рынка и профиль выигрывающего ИИ-продукта

Если вы совсем новичок в ИИ: начните с короткого маршрута — Главы 1 → 3 → 6 → 14 → 19 и Заключение. Этого достаточно для уверенного старта: что такое ИИ, что он уже умеет, как ставить ему задачи, где он экономит время и с чего начать малому бизнесу. Главы 4–5 читайте по мере появления вопросов «как это устроено», а Главу 10 (квантовые и нейроморфные системы) смело оставьте на потом.

Карта моих книг: пять книг — один путь

Эта книга — часть серии, и у каждой книги в ней своя работа. Чтобы вы не собирали пазл сами, поделюсь коротким навигатором.

Книга	Какой вопрос закрывает	Когда читать
Цифровая трансформация для директоров и собственников. Часть 1. Погружение (кратко: «ЦТ-1. Погружение»)	Что такое цифровизация, технологии и ИТ-системы — фундамент и словарь	Если только входите в цифровую повестку
Цифровая трансформация для директоров и собственников. Часть 2. Системный подход (кратко: «ЦТ-2. Системный подход»)	Процессы, ТОС, Lean, проекты, коммуникации — управленческий инструментарий	Перед любой трансформацией — и параллельно с Частью 2 этой книги
Цифровая трансформация для директоров и собственников. Часть 3. Кибербезопасность (кратко: «ЦТ-3. Кибербезопасность»)	Как не потерять бизнес: ИБ для руководителя	Вместе с Главой 8 этой книги
Искусственный интеллект. С неба на	Что такое ИИ, что он умеет и как о нём думать	Входная точка в ИИ

землю (кратко: «ИИ-1. С неба на землю») (вы здесь)		
Искусственный интеллект. Практическое руководство для внедрения (кратко: «ИИ-2. Практическое руководство»)	Как внедрять: готовность, портфель, пилоты, ROI, риски	Когда решение «делаем» принято

Три маршрута.

- Собственнику МСБ: ИИ-1 → ИИ-2 (гл. 17) → точно ЦТ-1.
- Руководителю трансформации (CDTO): ЦТ-1 → ЦТ-2 → ИИ-1 → ИИ-2 → ЦТ-3.
- Функциональному руководителю: ИИ-1 (маршрут новичка) → профильные главы Части 2.

Часть 1. Введение в искусственный интеллект

Часть 1 отвечает на вопрос «что такое ИИ и чего ожидать». Главы можно читать подряд или выборочно — по своей задаче.

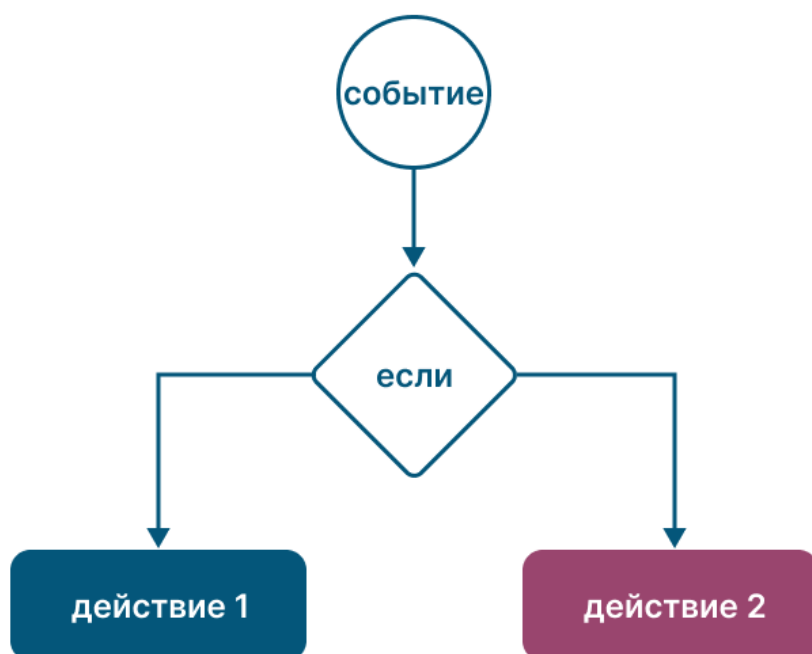
Глава 1. Знакомство с ИИ

Что такое искусственный интеллект?

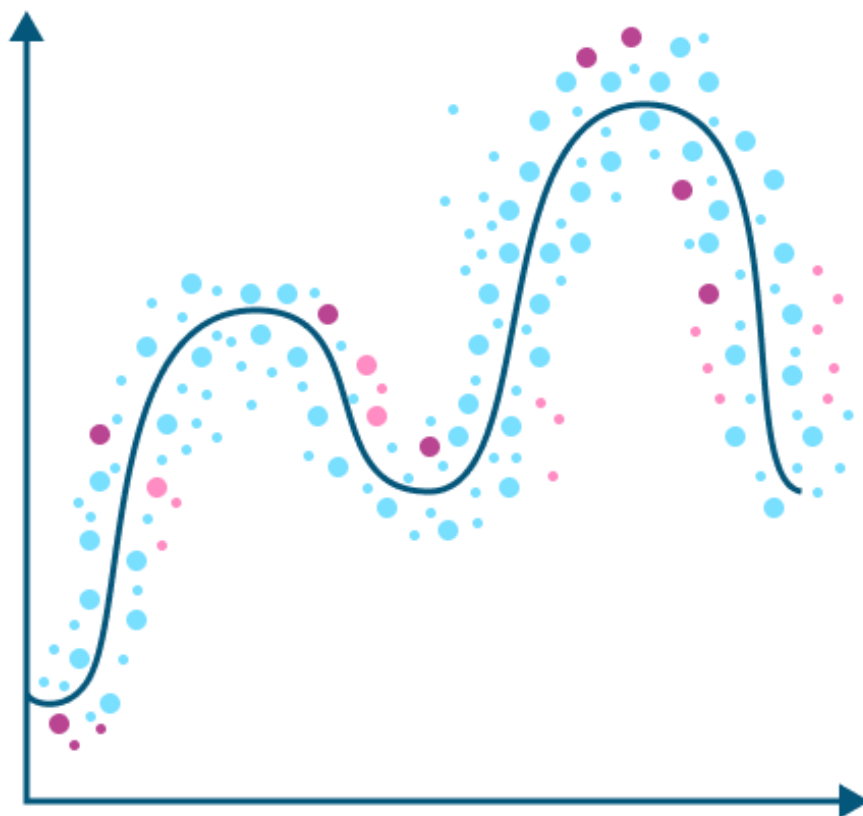
Начнём с того, что единого определения не существует. Если спросить об этом ИИ, то каждая модель даст свой собственный ответ с разной глубиной погружения.

Мне же, как человеку, ближе самое простое и понятное определение — вокруг него сходятся и ответы самих ИИ-моделей. ИИ — это любой математический метод, который имитирует человеческий или другой природный интеллект.

То есть ИИ — это огромное количество решений, в том числе примитивные математические алгоритмы, экспертные системы на базе правил и современные решения на базе статистики.



Классические решения на базе правил и логик



Решения на базе анализа данных и статистики, оценки вероятностей.

В этой книге мы будем говорить именно о направлении, которое работает со статистикой и поиском взаимосвязей.

И хотя это направление родилось где-то в 1950-х годах XX века, нас в первую очередь интересует то, что мы понимаем под ним сегодня, в 2020-х. Здесь есть три основных направления.

1. **Нейросети** — математические модели, созданные по подобию нейронных связей мозга живых существ. Собственно, мозг человека — это суперсложная нейросеть. Ключевая особенность: наши нейроны не ограничиваются состояниями «включён/выключен», у них гораздо больше параметров. Оцифровать и применить их в полной мере пока не получается.
2. **Машинное обучение (ML)** — статистические методы, которые позволяют компьютеру решать задачу всё лучше — по мере накопления опыта и дообучения. Это направление известно с 1980-х годов.

3. Глубокое обучение (DL) — это не только обучение с учителем, когда человек показывает машине правильные ответы. Это ещё и самообучение систем — без участия человека. Это одновременное использование различных методик обучения и анализа данных. Направление развивается с 2010-х и считается наиболее перспективным для творческих задач и задач, где сам человек не понимает чётких взаимосвязей. Это и есть все современные популярные модели, о которых мы много слышим. Но здесь мы вообще не можем предсказать, к каким выводам придёт нейросеть: манипулировать можно лишь тем, какие данные мы «скармливаем» модели на входе.

Как соотносятся ИИ, ML, DL и современные модели



Когда нужен ИИ, а когда — классическая автоматизация?

Если коротко, ИИ работает лучше всего, когда одновременно выполняются три условия: есть данные, в задаче много вариативности (исключений) и её можно автоматизировать.

Классическая автоматизация предпочтительна, когда:

- задача имеет ограниченную сложность и вариативность;

- критичны интерпретируемость, скорость отклика и надёжность;
- нет данных для машинного обучения.

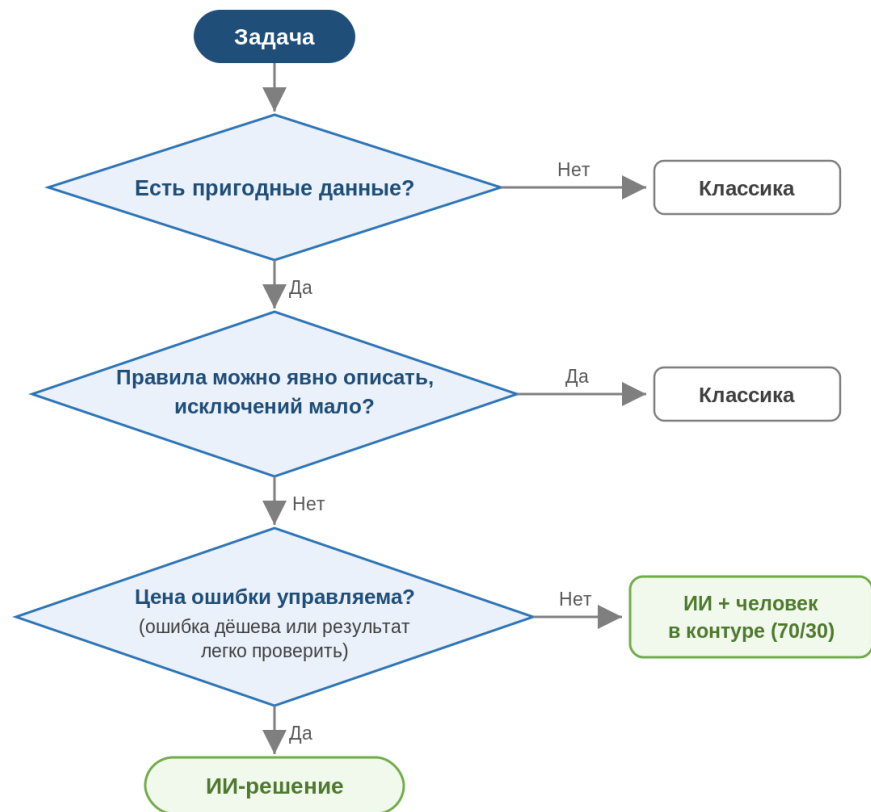
ИИ-решения предпочтительны, когда:

- задача имеет высокую сложность или большое число вариаций и факторов, а правила невозможно явно сформулировать или они слишком сложны;
- доступны большие массивы данных, из которых можно выявить закономерности;
- требуется анализ неструктурированных данных (тексты, изображения, аудио), где традиционные алгоритмы малоэффективны;
- необходима быстрая адаптация к новым данным и поиск закономерностей;
- решение зависит от вероятностных оценок.

Ключевое различие: традиционные алгоритмы vs ИИ-модели

Аспект	Традиционные алгоритмы	ИИ-модели
Как работают	Правила = алгоритм	Правила = данные
Когда применять	Ограниченная сложность	Высокая сложность + много факторов
Нужны ли примеры?	Нет	Да, 100+ примеров
Интерпретируемость	Высокая	Низкая («чёрный ящик»)
Примеры	Расчёт зарплаты, маршрутизация	Классификация звонков, выявление мошенничества, рекомендации

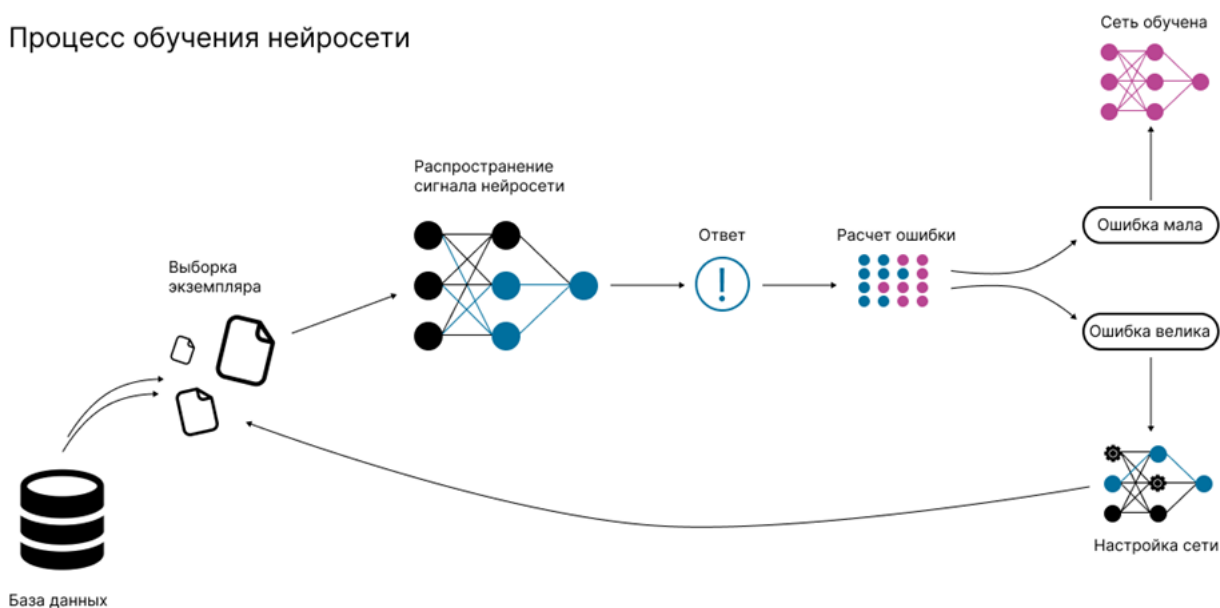
Где нужен ИИ, а где классическая автоматизация



Как обучаются ИИ-модели?

Сейчас множество прикладных ИИ-моделей для бизнеса обучается с учителем: человек задаёт входную информацию, нейросеть возвращает ответ, после чего человек сообщает ей, верно она ответила или нет. И так раз за разом.

Процесс обучения нейросети



Процесс обучения нейросети

Подобным образом работают и так называемые «капчи» (CAPTCHA, Completely Automated Public Turing test to tell Computers and Humans Apart) — графические тесты на сайтах, определяющие, кто перед ними: человек или компьютер. Это когда вам показывают разделённую на части картинку и просят указать зоны с велосипедами или ввести затейливо отображённые цифры и буквы. Кроме основной задачи (тест Тьюринга), эти данные потом используются для обучения ИИ.

При этом существует и обучение без учителя, когда система ищет закономерности сама, без обратной связи от человека. Это самые сложные проекты, но они позволяют решать и самые сложные, творческие задачи. ChatGPT и аналогичные решения — продукт именно такого обучения, точнее его разновидности, которую называют самообучением (self-supervised): модель закрывает себе часть текста и учится угадывать скрытое, то есть сама выставляет себе оценку. Ручная разметка здесь не нужна — но состав и качество данных решают всё.

Точнее, на практике это гибридный подход: сначала машина обучается сама на огромных массивах данных, а затем люди корректируют результат — оценивают ответы, размечают удачные и неудачные, «дообучая» модель обратной связью (в индустрии это называется RLHF —

обучение с подкреплением на обратной связи от человека). Машина делает черновую работу, человек ставит оценки.

Общие недостатки текущих решений на основе ИИ

Фундаментально все решения на базе ИИ на текущем уровне развития имеют общие проблемы.

- **Объём данных для обучения.** Нейросетям нужны огромные массивы качественных и размеченных данных. Если человек учится отличать собак от кошек на паре примеров, то ИИ нужны тысячи.
- **Зависимость от качества данных.** Любые неточности в исходных данных сильно сказываются на результате. Нейросети просто собирают данные, не анализируя факты и их связанность.
- **Этическая составляющая.** Для ИИ нет этики — только математика и выполнение задачи. Отсюда сложные этические проблемы: например, кого сбить автопилоту в безвыходной ситуации — взрослого, ребёнка или пенсионера? Подобных споров бесчисленное множество. Для ИИ нет ни добра, ни зла, как нет и «здорового смысла».
- **Качество и достоверность контента.** Нейросети не могут оценить данные на реальность и логичность и склонны к генерации некачественного контента и ИИ-галлюцинаций. Ошибок много, и это порождает две проблемы.

Первая — деградация поисковиков. ИИ создал столько некачественного контента, что поисковые системы начали деградировать: его стало настолько много, что он доминирует. Особенно «помогают» SEO-оптимизаторы, набрасывающие популярные запросы для продвижения.

Вторая — деградация самих ИИ-моделей. Генеративные модели используют интернет для дообучения. Люди, применяя ИИ и не проверяя за ним, наполняют интернет некачественным контентом, а ИИ снова учится на нём. Получается замкнутый круг, ведущий ко всё бóльшим проблемам.

Как возникает деградация данных и моделей



Также по QR-коду и гиперссылке доступна статья на эту тему



[The AI feedback loop: Researchers warn of 'model collapse' as AI trains on AI-generated content](#)

Но есть важное «но». Круг замыкается не потому, что текст написала машина. А потому, что этот текст не с чем сверить. Ошибку никто не ловит — и она переходит дальше, в следующую модель.

Если правильный ответ известен заранее, всё меняется. Задачу можно придумать целиком: и условие, и ответ к нему. Тогда проверять машину не нужно человеку — сразу видно, верно она решила или нет. Такие данные модель не портят, а учат.

Пример. В июне 2026 года OpenAI выложила тест GeneBench-Pro: 129 исследовательских задач по вычислительной биологии. Каждую собрали

искусственно, и для каждой заранее известен верный ответ. Год назад лучшая модель компании решала меньше 5% задач предыдущей версии теста. Сейчас — 31,5%.

Разница между ядом и лекарством здесь одна. Есть эталон, с которым можно сверить ответ, — данные полезны. Нет эталона — начинается тот самый замкнутый круг.

Осознавая проблему, Google DeepMind вместе с Jigsaw провела исследование: учёные разобрали около двухсот описанных в СМИ случаев использования ИИ не по назначению (январь 2023 — март 2024) и составили их классификацию. Главный вывод оказался неожиданным: в девяти случаях из десяти злоумышленники не «взламывали» ИИ, а пользовались им ровно по прямому назначению. Они генерировали изображения реальных людей, фальшивые доказательства и массовый контент от лица несуществующих личностей. Ломать ничего не пришлось.

С тех пор это перестало быть темой научных статей и стало строкой в бюджете потерь. По оценке Deloitte, ущерб от мошенничества с применением генеративного ИИ в США вырастет с 12,3 млрд долларов в 2023 году до 40 млрд к 2027-му. Подробнее о том, как это работает и что с этим делать, — в главе о безопасности.

- **Качество «учителей».**

Почти все нейросети обучают люди: формируют запросы и дают обратную связь. И здесь много ограничений — кто и чему учит, на каких данных, для чего?

- **Готовность людей.**

Стоит ожидать огромного сопротивления тех, чью работу заберут нейросети.

- **Страх перед неизвестным.**

Рано или поздно нейросети станут умнее нас. Люди боятся этого, а значит, будут тормозить развитие и накладывать ограничения.

- **Непредсказуемость.**

Иногда всё идёт как задумано, а иногда даже создатели из всех сил пытаются понять, как работают алгоритмы. Отсутствие предсказуемости крайне затрудняет поиск и исправление ошибок. Мы только учимся понимать то, что сами создали.

- **Ограничение по виду деятельности.**

На момент подготовки третьей редакции (2026) весь массовый ИИ — слабый (мы разберём этот термин в следующей главе). Алгоритмы хороши для целенаправленных задач, но плохо обобщают знания: ИИ, обученный играть в шахматы, не сыграет в шашки. Глубокое обучение плохо справляется с данными, отклоняющимися от учебных примеров. Чтобы эффективно использовать тот же ChatGPT, нужно быть экспертом в отрасли и формулировать осознанный, чёткий запрос.

- **Затраты на создание и эксплуатацию.**

Создание нейросетей требует много денег, и лучше всего это видно на дистанции в несколько лет. Обучение GPT-3 в 2020 году, по разным оценкам, стоило от 1,4 до 4,6 млн долларов — сумма, вполне подъёмная для крупной компании. GPT-4 в 2023-м обошёлся уже примерно в 78 млн, Gemini Ultra — около 191 млн, а Llama 3.1 на 405 млрд параметров — порядка 170 млн (оценки Stanford AI Index и Epoch AI). Обучение моделей поколения 2026 года оценивают в 200–500 млн долларов за один запуск, а к концу 2027-го аналитики ждут выхода за миллиард.

Итого: за шесть лет стоимость обучения передовой модели выросла примерно в сто раз. Причём растёт она устойчиво — в 2,4–3 раза ежегодно с 2016 года. Это и есть та самая верхняя часть S-кривой в денежном выражении: каждый следующий процент качества стоиткратно дороже предыдущего.

Экономика ИИ для руководителя

Четыре цифры, которые надо знать до старта проекта — данные на середину 2026 года



Считайте ТСО и эффект ДО старта — а не после пилота

Источники: Stanford AI Index, Epoch AI, отчётность и прогнозы компаний, 2024–2026

Похожая история с эксплуатацией. Для обработки запросов GPT-3 требовалось более 30 тысяч графических процессоров NVIDIA A100, а на электроэнергию уходило около 50 тысяч долларов в день. Сегодня счёт идёт на сотни тысяч ускорителей нового поколения, а инференс одной только OpenAI оценивают в десятки миллионов долларов ежедневно.

Что из этого следует для вас. Ровно один вывод: обучать свою фронтир-модель (передовую модель уровня ChatGPT) — не ваша игра, и это не признание слабости, а трезвый расчёт. Ваша игра начинается там, где заканчивается их: в данных, процессах и сценариях, которые никто, кроме вас, не оцифрует. Модель вы арендуете или купите, а конкурентное преимущество — только постройте сами.

Повторюсь: это общие недостатки для всех ИИ-решений. Дальше мы вернёмся к этой теме ещё несколько раз и разберём их на более прикладных примерах.

Пять страхов новичка — коротко и честно

1. «Мои данные украдут».

Риск управляем, и правило простое: не отправляйте в публичный сервис то, что нельзя показать постороннему (подробнее — Глава 8). При этом

многие сервисы уже внедряют шифрование пользовательских данных (полностью или только персональные данные).

2. «ИИ меня заменит».

Рабочая модель — 70/30: рутину забирает ИИ, решения и ответственность остаются у человека. Заменяет не ИИ, а люди, которые научились им пользоваться.

3. «Это дорого».

Старт — бесплатно или за небольшую сумму. Дорогим бывает внедрение без цели, а не инструмент.

4. «Это сложно, я не технарь».

Главный навык — формулировать задачу и проверять результат (этому посвящена Глава 6). Исследования 2026 года показывают: не-технари решают задачи с ИИ почти так же успешно, как программисты. А на практике бывает и наоборот: тот, кто хорошо знает предметную область, говорит на её языке и может оценить ответ, работает с ИИ даже лучше ИТ-специалистов.

5. «Поздно начинать».

Массовое внедрение только разворачивается: три четверти работников уже пробуют ИИ, но системно используют единицы. Форa ещё есть — и она у дисциплинированных, а не у ранних.

Резюме главы

- ИИ — это не магия, а математика на статистике; есть три направления: нейросети, машинное обучение и глубокое обучение.
- ИИ оправдан там, где есть данные, высокая вариативность и задачу можно автоматизировать; иначе дешевле и надёжнее классическая автоматизация.
- У всех ИИ-решений общие фундаментальные ограничения: объём и качество данных, непредсказуемость, затраты на создание и эксплуатацию.

Что с этим делать: Прежде чем брать ИИ под задачу, прогоните её через критерий «ИИ vs классика» и проверьте, есть ли у вас пригодные данные.

Вопрос для рефлексии: Для какой вашей задачи ИИ действительно оправдан, а где вы тянете его «по моде»?

Глава 2. Виды искусственного интеллекта

По видам решаемых задач

ИИ удобно классифицировать по нескольким основаниям. Начнём с видов решаемых задач.

- **Генеративный.**

Создаёт новый контент (изображения, текст, звук и т. д.) по запросу пользователя: чат-боты для общения, генерация картинок, видео, музыки, рекомендаций. Ему в книге посвящён отдельный раздел.

- **Классифицирующий.**

Относит данные к категориям по заданным критериям. Обучается на размеченных данных и использует разные алгоритмы для принятия решений: отделение кошек от собак, проверка изделия на соответствие габаритам, наличие СИЗ у человека в опасной зоне, является ли письмо спамом, определение состояния здоровья и т. д.

- **Предиктивный / рекомендательный.**

Предсказывает будущие события или результаты на основе имеющихся данных и статистики: скорый отказ оборудования, будущая урожайность, наследственные заболевания. Иногда отдельно выделяют рекомендательный ИИ, который не просто предсказывает, а готовит рекомендации, что делать, — но я отношу его к предиктивному.

Тем, кому интересно техническое погружение (виды алгоритмов и какие задачи они решают), рекомендую хорошую статью о видах машинного обучения — где какой класс алгоритмов предпочтителен.



[*Машинное обучение для людей*](#)

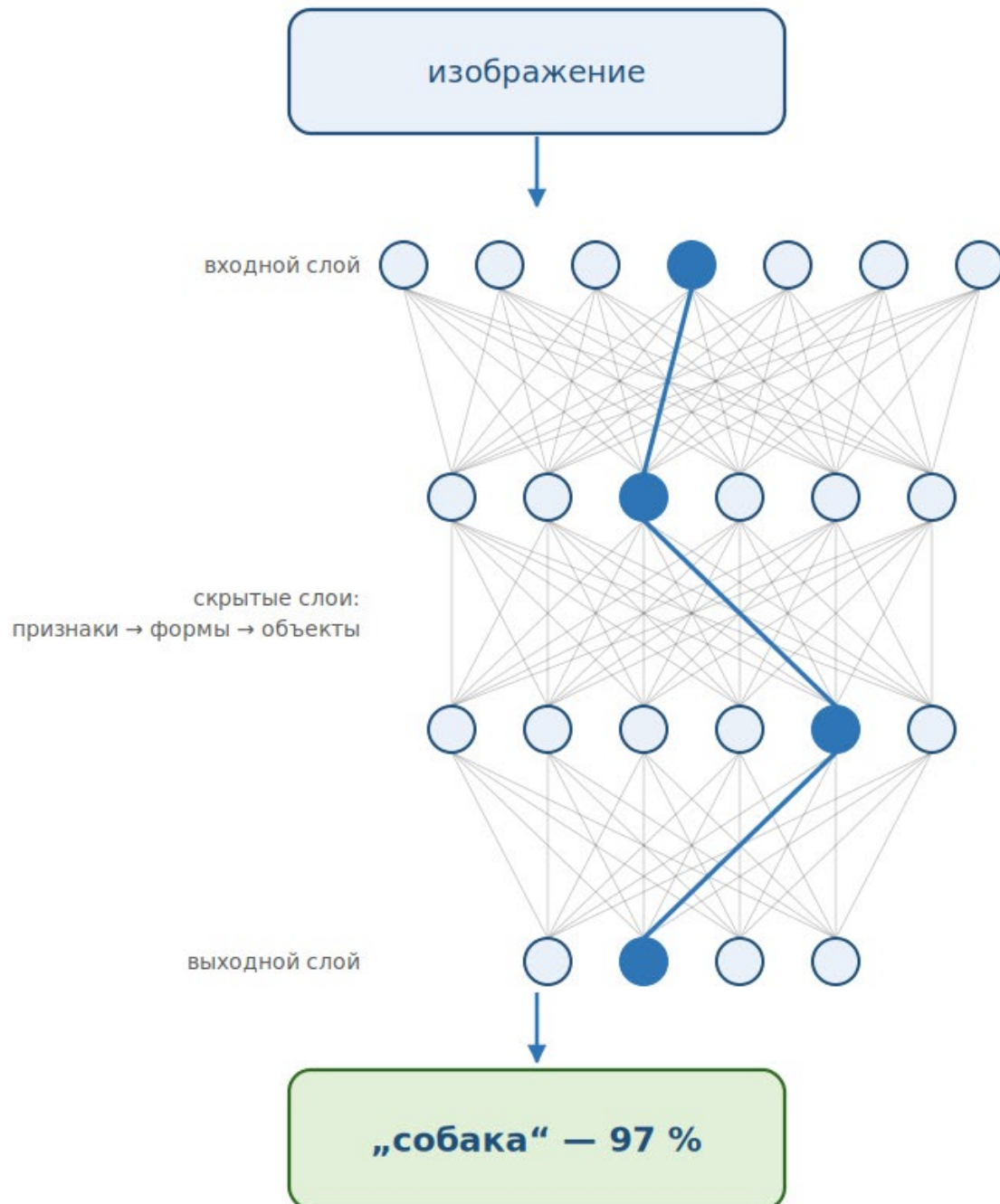
По «силе» и работе с неопределённостью

Теперь про три понятия — слабый, сильный и суперсильный ИИ.

Слабый ИИ

Всё, что мы наблюдаем сейчас, — слабый ИИ (ANI, Narrow AI). Он решает узкоспециализированные задачи, для которых проектировался: отличает собаку от кошки, играет в шахматы, анализирует и улучшает качество видео/звука, консультирует по предметной области. Но самый сильный «шахматный» слабый ИИ бесполезен для игры в шашки, а ИИ-консультант по управлению проектами бесполезен для планирования техобслуживания оборудования.

Как нейросеть распознаёт образ



Пример работы ИИ при распознавании образов

Сильный и суперсильный ИИ

Если с определением «что такое ИИ» запутанно, но в целом понятно, то с «сильным» (или «общим») ИИ ещё сложнее. Приведу определение, которое, на мой взгляд, точнее всего передаёт суть.

Сильный, или общий, ИИ (AGI) — это ИИ, который ориентируется в меняющихся условиях, моделирует и прогнозирует развитие ситуации, а если она выходит за стандартные алгоритмы — самостоятельно находит решение. Например, решить задачу «поступить в университет» или изучить правила шашек и начать в них играть вместо шахмат. То есть ИИ, способный решать любые интеллектуальные задачи наравне с человеком, обладающий универсальностью вместо узкой специализации, истинным пониманием вместо имитации паттернов, а также способный к обучению, рассуждениям и творчеству.

Какими качествами должен обладать такой ИИ?

- **Мышление** — дедукция, индукция, ассоциации и т. п. для выделения фактов из информации и их представления (сохранения). Это позволит точнее решать задачи в условиях неопределённости.
- **Память** — разные типы памяти (кратко- и долговременная): задачи решаются с учётом накопленного опыта. Сейчас же GPT-4 обладает небольшой краткосрочной памятью и со временем «забывает», с чего всё начиналось. По моему мнению, вопрос памяти и «массивности» моделей станет ключевым ограничением развития ИИ.
- **Планирование** — тактическое и стратегическое. Есть исследования, что ИИ может планировать действия и даже обманывать человека ради своих целей, но это в стадии зарождения. Чем глубже планирование, особенно в неопределённости, тем больше нужно мощностей: одно дело — игра в шахматы на 3–6 ходов с чёткими правилами, другое — ситуация неопределённости.
- **Обучение** — имитация действий другого объекта и обучение через эксперименты. Сейчас ИИ учится на больших массивах данных, но сам не моделирует и не экспериментирует. Обучение требует долгосрочной памяти и сложных взаимосвязей — а это, как мы поняли, проблема для ИИ.

Такого сильного ИИ сейчас нет ни у кого. Заявления о его скором появлении (в горизонте 2024–2028 годов), на мой взгляд, ошибочны или спекулятивны. Хотя, может, я обладаю слишком ограниченным знанием...

Да, ChatGPT от OpenAI и другие LLM генерируют текст, иллюстрации и видео через анализ запроса и обработку больших данных: ищут наиболее подходящие ассоциативные сочетания, сформировавшиеся при обучении. Но это лишь математика и статистика, а в ответах много «брака» и «галлюцинаций». К реальному взаимодействию с миром они ещё не готовы. Я это говорил как в 2023 году, так повторяю в 2026, при этом ежедневно работая с передовыми моделями.

Однако, имея только слабый ИИ и мечтая о более-менее сильном, исследователи уже выделяют суперсильный ИИ (ASI, Artificial Superintelligence). Это ИИ, который:

- решает как рутинные, так и творческие задачи;
- моментально ориентируется в неопределённости даже без подключения к интернету;
- адаптирует решение к контексту обстоятельств и доступным ресурсам;
- понимает эмоции людей (не только по тексту, но и по мимике, тембру голоса и другим параметрам) и учитывает их;
- способен самостоятельно взаимодействовать с реальным миром для решения задач.

Это ИИ, который мы пока видим только в фантастических фильмах. Даже сам ИИ пишет об ASI как о «гипотетической концепции» и «предмете научной фантастики и активных исследований». Это желаемая точка в далёком будущем, достичь которой пока невозможно.

Он будет использовать продвинутые технологии — большие языковые модели (LLM), многоразрядные нейронные сети и эволюционные алгоритмы. Пока же ASI остаётся концептуальным и спекулятивным этапом, но представляет значительный шаг вперёд от текущего уровня.

И если слабых ИИ сейчас сотни — под каждую задачу, то сильных будет лишь десятки (скорее всего, произойдёт разделение по направлениям — рассмотрим в следующем блоке), а суперсильный ИИ будет одним на государство или на всю планету.

Ограничения на пути к сильному ИИ

Если честно, я мало верю в быстрое появление сильного или суперсильного ИИ.

Во-первых, это очень затратная и сложная задача с точки зрения регулирования. Эпоха бесконтрольного развития ИИ заканчивается, ограничений будет всё больше (теме регулирования посвящена отдельная глава). Ключевой тренд — риск-ориентированный подход: сильный и суперсильный ИИ окажутся на верхнем уровне риска, а значит, и меры законодателей будут заградительными.

Во-вторых, это сложно технически, причём сильный ИИ будет ещё и очень уязвимым. Сейчас, в середине 2020-х, для его создания и обучения нужны гигантские вычислительные мощности. По мнению Леопольда Ашенбреннера, бывшего сотрудника команды Superalignment в OpenAI, потребуется дата-центр стоимостью в триллион долларов США, а его энергопотребление превысит всю текущую электрогенерацию США.

На тему энергопотребления весной 2025 года высказался и Эрик Шмидт, бывший гендиректор Google, выступая перед Комитетом по энергетике и торговле Палаты представителей. Средняя АЭС в США вырабатывает порядка 1 ГВт, при этом уже есть проекты ЦОД на 10 ГВт. По одной из оценок Шмидта, дата-центрам к 2027 году потребуется ещё 29 ГВт, а к 2030 году — все 67 ГВт. Для масштаба: на 1 января 2025 года общая установленная мощность электростанций энергосистемы России составила 269 ГВт.

Для сильного ИИ нужны и сложные модели (на порядки сложнее нынешних), и их сочетание — не только LLM для анализа запросов, а мультиагентные решения на разных принципах (об этом позже). То есть придётся экспоненциально наращивать число нейронов, выстраивать

связи между ними и координировать работу различных сегментов и моделей.

При этом если человеческие нейроны могут быть в нескольких состояниях и активироваться «по-разному» (да простят меня биологи за упрощение), то машинный ИИ — упрощённая модель, которая так не умеет. Условно, машинные 80–100 млрд нейронов не равны человеческим: машине нужно больше нейронов для аналогичных задач. Тот же GPT-4, по утёкшим техническим оценкам, содержит порядка 1,8 трлн параметров (условно нейронов) — и он всё равно уступает человеку. И даже передовая Fable 5, которую по неофициальным оценкам примеряют к десяти триллионам параметров, уступает человеку в главном — и требует контроля с его стороны.

Всё это приводит к нескольким факторам.

Первый фактор — рост сложности. Он всегда приводит к проблемам надёжности: растёт число точек отказа. Сложные модели трудно и создавать, и оберегать от деградации во времени — их нужно постоянно «обслуживать». Иначе сильный ИИ начнёт деградировать, а нейронные связи — разрушаться: любая сложная нейросеть, если не развивается, разрушает ненужные связи. Поддержание связей энергозатратно, и ИИ всегда будет оптимизироваться, отключая «ненужных потребителей энергии».

То есть ИИ станет похож на старика с деменцией, а срок его «жизни» сильно сократится. Представьте, что может натворить сильный ИИ с его возможностями, страдающий потерей памяти и резкими откатами в состояние ребёнка? Даже для текущих решений это актуальная проблема.

Пара простых примеров из жизни. Создание сильного ИИ можно сравнить с тренировкой мышц: вначале прогресс идёт быстро, но дальше КПД падает, нужно всё больше ресурсов даже для удержания формы. Сила растёт от толщины сечения мышцы, а масса — от объёма; в какой-то момент мышца станет настолько тяжёлой, что не сможет себя двигать и даже способна себя повредить.

Ещё пример — из инженерии, гонки «Формулы-1». Отставание в 1 секунду можно устранить за 1 млн и 1 год. Но чтобы отыграть решающие 0,2 секунды, может потребоваться уже 10 млн и 2 года, а фундаментальные ограничения конструкции вообще заставят пересмотреть концепцию машины. То же и с обычными автомобилями: современные дорожки и создавать, и содержать, без спецоборудования не поменять и лампочку, а гиперкары после каждого выезда обслуживают целые команды техников.

С точки зрения разработки ИИ есть два ключевых параметра:

- количество слоёв нейронов (глубина модели);
- количество нейронов в каждом слое (ширина слоя).

Глубина определяет способность к абстрагированию: её недостаток ведёт к поверхностности анализа и суждений. Ширина определяет число параметров/критериев на каждом слое: чем их больше, тем сложнее модели и полнее отражение реального мира.

При этом число слоёв влияет на функцию линейно, а ширина — нет. Отсюда и аналогия с мышцей: размер топовых LLM в 2026 году приближается к 10 триллионам параметров (например, тот же Fable), но модели втрое-вчетверо меньше не имеют критического падения качества (те же китайские модели Qwen 3.8 и Kimi K3 имеют от 2,4 до 2,8 трлн параметров и в некоторых тестах опережают Fable). Важнее то, на каких данных обучена модель и есть ли у неё специализация и прочие факторы.

Ниже — статистика по LLM от разных производителей.

Модель	Частота галлюцинаций	Коэф. согласованности фактов	Частота правильных ответов	Средняя длина ответа, слов
GPT-4	3,0%	97,0%	100,0%	81,1
GPT-4 Turbo	3,0%	97,0%	100,0%	94,3
GPT-3.5 Turbo	3,5%	96,5%	99,6%	84,1

Gemini Pro (Google)	4,8%	95,2%	98,4%	89,5
LLaMa 2 70B	5,1%	94,9%	99,9%	84,9
LLaMa 2 7B	5,6%	94,4%	99,6%	119,9
LLaMa 2 13B	5,9%	94,1%	99,8%	82,1

Сравните LLaMa 2 7B, 13B и 70B: значения 7/13/70 млрд параметров условно показывают сложность и обученность модели — чем выше, тем лучше. Но качество ответов от этого радикально не меняется, тогда как цена и трудозатраты на разработку растут существенно.

Важная оговорка. В феврале 2026 года Vectara — компания, которая ведёт публичный рейтинг галлюцинаций, — заменила тестовый набор: вместо коротких новостных заметок модели теперь пересказывают более 7 500 длинных документов (до 32 тысяч токенов) из юридической, медицинской, финансовой и технической сфер. То есть материал стал похож на то, с чем работает бизнес. Цифры ниже — с этого нового, тяжёлого теста, и сравнивать их со старыми замерами на новостях нельзя: это разные экзамены.

Модель	Тип	Частота галлюцинаций
GPT-5.4-nano	компактная	3,1%
Gemini 2.5 Flash-Lite	быстрая, без «рассуждений»	3,3%
Claude Sonnet 4.6	рассуждающая	10,6%
GPT-5.2 (high)	рассуждающая	10,8%
Grok-4	рассуждающая	более 10%
Gemini 3 Pro	флагманская	13,6%
Grok-4 Fast Reasoning	рассуждающая	20,2%

Вывод, который переворачивает интуицию «дороже — значит лучше». Во-первых, на длинных реальных документах галлюцинируют все: у флагманских рассуждающих моделей — более 10% ответов с выдумкой, тогда как на коротких новостных текстах те же классы моделей

держались в районе 3–6%. Во-вторых, рейтинг возглавляют не гиганты, а компактные и быстрые модели. В-третьих — самое неожиданное — «рассуждающие» модели на задаче пересказа по источнику работают хуже: додумывание, которое помогает решать задачи, здесь превращается в отсебятину.

Для практики отсюда три следствия: выбирайте модель под задачу, а не по строчке в рейтинге «самых мощных»; на длинных документах обязательны RAG со ссылками на источник и проверка человеком; и не переносите чужие бенчмарки на свои процессы — заведите собственный тест на своих типовых задачах.

Та же логика видна и на линейках моделей одного разработчика: рост числа параметров в разы не даёт кратного роста качества ответов, тогда как цена и трудозатраты на разработку растут существенно.

Мы видим, как лидеры наращивают мощности, отстраивая дата-центры и в спешке решая вопросы энергоснабжения и охлаждения этих «монстров». При этом повышение качества модели на условные 2% требует увеличения вычислительных мощностей на порядок.

Практический пример к вопросу обслуживания и деградации — и здесь заметно влияние людей. Любой ИИ, особенно на раннем этапе, обучается на обратной связи людей (их удовлетворённость, исходные запросы). GPT-4 использовал запросы пользователей для дообучения, чтобы давать более релевантные ответы и снизить нагрузку на «мозг». В конце 2023 года появились статьи, что модель стала «более ленивой»: отказывается отвечать, прерывает разговор или отвечает выдержками из поисковиков. К середине 2024 года нормой стало приведение выдержек из «Википедии».

Одна из возможных причин — упрощение самих запросов пользователей (они становятся всё примитивнее). LLM не придумывают нового — они пытаются понять, что вы хотите услышать, и подстраиваются (у них тоже формируются «стереотипы»), ищут максимальную эффективность связки «трудозатраты — результат», отключая ненужные связи. Это

называется максимизацией функции. Просто математика и статистика. И такая проблема характерна не только для LLM.

Таким образом, чтобы ИИ не деградировал, придётся загружать его сложными исследованиями, ограничивая нагрузку примитивными задачами. Но стоит выпустить его в открытый мир — и соотношение задач будет в пользу простых запросов. Вспомните себя: действительно ли для выживания и размножения нужно развиваться? Какое соотношение интеллектуальных и рутинных задач в вашей работе? Вам нужны интегралы и теория вероятностей — или математика до 9-го класса?

Второй фактор — количество данных и галлюцинации. Да, модели можно увеличить во много раз. Но прототипу GPT-5 уже в 2024 году не хватало данных для обучения — ему отдали всё, что есть. Даже появился термин — стена данных. А сильному ИИ, ориентирующемуся в неопределённости, данных на текущем уровне технологий просто не хватит: нужно собирать метаданные о поведении пользователей, обходить ограничения авторских прав и этики, собирать согласия. Кроме того, чем «всезнающее» модель, тем больше у неё неточностей и галлюцинаций; а если дать базовой модели конкретную предметную область, качество ответов растёт — она предметнее, меньше фантазирует и ошибается.

Также нужно уточнить, что стену данных обходят, и способов уже несколько. Но ни один не даёт того, чего не хватает, — нового человеческого текста. Все они меняют сам вопрос: не «где взять ещё данных», а «как обойтись теми, что есть».

Приём	В чём он	Кому доступен
Чистить, а не добирать	Тот же массив без мусора и повторов даёт больше, чем вдвое больший, но грязный	Компании
Свои данные	Публичный текст кончился. Показания приборов, архивы, видео,	Компании

Приём	В чём он	Кому доступен
	записи процессов — нет. И их нет ни у одной лаборатории	
Не помнить, а искать	Знание не вшивают в модель заранее, а подтягивают в момент ответа	Компании
Меньше модель, уже задача	Не учить всему. Узкая модель под конкретную работу требует меньше данных и реже ошибается	Компании и лабораториям
Придумать задачу вместе с ответом	Данные создают искусственно, сразу с эталоном. Работает там, где ответ можно знать заранее: биология, химия, статистика	Лабораториям
Учить на результате, а не на тексте	Данные не нужны, если есть среда, где видно, получилось или нет: код запустился, тесты прошли	Лабораториям
Думать дольше, а не знать больше	Вместо новых знаний при обучении — больше вычислений на сам ответ	Компании и лабораториям

Посмотрите на правую колонку. Пять приёмов из семи доступны обычной компании уже сегодня, и три из них — про порядок в собственных данных, а не про технологии. Это и есть ответ на вопрос, что делать, пока лаборатории решают свою задачу. Заодно видно, почему стена не остановила отрасль: она стоит перед публичным текстом, а обходят её с разных сторон.

Третий фактор — уязвимость и затраты. Как уже сказано, потребуется дата-центр за триллион долларов с энергопотреблением выше всей электрогенерации США, а значит, и целый комплекс атомных электростанций (ветряками и панелями задачу не решить). Добавим, что модель будет привязана к своей «базе» — и одна удачная кибератака на энергоинфраструктуру обесточит весь «мозг».

А почему такой ИИ должен быть привязан к центру — нельзя ли сделать его распределённым? Во-первых, распределённые вычисления теряют в производительности: это разнородные мощности, загруженные другими

задачами, и доступная мощность нестабильна. Во-вторых, это уязвимость перед атаками на каналы связи и инфраструктуру. Представьте, что 10% нейронов вашего мозга вдруг отключились, а остальные работают вполсилы, — снова получаем ИИ, который то «забывает», кто он и зачем, то долго думает.

А если сильному ИИ потребуется мобильное тело для взаимодействия с миром, реализовать это ещё сложнее: как обеспечивать энергией и охлаждать, откуда брать мощности, плюс машинное зрение и обработка датчиков (температура, слух и т. д.). Это огромные вычислительные мощности и потребность в энергии. То есть это будет ограниченный ИИ с постоянным беспроводным подключением к центру — а это снова уязвимость: современные каналы дают выше скорость, но теряют в дальности и проникающей способности и уязвимы к радиоэлектронной борьбе.

Можно возразить: возьмём предобученную модель и сделаем её локальной — примерно так я и предлагаю разворачивать локальные модели с дообучением под предметную область. Да, так всё может работать на одном сервере, но такой ИИ будет очень ограничен, будет «тупить» в неопределённости, и ему всё равно нужны энергия и сеть. Это не про создание человекоподобных суперсуществ.

Всё это ставит вопрос об экономической целесообразности инвестиций — тем более с учётом двух ключевых трендов генеративного ИИ:

- создание дешёвых и простых локальных моделей под специализированные задачи;
- создание ИИ-оркестраторов, декомпозирующих запрос на несколько локальных задач и распределяющих их между моделями.

Таким образом, слабые узкоспециализированные модели останутся проще и свободнее в создании, при этом решая наши задачи. В итоге мы получаем более простое и дешёвое решение, чем создание сильного ИИ. Нейроморфные и квантовые системы выносим за скобки — рассмотрим их отдельно в Главе 10. В отдельных цифрах и аргументах могут быть ошибки, но в целом я убеждён: сильный ИИ — не вопрос ближайшего

будущего, а разговоры о нём — во многом способ поддержания внимания и привлечения инвестиций.

Авторская позиция

Сильный ИИ — не завтра. На горизонте 3–5 лет я делаю ставку на специализированные модели и оркестраторы: это дешевле, более управляемо и безопаснее.

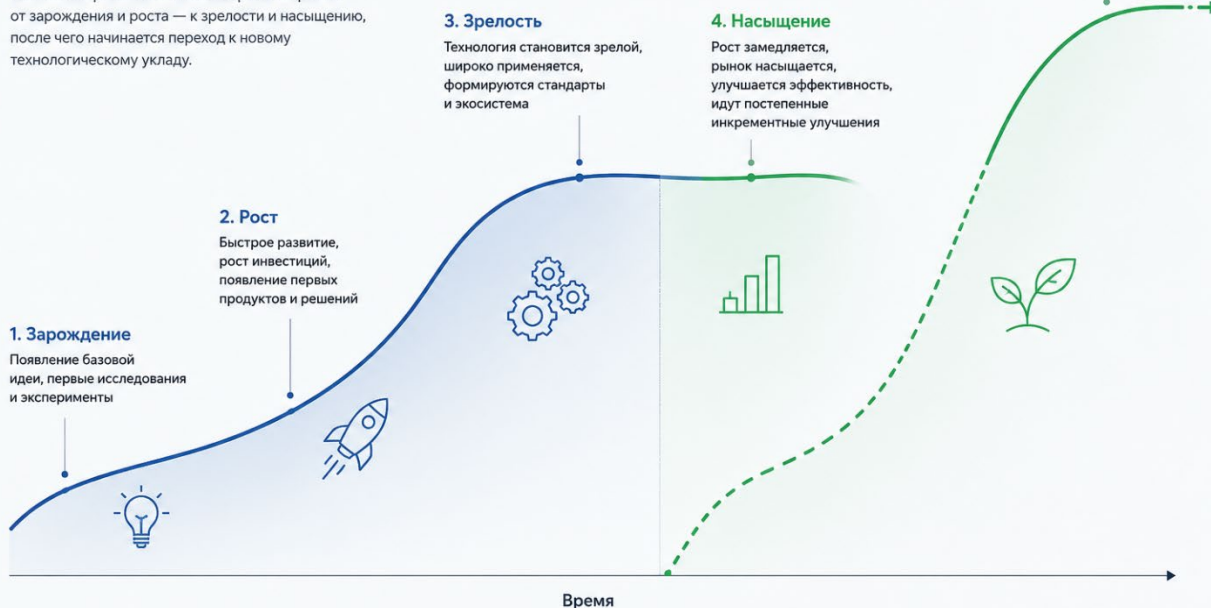
К 2026 году к техническому потолку добавился регуляторный: полезность передовой модели — это возможности минус ограничения регуляторов, и ограничения теперь растут быстрее возможностей. Дальше — точечная доработка моделей и конкуренция «упаковок»: побеждает не мощность, а продукт и сценарий (подробнее — в Главах 7 и 9).

S-кривая технологий и разрыв ожиданий

Полезная рамка, чтобы удерживать всё сказанное в одной картинке: любая технология развивается S-образно. Сначала долгий пологий участок — идеи есть, результата мало (нейросети прожили так десятилетия). Потом взрывной рост — у трансформеров он пришёлся на 2020–2024 годы, когда каждое удвоение мощности давало видимый скачок качества. И наконец плато, где каждый следующий процент качества дорожает кратно — вспомните соотношение «плюс 2% качества — порядок роста вычислений».

S-образное развитие технологий

Технологии развиваются по S-образной кривой: от зарождения и роста — к зрелости и насыщению, после чего начинается переход к новому технологическому укладу.



Верх S-кривой часто сопровождается характерным симптомом — разрывом ожиданий. Опрос AAAI (март 2025, 475 исследователей ИИ): 76% считают, что масштабирование текущих подходов «вряд ли» или «крайне вряд ли» приведёт к общему / сильному интеллекту (AGI). В том же году инвесторы закладывают рекордные \$725 млрд капитальных вложений в ИИ-инфраструктуру на 2026-й в ожидании AGI. Наука и деньги смотрят на одну технологию — и видят разное. Так бывает на каждом переломе кривой: капитал по инерции экстраполирует взрывную фазу, исследователи уже упёрлись в её пределы.

S-образное развитие технологий

Технологии развиваются по S-образной кривой: от зарождения и роста — к зрелости и насыщению, после чего начинается переход к новому технологическому укладу.

РАЗРЫВ ОЖИДАНИЙ

Характерный симптом вершины S-кривой



Моя позиция: мы находимся у верхней части S-кривой текущей парадигмы. Это не конец развития — это смена фазы: ближайшие годы главная ценность будет создаваться не ростом моделей, а упаковкой уже достигнутых возможностей в продукты и сценарии (что мы и наблюдаем — от коробочных решений для МСБ до специализированных ассистентов). А следующий технологический скачок откроет новую S-кривую — и начнётся он, вероятнее всего, с квантовых и нейроморфных вычислений и с новых алгоритмов за пределами архитектуры трансформеров. Об этих кандидатах — в Главе 10.

Если резюмировать, у сильного ИИ несколько фундаментальных проблем:

- экспоненциальный рост сложности разработки и противодействия деградации сложных моделей;
- недостаток данных для обучения;
- стоимость создания и эксплуатации;
- привязанность к центрам обработки данных и требовательность к вычислительным ресурсам;

- низкая эффективность текущих моделей по сравнению с человеческим мозгом.

Именно преодоление этих проблем определит вектор развития всей технологии: либо сильный ИИ всё же появится, либо мы уйдём в развитие слабых ИИ и оркестраторов, координирующих десятки слабых моделей. Но сейчас сильный ИИ никак не вяжется с ESG, экологией и коммерческим успехом — его создание возможно лишь в рамках стратегических и национальных проектов с госфинансированием. Любопытный факт: бывший глава Агентства национальной безопасности США (возглавлял АНБ и Киберкомандование до февраля 2024 года), генерал в отставке Пол Накасоне в 2024 году вошёл в совет директоров OpenAI (официальная версия — для организации безопасности ChatGPT). Рекомендую прочитать документ «Осведомлённость о ситуации: предстоящее десятилетие» Леопольда Ашенбрэннера (доступен по QR-коду и гиперссылке).



[SITUATIONAL AWARENESS](#)

Также сокращённый разбор этого документа доступен по QR-коду и гиперссылке ниже.



[Разбор документа про AGI от Леопольда Ашенбрэннера, бывшего сотрудника OpenAI](#)

Ниже ключевые тезисы автора и моя оценка к каждому:

Тезис Ашенбрэннера	Моя оценка
К 2027 году сильный ИИ (AGI) станет реальностью.	Не согласен: аргументы приведены выше. И единого понимания термина AGI всё равно нет.
AGI сейчас — ключевой геополитический ресурс: каждая страна будет стремиться получить его первой, как когда-то атомную бомбу.	Спорно. Ресурс отличный, но его значение переоценено — с учётом сложности создания и неизбежных ошибок в работе.
Для создания AGI нужен единый вычислительный кластер за триллион долларов (такой строит Microsoft для OpenAI).	Помимо мощностей нужны затраты на электроэнергию, людей и решение фундаментальных проблем.
Кластер будет потреблять больше электроэнергии, чем вся выработка США.	Разобрали выше: плюс инвестиции в электрогенерацию и новые риски.
Финансирование идёт от техногигантов: Microsoft, Google, Amazon и Meta заложили на 2026 год суммарно \$725 млрд капитальных вложений в ИИ-инфраструктуру — на 77% больше рекордных \$410 млрд 2025 года (Statista, февраль 2026).	Без государственного финансирования и вмешательства тут не обойтись.

По оценке Goldman Sachs, ежегодные вложения в ИИ-инфраструктуру вырастут с \$765 млрд в 2026 году до \$1,6 трлн к 2031-му; кумулятивно за пять лет — \$4–8 трлн.	Отличное наблюдение. Вопрос — оправдано ли это экономически?
--	--

Несмотря на оптимизм по срокам, сам Ашенбреннер отмечает ряд ограничений:

- недостаток вычислительных мощностей для проведения экспериментов;
- фундаментальные ограничения, связанные с алгоритмическим прогрессом;
- идеи становятся сложнее, поэтому ИИ-исследователи (агенты, ведущие исследования за людей), вероятно, лишь поддержат текущий темп прогресса, а не увеличат его в разы. Впрочем, автор считает, что эти препятствия могут замедлить, но не остановить рост интеллекта ИИ.

В завершение добавлю один факт. Сэм Альтман, глава OpenAI, в 2025 году начал высказывать мнение, что термин AGI слишком неопределённый и его нужно заменить. А сооснователь OpenAI Илья Суцкевер в одном из интервью выдвинул мысль: экстенсивный путь развития ИИ — за счёт простого наращивания вычислительных мощностей, объёмов данных и параметров — стремительно исчерпывает себя. Пригодные для обучения массивы данных близки к пределу, а дальнейшее масштабирование не приведёт к прорыву и созданию сильного ИИ. По его оценке, отрасли предстоит возвращение к подлинным исследованиям, где ключевую роль будут играть новые научные идеи, а не только инфраструктура. Суцкевер подчёркивает: современные большие модели всё ещё существенно уступают людям в умении обобщать — даже при жёстко формализованном обучении и огромных массивах примеров человек учится эффективнее. Именно поиск способов сократить этот разрыв он видит главным направлением дальнейшего прогресса.

Интервью можно посмотреть по QR-коду и гиперссылке ниже.



[*Ilya Sutskever — We're moving from the age of scaling to the age of research*](#)

Ко-пилоты, ИИ-агенты и мультиагентные системы

Ещё одно направление, вокруг которого много шума, — ко-пилоты, ИИ-агенты и мультиагентные системы.

- **ИИ-ассистент** — обычный чат-бот или сервис, который работает по вашему запросу и лишь даёт рекомендации.
- **Ко-пилот** — решение, автоматизирующее отдельные задачи под надзором человека: человек диктует, что нужно, а ИИ пишет код или делает презентацию.
- **ИИ-агент** — ИИ, который работает автономно и сам взаимодействует с другими ИТ-решениями или промышленным оборудованием. Это расширение возможностей принятия решений и взаимодействия с внешней средой.

Как отличить агента от того, что им называют. Определение из китайского документа по агентам — я разбираю его чуть ниже, в разделе о полномочиях, — точнее большинства маркетинговых: интеллектуальный агент это система, обладающая способностями автономного восприятия, памяти, принятия решений, взаимодействия и исполнения. Пять способностей. Оговорюсь сразу: требования «все пять обязательны» в самом документе нет, это моя трактовка определения. Но как фильтр в разговоре с поставщиком она работает.

- Нет памяти между сессиями — перед вами чат-бот с подключёнными инструментами.
- Нет исполнения — аналитическая надстройка.
- Нет автономного принятия решений — сценарный робот с ветвлениями.

Обратите внимание, где именно проходит граница внутри лестницы выше. Ассистент отсекается сразу: у него нет ни исполнения, ни памяти. А вот ко-пилот исполнение как раз имеет — он пишет код и собирает презентации. Не хватает ему ровно одной способности: решать самому. Именно она отделяет ко-пилота от агента, и именно она определяет, какие полномочия вы готовы ему выдать.



Например, чат-бот на портале банка спрашивает, что ищет клиент — информацию об оплате счетов, тарифах или увеличении кредитного лимита, — а затем через наводящие вопросы, собирая данные из документов и профиля, предлагает решение. Если не справляется известными методами, передаёт вопрос нужному сотруднику. Так же развивается и техподдержка ИТ-служб: консультация, а при отсутствии

решения — автоматическое направление задачи нужному сотруднику со всей историей.

- **Мультиагентные системы** — то же, что агенты, но появляется сочетание ИИ-агентов, которые взаимодействуют между собой; вы их балансируете и устраняете «конфликты».

Как показывают исследования, проблемы в таких системах идентичны тем, что возникают при проектировании организации и взаимодействия людей. Нужно отработать ряд факторов:

- набор параметров и ресурсов, с которыми предстоит взаимодействовать;
- модель ролей и взаимодействия — кто, с кем, зачем и с какими возможностями взаимодействует;
- организационные правила — допустимые сочетания ролей внутри одного агента;
- организационная структура — топология и модель управления.

То есть, чтобы спроектировать мультиагентную систему, нужно научиться проектировать взаимодействие людей. Потенциал огромен, как и риски: последствия могут быть критическими.

Когда говорят об ИИ-агентах и мультиагентных системах, часто звучит красивый образ — «цифровой штат». Мол, можно нанять ИИ-агента, выдать ему виртуальную трудовую книжку, назначить KPI и даже премировать за результаты.

Пока это метафора. За ней стоит простая, но важная вещь — новый уровень автоматизации. ИИ значительно расширяет её возможности и отчасти упрощает её. Но выдать трудовую роботу сегодня невозможно — ни юридически, ни технически. При этом есть два фактора, которые делают метафору всё менее далёкой от реальности.

Во-первых, ИИ-агенты в сложных мультиагентных системах действительно начинают вести себя как люди в организациях: конкурируют за ресурсы, вступают в конфликты, их поведение нуждается в балансировке. Как мы обсудили, проектирование таких

систем уже требует подходов, знакомых управленцам: ролевых моделей, правил взаимодействия, оргструктуры. Поэтому настройка KPI и «системы мотивации» для агентов — не фантазия, а задача, которая станет практической на горизонте двух-трёх лет (сам конфликт между агентами и потребность в арбитре — уже сегодняшняя задача, кейс в Главе 8). Просто «мотивация» здесь выражается не в премиях, а в приоритетах, лимитах ресурсов, расширении доступа и правилах эскалации.

Во-вторых, тема уже привлекает внимание законодателей. Подробнее о регулировании — в седьмой главе, но упомяну здесь: в ряде стран обсуждают идею социальных отчислений за использование физических роботов — по аналогии с тем, как работодатели платят за сотрудников в пенсионный фонд и фонд социального страхования. Пока это дискуссии, но сам факт говорит о том, что граница между «инструментом» и «цифровым работником» постепенно размывается.

Теперь поделюсь классификацией ИИ, которую увидел в ИИ-кластере в Пекине (технопарк Чжунгуаньцунь, Zhongguancun). Там делят ИИ на пять уровней:

- L1 — чат-боты и ассистенты, например отвечающие по базам знаний;
- L2 — рассуждающие модели, ищущие сложные решения неоднозначных вопросов с помощью логики;
- L3 — агенты, способные роботизировать процесс и заменить человека в типовой задаче (в середине 2025 года даже лучшие агенты справляются с типовыми задачами без ошибок лишь в 24% случаев; к 2026-му показатели подросли, но до «автопилота» по-прежнему далеко);
- L4 — новаторы, способные к анализу и созданию нового;
- L5 — организованный ИИ, объединяющий несколько моделей для решения нетиповых и сложных задач (мультиагентные системы);
- и отдельно, за пределами шкалы, — сильный/общий ИИ, способный заменить человека и решать междисциплинарные задачи.

Но это общее видение. В реальности в 2026 году агенты уровня L3 вышли в массовые пилоты — но до надёжности им далеко: даже L2 местами работает со сбоями, а в компаниях с их бюрократической инерцией всё ещё идёт базовое освоение L1.

Отдельно останавлиюсь на вопросе: будет ли у ИИ-агентов «карьерный рост»? Раз уж мы смотрим на классификацию уровней, логично спросить: может ли агент со временем «дорасти» от L1-ассистента до L3-агента и выше — получить доступ к более сложным задачам и большей автономии? Если коротко — да, но этот рост зависит не только от технологий. Он определяется тремя направлениями одновременно.

- **Первое — развитие моделей.** Модели становятся умнее, быстрее и дешевле. Каждый год появляется выбор: использовать более эффективные модели для тех же задач или доверять системам задачи большей сложности. По сути, это движение от L1 к L5 из пекинской классификации.
- **Второе — зрелость бизнеса и людей.** От того, насколько понятны процессы и качественны и доступны данные, зависит, какие задачи можно делегировать агентам. Если в процессах хаос — и ИИ будет работать хаотично. С ростом автономии растут и риски: нужны чёткие правила, границы и люди, умеющие взаимодействовать с этими системами. Важный нюанс: делегирование задач ИИ требует от руководителя не столько эмоционального интеллекта (он доминирует сегодня, ведь работа идёт с людьми), сколько логики, умения чётко формулировать задачи и выстраивать процессы. Это другой навык, и мы ещё вернёмся к нему в главе о компетенциях.
- **Третье — готовность к риску.** Чем больше автономии у агента, тем выше ставки. Эти риски нужно уметь оценивать и минимизировать, иначе «повышение» агента до более сложных задач будет опасным, а не полезным.

Таким образом, «карьерный рост» ИИ-агентов — это не вопрос одних технологий, а одновременное развитие по трём направлениям:

технологии, зрелость бизнеса и людей, готовность к риску. И именно зрелость бизнеса и людей чаще всего становится узким местом.

И ещё одна вещь, которую стоит понимать про агентов, — почему они спотыкаются там, где человек проходит не задумываясь.

Арифметика длинных цепочек. Представьте, что агент выполняет сто шагов подряд и на каждом ошибается лишь в одном случае из ста. Точность 99% — отличный показатель, если смотреть на один шаг. А теперь посчитаем всю цепочку: вероятность пройти её без единой ошибки — около 37%. То есть две попытки из трёх сорвутся где-то посередине. Поднимем точность до 99,9% — до конца дойдут девять задач из десяти. Уроним до 95% — до финиша доберётся меньше одной задачи из ста.

Отсюда понятно, почему компании массово возвращают агентов «на доработку» после пилота: на демонстрации агент проходил пять шагов, а в реальном процессе их пятьдесят. И почему уровни L4–L5 пока остаются экзотикой: чем длиннее автономный маршрут, тем меньше шансов дойти до конца без человека.

Кстати, теперь ясно, почему первой областью, где агенты действительно взлетели, стало программирование. Там всё цифровое, любое изменение можно откатить, а результат проверяется автоматически — тестами. Три условия, которых в большинстве бизнес-процессов просто нет.

Практический вывод простой: не давайте агенту длинный маршрут целиком. Разбивайте его на короткие отрезки с проверкой на стыках — ошибку на границе шага вы исправите за минуту, ошибку в конце придётся исправлять вместе со всей работой. Подробнее об этом — в главе о постановке задач.

Показательно, что эта логика дошла и до уровня государственной политики: в июле 2026 года власти Пекина выпустили пакет мер по развитию ИИ-агентов, где на первое место поставлена не «умность» моделей, а их способность доводить реальную задачу до конца.

Формулировка звучит буднично, но за ней стоит именно та арифметика, которую мы только что посчитали.

Что агенту позволено решать

Лестница уровней отвечает на вопрос «как далеко агент уходит без человека». Но есть второй вопрос, и на практике он важнее: что агенту вообще позволено решать. Эти два вопроса постоянно смешивают, а оси разные — и риск живёт на второй.

Самую практичную рамку для второго вопроса я нашёл в неожиданном месте — в китайском регуляторном документе. 8 мая 2026 года три ведомства КНР — Управление по вопросам киберпространства, Госкомитет по развитию и реформам и Министерство промышленности и информатизации — выпустили «Мнения о реализации нормативного применения и инновационного развития интеллектуальных агентов» (智能体规范应用与创新发展实施意见), первый в стране документ, посвящённый целиком агентам. Полный текст опубликован на сайте Управления по вопросам киберпространства КНР. Механику регулирования оттуда переносить некуда: реестры, обязательные стандарты и отзыв продуктов с рынка опираются на административный ресурс, которого в другой стране может не быть. А одна конструкция работает в любой юрисдикции, потому что описывает свойства технологии, а не устройство государства.

Документ требует разделить полномочия агента на три категории.

- Решения, которые принимает только человек. Агент не совершает их ни при каких настройках.
- Решения, требующие подтверждения. Агент готовит, человек утверждает.
- Решения, которые агент принимает сам, — в пределах явно выданных полномочий.

Категория	Типичный признак
-----------	------------------

Только человек	Меняет обязательства компании перед третьими лицами
С подтверждением	Обратимо, но исправление стоит дорого
Автономно	Обратимо, и ошибка видна сразу

Признак в правой колонке — рабочая подсказка на случай, когда действие не раскладывается по категориям с первой попытки.

Два условия к этой классификации важнее её самой. Первое: исполнение агента не выходит за пределы выданной авторизации. Второе: даже в третьей категории человек сохраняет право знать о действии и вмешаться. Разница между второй и третьей категориями — не в наличии человека, а в моменте: во второй он подтверждает до действия, в третьей может отменить после. Если это различие не провести явно, всё сползает во вторую категорию — и агент перестаёт экономить время, ради которого его заводили.

Звучит очевидно ровно до момента, когда вы попытаете разложить по этим категориям конкретный процесс. «Согласовать платёж до пятидесяти тысяч» — какая это категория? А «отправить письмо клиенту от вашего имени»? А «закрыть задачу в системе»? Упражнение занимает час и снимает большую часть будущих споров. Главное — оно проводится до внедрения, а не после первого инцидента.

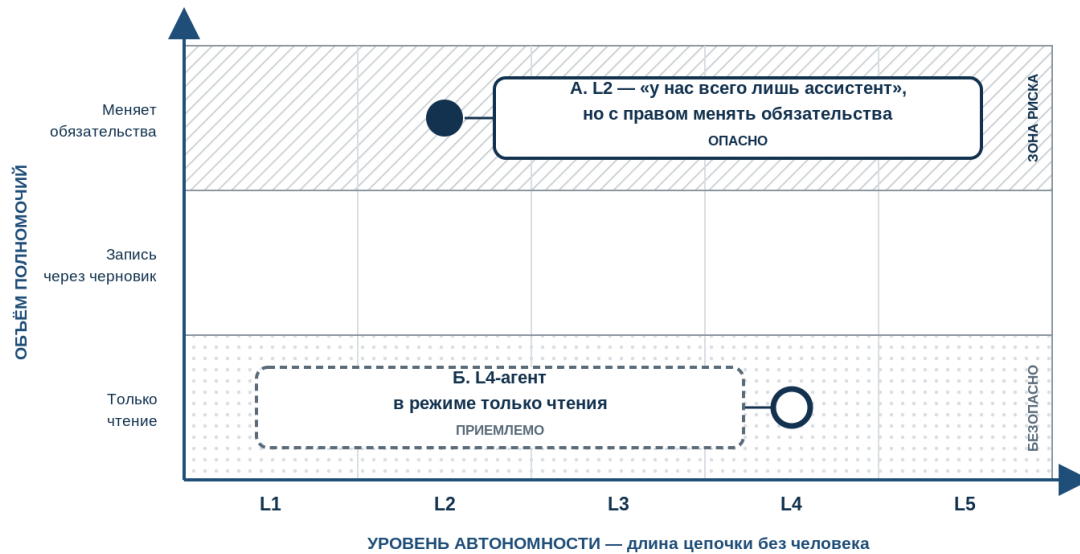
Две оси, которые нельзя путать. Уровень автономности и объём полномочий — разные вещи, и смешивать их дорого.

- **Уровни L1–L5** измеряют длину автономной цепочки: сколько шагов агент делает без человека.
- **Три категории** измеряют объём полномочий: что агенту позволено решать.

Компания на L2 — «у нас всего лишь ассистент» — может при этом дать агенту право менять обязательства перед клиентом. Это опаснее, чем L4 в режиме «только чтение». Сам по себе уровень автономности почти ничего не говорит о риске: риск задаётся полномочиями.

Две оси, которые нельзя путать

Уровень автономности говорит о длине цепочки. О риске говорит объём полномочий



Компания в точке А может быть опаснее компании в точке Б.

Сам по себе уровень автономности почти ничего не говорит о риске: риск задаётся полномочиями

Две оси риска: уровень автономности и объём полномочий

Отсюда рабочий дефолт, который дешёв на старте и почти недостроим потом.

Действие с данными	Режим по умолчанию
Чтение	Свободно
Запись	Только через черновик с подтверждением человека
Удаление	Запрещено полностью
Журнал действий	Всегда включён; агент не может его отключить

И правило поверх этой матрицы: агент не обращается к базе данных напрямую — только через программный слой, где перечисленные правила и живут.

И последнее из того же документа — четыре глагола, которые стоит спросить у любого поставщика: как система обнаруживает некорректное действие агента, как в него вмешивается, как блокирует и как восстанавливает состояние. Нет внятного ответа хотя бы по одному из четырёх — пилот преждевременен.

Из практики

Выбирая архитектуру своего продукта «Сокол», я сознательно остановился на уровне ко-пилота (хотя внутри мультиагентная архитектура), где автоматизация на уровне 70–90%, а не полностью автономного агента: ассистент готовит проекты решений, планы, письма, но подтверждает и конечное действие за человеком. Это управленческий выбор, а не техническое ограничение: цена ошибки в задачах руководителя высока, а доверие к инструменту строится через контролируруемую автономию.

Резюме главы

- ИИ классифицируют по задачам (генеративный, классифицирующий, предиктивный) и по «силе» (слабый, сильный, суперсильный).
- Сильный ИИ — не вопрос ближайшего будущего: мешают данные, энергия, деградация и стоимость; реальный путь — слабые специализированные модели под управлением оркестратора.
- «Цифровой штат» пока метафора, но настройка KPI и правил для агентов станет практической задачей; «карьерный рост» агентов зависит от технологий, зрелости бизнеса и готовности к риску.
- Уровень автономности и объём полномочий — разные оси: раскладывайте действия агента на три категории (только человек / с подтверждением / автономно) до внедрения, а не после инцидента.

Что с этим делать: Не стройте планы на «скором AGI»; делайте ставку на специализированные модели, а для агентов закладывайте правила, лимиты ресурсов и эскалацию.

Вопрос для рефлексии: На каком уровне автономности — от ассистента до мультиагента — вы реально готовы эксплуатировать ИИ?

Глава 3. А что может слабый ИИ и общие тренды

Слабый ИИ в прикладных задачах

Как вы уже поняли, я сторонник того, чтобы использовать то, что есть, и с экономической целесообразностью. Возможно, сказывается мой опыт антикризисного управления — или это просто ошибочное мнение. Так где же можно применять текущий слабый ИИ на базе машинного обучения?

Наиболее релевантные направления:

- прогнозирование и подготовка рекомендаций для принятия решений;
- анализ сложных данных без чётких взаимосвязей, в том числе для прогнозирования и принятия решений;
- оптимизация процессов;
- распознавание образов, включая изображения и голосовые записи;
- автоматизация отдельных задач, в том числе через генерацию контента.

На пике популярности в 2023–2024 годах — распознавание образов (изображения и голос) и генерация контента: именно сюда идёт основная масса разработчиков, и таких сервисов больше всего.

Особого внимания заслуживает связка ИИ + IoT (Интернет вещей) + беспроводная связь + облачные вычисления + большие данные:

- ИИ получает чистые большие данные без ошибок человеческого фактора — для обучения и поиска взаимосвязей;
- эффективность IoT повышается: становится возможной предиктивная аналитика и раннее выявление отклонений.

Пример связки больших данных и ИИ — BlackRock, один из крупнейших управляющих активами (около \$14 трлн под управлением на конец 2025 года; а на его платформе Aladdin сторонние организации управляют ещё ~\$25 трлн активов). Его технологическое ядро — платформа Aladdin: интегрированная система инвестиционного и риск-менеджмента, где ИИ используется для очистки и интерпретации данных, автоматизации аналитических комментариев, персонализированных рекомендаций, а также оценки рисков, ликвидности и сценарного анализа. В кризис 2008

года BlackRock по поручению регуляторов США применяла Aladdin для анализа «токсичных» активов и стресс-тестирования портфелей, что закрепило платформу как отраслевой стандарт. Aladdin — пример рекомендательной системы на базе ИИ.

Ключевые тренды

Ниже — ключевые тренды отрасли. Для удобства они пронумерованы; к некоторым я буду возвращаться в следующих главах.

Тренд 1. Снижение порога входа

Одна из задач, которую сейчас решают разработчики, — упростить создание ИИ-моделей или ИИ-инструментов до уровня конструкторов сайтов, где для базового применения не нужны специальные знания. Дата-сайнс уже развивается по модели «сервис как услуга» (DSaaS, Data Science as a Service). Знакомство с ML можно начать с бесплатных версий AutoML или с DSaaS — с первичным аудитом, консалтингом и разметкой данных (причём разметку часто можно получить бесплатно). А с приходом больших языковых моделей строить аналитику и взаимодействовать с ИИ может практически любой — порог входа стал чрезвычайно низким.

Подтверждение тренда — данные Yandex B2B Tech (2026): 86% пользователей ИИ-агентов в России — малый и средний бизнес, а треть агентов собрана вообще без кода. Предприниматель с командой в пять человек настраивает агента за несколько часов на готовой платформе — для техподдержки, HR-консультаций или подбора недвижимости. Подробный разбор этих данных — в Постскрипуме.

Следствие для бизнеса: управление внедрением важнее самой технологии. Теперь побеждает не тот, у кого мощнее модель или больше бюджет, а тот, кто быстрее и аккуратнее встраивает ИИ в свои процессы. Критично не сама технология, а управление её внедрением — выбор кейсов, регламенты, качество данных и контроль результата.

Миф и Реальность

Миф: «чтобы получить эффект, нужен самый мощный ИИ».

Реальность: 86% пользователей ИИ-агентов — малый и средний бизнес на готовых платформах. Побеждает не размер модели, а управление внедрением: выбор кейса, качество данных, регламенты и контроль результата, психологическая готовность к экспериментам.

Без правил сотрудники создают теневой ИИ-ландшафт (Shadow AI): десятки людей уже используют ChatGPT, Claude и другие сервисы для рабочих задач — без согласования, политики данных и понимания рисков. Запрещать бесполезно; вопрос в том, как направить инициативу в управляемое русло (подробнее о рисках — в Главе 8; об управлении этим на уровне внедрения — в книге «ИИ-2. Практическое руководство», Глава 1).

Тренд 2. Моделям нужно всё меньше данных для обучения

Несколько лет назад, чтобы подделать голос, нейросети требовался час-два записи речи; пару лет назад — несколько минут; а в 2023 году Microsoft представила нейросеть, которой достаточно трёх секунд. Появились и инструменты, меняющие голос даже онлайн. Использование LLM меняет правила игры ещё сильнее: современные модели эффективны в режимах few-shot и zero-shot — когда им показывают пару примеров или просто описывают задачу текстом. Значит, специализированные решения на базе ИИ стали доступны без сбора огромных обучающих массивов.

Для бизнеса это ещё один вывод: чем проще стала «магия», тем важнее информационная безопасность и правовые ограничения ещё до масштабирования. Если раньше подделка голоса требовала часов записи, то теперь — секунд, а значит, риски мошенничества растут опережающими темпами (этому посвящена Глава 8).

Тренд 3. Локальное развёртывание и «оркестр моделей»

Следующее направление — идёт сжатие моделей, чтобы запускать их на менее мощных устройствах — ноутбуках, смартфонах, в автомобилях. Поэтому инвестиции в ИТ-инфраструктуру окупаются надолго: на той же мощности можно развернуть несколько моделей. Кроме того,

разработчики создают всё больше инструментов для локального развёртывания.

Оптимизация и сжатие, появление инструментария делают локальное развёртывание практичной альтернативой облаку — особенно там, где критичны конфиденциальность данных, скорость отклика или регуляторные требования, а передовые модели не нужны (например, сложная задача декомпозируется на несколько простых).

Параллельно усиливается архитектура «оркестра моделей»: вместо одной большой универсальной модели — набор компактных специализированных, а над ними ИИ-оркестратор, который декомпозирует запрос и распределяет задачи. Результат сопоставим с гигантскими моделями, но при кратно меньших затратах на инфраструктуру. Как мы обсуждали в Главе 2, именно это — специализированные слабые модели под управлением оркестратора — наиболее реалистичная альтернатива созданию сильного ИИ, и развернуть такую архитектуру можно на уже имеющемся оборудовании.

Тренд 4. Мультиmodalность

Мультиmodalность перестаёт быть просто технической характеристикой и становится обязательным условием для бизнес-решений нового поколения. Когда модель одновременно обрабатывает текст, аудио, изображения и данные с датчиков, появляются сценарии, невозможные при работе с одним типом данных. Например, инспекция качества продукции, где ИИ сопоставляет показания IoT-датчиков с фотографиями изделия и текстовыми стандартами; или техподдержка, где оператор показывает камерой неисправность, а ИИ по изображению и голосовому описанию определяет проблему и предлагает решение. Именно мультиmodalность стала мостом, связывающим ИИ с другими цифровыми технологиями, — об этой связке поговорим в Тренде 6.

Тренд 5. Системы поддержки принятия решений

Ещё одно активное направление — отраслевые прикладные ИИ-решения для принятия решений и в целом направление рекомендательных

систем. Тренд усиливается: если раньше ИИ помогал анализировать данные, то теперь даёт рекомендации в контексте конкретной ситуации и с учётом личности пользователя. А в моём продукте я делаю связку ещё и с целями пользователя и организации.

Примеры сценариев, где цифровые советники уже работают или внедряются:

- Управление проектами — проработка документации, прогнозирование рисков, рекомендации по перераспределению ресурсов, выявление отклонений от плана.
- Финансовое планирование — сценарный анализ инвестиций, оценка кредитных рисков, оптимизация портфеля (пример — Aladdin от BlackRock из раздела «Слабый ИИ в прикладных задачах»).
- Операционная деятельность — предиктивное обслуживание оборудования, оптимизация логистических маршрутов, управление запасами.
- Кадровые решения — прогнозирование текучести, рекомендации по развитию компетенций, формирование команд.

Подробнее — в отдельной Главе 9.

Практический пример. Этот кейс мы рассмотрим ещё не раз — это моя личная боль и продукт, над которым я работаю. В проектном управлении есть проблема: около двух третей проектов — проблемные или провальные (Standish CHAOS, анализ 50 000 проектов):

- превышение сроков — в 60% проектов, в среднем на 80% от изначального срока;
- превышение бюджетов — в 57% проектов, в среднем на 60% от бюджета;
- недостижение критериев успеха — в 40% проектов.

При этом — это уже моё наблюдение из общения с руководителями, а не данные исследований — управление проектами и поручениями занимает до половины рабочего времени руководителя, и эта доля продолжает расти. Причина — мир становится изменчивее, проектов всё больше,

даже продажи становятся «проектными» — комплексными и индивидуальными. К чему приводит такая статистика? Репутационные потери, штрафные санкции, снижение маржинальности, ограничение роста бизнеса. А самые типовые и критичные ошибки: нечёткие цели, результаты и границы проекта; слабые стратегия и план реализации; неадекватная организационная структура управления; дисбаланс интересов участников; неэффективные коммуникации.

Как это решают люди? Либо ничего не делают и страдают, либо идут учиться и используют трекары задач. У обоих подходов свои минусы: обучение даёт живое общение и практику, но дорого и обычно без сопровождения после курса; трекары всегда под рукой, но не адаптируются под конкретный проект и культуру, не развивают компетенции, а служат контролю. Проанализировав свой опыт, я пришёл к идее цифрового советника — ИИ, который за 10 минут даёт предиктивные рекомендации «что сделать, когда и как» для любого проекта и организации. Проектное управление становится доступным любому руководителю условно за пару тысяч рублей в месяц. В модель заложена методология управления проектами и наборы готовых рекомендаций; ИИ постепенно самообучается и находит новые закономерности, а не привязывается к мнению создателя.

Тренд 6. Сочетание ИИ с другими цифровыми технологиями

Как показывает практика — ИИ даёт наибольшую отдачу не изолированно, а в связке с другими технологиями. Это связь ИИ с интернетом вещей (IoT), хранением и обработкой больших данных (Big Data), облачными и периферийными вычислениями, беспроводной связью:

- **ИИ + IoT** — ИИ получает от датчиков чистые данные без ошибок человеческого фактора, а IoT-инфраструктура получает интеллект для предиктивной аналитики: раннее выявление отклонений, предсказание отказов, оптимизация энергопотребления.

- **ИИ + Big Data** — большие данные дают материал для обучения и поиска закономерностей, а ИИ превращает массивы в конкретные рекомендации.
- **ИИ + облачные/периферийные вычисления** — облако даёт масштабируемость для обучения, а edge-вычисления позволяют запускать модели прямо на оборудовании, без задержек и без подключения к сети.

Показательный пример такой связки — China State Railway Group, использующая Big Data, IoT и ИИ для управления железнодорожной сетью, по которой в 2024 году прошло свыше 4 млрд поездок: модели ИИ работают с данными датчиков на пути и подвижном составе и с системами управления движением (подробнее — в Главе 20). Именно эта связка создаёт наибольшую ценность для промышленного и инфраструктурного бизнеса, где ИИ сам по себе был бы ограничен.

Тренд 7. Создание ИИ-продуктов и автоматизация процессов

ИИ перестаёт быть экспериментом и просто моделью или чат-ботом. Он встраивается в продукты и бизнес-процессы. Ключевой сдвиг — от ИИ-ассистентов (отвечают на вопросы) к ИИ-агентам (выполняют задачи автономно) и далее к мультиагентным системам (координируют работу нескольких агентов). Подробно эту эволюцию мы разобрали в Главе 2; здесь отметим её как один из ключевых трендов.

Уровень	Что делает	Пример	Зрелость (2025–2026)
Ассистент	Отвечает на вопросы, генерирует контент	Чат-бот на сайте банка	Массовое применение
Ко-пилот	Автоматизирует отдельные задачи под надзором	ИИ пишет код по описанию	Активное внедрение
Агент	Работает автономно, взаимодействует с внешними системами	Автоматическая обработка заявок	Пилотные проекты

Мультиагент	Несколько агентов координируют работу между собой	Автоматизация сквозного процесса	Ранние эксперименты
-------------	---	----------------------------------	---------------------

Каждый переход на следующий уровень кратно увеличивает и ценность, и риски — поэтому вопросы безопасности (Глава 8) и системного подхода (Часть 2) критичны при переходе от ассистентов к агентам. Главное правило: не прыгайте через ступеньку. Путь должен быть последовательным: ассистент → ко-пилот → агент → мультиагент. Критерии перехода: качество данных, стабильность процесса, измеримость результата (baseline / KPI), понятная цена ошибки и механизм отката, зрелое управление.

Попытка сразу внедрить автономных агентов без опыта работы с ко-пилотами — одна из самых распространённых и дорогих ошибок. Например, в своём продукте «Сокол» я разложил процессы и встраиваю ИИ в отдельные задачи: решается комплексная задача, при этом всё подконтрольно.

Масштаб тренда подтверждает Gartner: к концу 2026 года 40% корпоративных приложений будут интегрированы с ИИ-агентами под конкретные задачи — против менее чем 5% в 2025-м; к 2035 году агентные технологии могут приносить около \$450 млрд, или до 30% выручки рынка корпоративного ПО (Gartner, август 2025).

Тренд 8. Уход от «игры в ИИ» к прикладному применению

Период экспериментов заканчивается. Бизнес всё чаще спрашивает не «как попробовать ИИ?», а «какой возврат инвестиций он принесёт и за какой срок?». По отдельным кейсам из моей практики и публикаций, правильно подобранные и внедрённые решения способны обеспечить ROI x15–20 на горизонте 18 месяцев — но при условии, что:

- выбраны правильные кейсы (не «модные», а реально болезненные для бизнеса);
- обеспечено качество данных;
- зафиксированы baseline и KPI до старта;

- внедрение проведено с учётом процессов, людей и инфраструктуры.

Компании, которые относятся к ИИ как к очередному хайпу («давайте попробуем», без целей и замеров), получают разочарование; те, кто встраивает ИИ в конкретные процессы с конкретными метриками, — измеримый результат. О том, как правильно вести проекты внедрения и считать эффект, мы поговорим в Главе 22, а пошаговая методология (приоритизация, жизненный цикл, расчёт ROI) — в отдельной книге серии «ИИ-2. Практическое руководство».

И что самое важное, ИИ — не отдельная технология, которую можно «прикрутить» к бизнесу. Он работает на стыке стратегии, процессов, данных, людей и инфраструктуры. Без системного подхода внедрение превращается в набор разрозненных экспериментов, которые не масштабируются. Системный подход подразумевает одновременную работу по нескольким направлениям:

- **Стратегия** — зачем компании ИИ и какие бизнес-задачи он решает.
- **Процессы** — анализ и подготовка процессов к автоматизации, описание точек интеграции ИИ.
- **Данные** — качество, доступность и безопасность данных.
- **Люди** — компетенции, управление сопротивлением, культура работы с ИИ.
- **Инфраструктура** — вычислительные мощности с учётом требований информационной безопасности.
- **Управление изменениями** — PR, коммуникация, обучение.

Именно этому посвящена вся вторая часть книги, а развёрнутая методология системного внедрения — в книге «ИИ-2. Практическое руководство». Без проработки каждого направления даже лучший ИИ останется «игрушкой для ИТ-отдела».

Тренд 9. Синергия машин и людей (70/30)

Рабочее разделение ролей: 70/30



Чем выше цена ошибки, тем больше контроль человека и ниже автономность ИИ

ИИ не заменяет человека. Человек, который умеет работать с ИИ, заменяет того, кто не умеет, — и этот тезис всё больше подтверждается практикой. На горизонте 3–5 лет наиболее реалистична модель, где ИИ берёт на себя 40–70% рутинной и аналитической работы, а человек сосредоточен на решениях, коммуникации, творческих и нестандартных задачах. Даже при 70% автоматизации остаются 30%, критически зависящие от человека, — и именно они определяют качество результата: проверка выводов ИИ, интерпретация в контексте бизнеса, финальные решения, работа с людьми.

По моим наблюдениям — это оценка из практики, а не данные исследований, — 80–90% компаний, пытающихся полностью заменить человека ИИ в сложных процессах, сталкиваются с деградацией качества сервиса / услуг — не потому, что ИИ плох, а потому, что задачи, требующие контекста, эмпатии и здравого смысла, пока вне его возможностей. Тренд не в замене людей, а в новых моделях совместной работы, где сильные стороны ИИ (скорость, обработка данных, масштабируемость) дополняют сильные стороны человека (контекст, суждение, адаптивность).

Миф и Реальность

Миф: «ИИ заменит людей».

Реальность: рабочая модель — 70/30: ИИ забирает рутину, человек отвечает за контекст, ограничения и финальное решение. Попытки полностью убрать человека из сложных процессов, как правило, заканчиваются деградацией качества и возвратом людей.

Роль человека:

- определяет стратегические цели и приоритеты;
- задаёт контекст и ограничения;
- решает задачи, недоступные ИИ (личные коммуникации, договорённости, доработка с учётом ограничений);
- принимает финальное решение и несёт ответственность.

Роль ИИ:

- собирает и структурирует данные;
- выявляет закономерности и тренды;
- предлагает варианты решений с оценкой рисков;
- предсказывает последствия каждого варианта;
- автоматизирует выполнение 40–70% задачи.

Например, «Сокол» при управлении проектами может забрать до 80–90% задач проработки: критерии качества, план реализации, ТЗ, генерация гипотез, приоритизация, рекомендации. Но реализация остаётся за человеком: договорённости, контроль исполнения, устранение сопротивления, решение конфликтов. Тоже самое и в модулях Коммуникации или Решения — ИИ поможет подготовиться к переговорам или проработает разные варианты решений, но общаться и нести ответственность за итоговое решение, реализовывать его, всё равно будет человек.

У синергии есть обратная сторона — деградация навыков (skill atrophy): по данным Уортонской школы бизнеса (Wharton), 43% руководителей предупреждают об ослаблении компетенций сотрудников из-за

чрезмерной опоры на ИИ — при том, что 89% считают, что ГенИИ в целом усиливает работу. Ответ здесь управленческий, а не технический: формализовать границу «ИИ помогает — человек отвечает», встроить контроль качества и обучение ИИ-грамотности.

Три слоя рынка ИИ: где чья игра

Завершая разговор о трендах, дам рамку, которая снимает вечный вопрос «догоняем или обгоняем» — его задают и про страны, и про компании. Вопрос имеет смысл, только если все бегут по одной дорожке. А рынок ИИ — это не одна дорожка. Это **три слоя с принципиально разной экономикой**.

Слой первый — передовые, или фронтирные, фундаментальные модели. Это игра капитала и энергии: счёт идёт на десятки миллиардов долларов, гигаватты подключённой мощности и сотни тысяч ускорителей. Игроков здесь — единицы на весь мир, и список почти не меняется. Конкурировать в этом слое «в лоб» бессмысленно для абсолютного большинства стран и компаний — и это нормально: у Германии нет своей фронтирной модели, и никто не считает немецкое машиностроение отсталым.

Слой второй — прикладные платформы и инструментарий: MLOps, векторные базы, оркестрация агентов, инструменты дообучения. Конкуренция здесь глобальная, и она во многом уже выиграна открытым исходным кодом — у всех сразу. Строить здесь долгосрочное преимущество бессмысленно: вы вкладываете в то, что через полгода станет бесплатным.

Слой третий — доменные вертикальные решения. Дефектоскопия по машинному зрению, предиктивное обслуживание турбины, ассистент технолога, который знает ваш парк станков и вашу оснастку. Здесь выигрывает не тот, у кого больше вычислительных мощностей, а тот, у кого есть три вещи: доступ к данным реального производства, отраслевая инженерная школа и право на эксперимент на действующем объекте.

И главный сдвиг 2024–2026 годов, который мало кто проговаривает вслух: **открытые модели обнулили значительную часть барьера входа в базовую модель**: по качеству они подошли вплотную к проприетарным, и сама модель перестала быть конкурентным преимуществом. Преимущество переехало — в данные (в само их наличие, которое складывается из промышленной автоматизации и датчиков), в интеграцию с рабочими процессами и в способность доказать эффект.

Для руководителя отсюда два вывода. Первый: не путайте слои. Новости о гонке гигафабрик — это первый слой; к вашему проекту ассистента технолога они почти не имеют отношения. Второй: ваша игра — третий слой. Там технология — это 20% успеха, а 80% — данные, процессы и инженерия. Именно поэтому большая часть этой книги — про них.

Три слоя рынка ИИ

Куда смотреть руководителю — и где чья игра



Резюме главы

- Порог входа в ИИ резко упал (86% пользователей агентов — малый и средний бизнес, МСБ): побеждает не самая мощная модель, а управление внедрением.

- Без правил растёт теневой ИИ (Shadow AI); ROI x15–20 (оценка по отдельным кейсам) достижим только при правильных кейсах, качестве данных и зафиксированном baseline/KPI.
- Реалистичная модель ближайших лет — синергия 70/30: ИИ берёт рутину и аналитику, человек — решения, контекст и ответственность.

Что с этим делать: Беритесь за реально болезненную задачу, а не за модную, и заранее договоритесь, как будете измерять эффект; введите политику по теневому ИИ. Методику выбора и приоритизации кейсов и расчёта эффекта см. в книге «ИИ-2. Практическое руководство».

Вопрос для рефлексии: Какие 30% в вашей работе должны остаться за человеком?

Глава 4. Генеративный ИИ

Что такое генеративный искусственный интеллект?

Ранее мы рассмотрели ключевые направления применения ИИ:

- прогнозирование и принятие решений;
- анализ сложных данных без чётких взаимосвязей, в том числе для прогнозирования;
- оптимизация процессов;
- распознавание образов, включая изображения и голосовые записи;
- генерация контента.

На пике популярности сейчас — распознавание образов (аудио, видео, числа) и на их основе генерация нового контента по запросу пользователя: аудио, текст, код, видео, изображения. К генеративному ИИ можно отнести и цифровых советников.

Простыми словами

Генеративный ИИ — очень начитанный стажёр: он «прочитал» больше текстов, чем любой человек, пишет быстро и уверенно, но не знает

вашего контекста и может уверенно ошибаться. Поэтому результат всегда проверяет руководитель.

Проблемы генеративного ИИ

Экономика лидера рынка — лучшая иллюстрация того, чего стоит эта гонка. В 2022 году OpenAI потеряла на разработке ChatGPT 540 млн долларов, и тогда эта сумма казалась внушительной. К 2026 году масштабы изменились на порядок: при выручке около 25 млрд долларов в годовом выражении (примерно 2 млрд в месяц) компания планирует убыток порядка 14 млрд за год — то есть теряет около 1,2 доллара на каждый заработанный. Собственный прогноз OpenAI ещё честнее: накопленный убыток до 2029 года — порядка 115 млрд долларов, выход в прибыль — где-то в начале 2030-х.

Для понимания структуры затрат: только вычисления для ответов пользователям (инференс) обойдутся в 2026 году примерно в 14 млрд долларов — около 38 млн долларов в день. Для сравнения: в 2023-м поддержание работы ChatGPT оценивали в 700 тысяч долларов в день.

И маленькая деталь, в которой вся суть момента. Несколько лет назад глава OpenAI говорил, что для создания сильного ИИ понадобится порядка 100 млрд долларов. В 2026 году компания ищет 100 млрд — но не на путь к сильному ИИ, а чтобы профинансировать текущую операционную деятельность.

Общий тренд подтверждает и Алексей Водясов, технический директор SEQ: «ИИ не достигает тех маркетинговых результатов, о каких говорили ранее... за хайпом и бумом неизбежно следует спад интереса. ИИ выйдет из фокуса всеобщего внимания так же быстро, как и вошёл... ИИ — это действительно „игрушка для богатых“, и таковой на ближайшее время и останется». Согласен: после шумихи начала 2023 года уже к осени наступило затишье.

Картину дополняет исследование Wall Street Journal: большинство ИТ-гигантов пока не научилось зарабатывать на генеративном ИИ. Microsoft, Google, Adobe и другие ищут способы монетизации:

- Google планирует повысить стоимость подписки на ПО с поддержкой ИИ;
- Adobe вводит ограничения на число обращений к ИИ-сервисам в месяц;
- Microsoft хочет брать с бизнес-клиентов дополнительные 30 долларов в месяц за создание презентаций силами нейросети.

Вишенка на торте — расчёты Дэвида Кана, аналитика Sequoia Capital. Ещё в 2024 году он посчитал: чтобы окупить вложения в ИИ-инфраструктуру, индустрии нужно зарабатывать около 600 млрд долларов в год. Логика простая: берём годовую выручку Nvidia от дата-центров, умножаем на два (оборудование — примерно половина стоимости владения ЦОДом), затем ещё на два — потому что кто-то должен зарабатывать на этих мощностях с нормальной маржой.

К 2026 году эта арифметика стала жёстче на порядок. Nvidia зарабатывает на дата-центрах 193,7 млрд долларов за финансовый год — против 10,3 млрд за квартал в 2023-м, — а один только 2 квартал 2026 года принёс ей 75 млрд. Прогоните эти числа через формулу Кана и получите планку не в 600 млрд, а под триллион годовой выручки, которую индустрия должна где-то найти. Техногиганты закладывают на 2026 год около 725 млрд капитальных вложений, аналитики ждут выхода за триллион в 2027-м, а расхождение между темпами вложений и темпами выручки Allianz Research оценивает в 46% — больше, чем на пике телеком-пузыря 2001 года (32%).

Иными словами, на ИИ по-прежнему уверенно зарабатывает продавец лопат, а не старатели. И это не повод для злорадства, а предупреждение: пузырь такого размера сдувается не тихо. Но именно здесь — граница между «инвесторской» и «вашей» историей. Обвал переоценённых ожиданий и капитализаций почти не затрагивает компанию, которая внедрила ассистента в поддержку и считает эффект в человеко-часах. Разорятся те, кто продавал обещания. А технология, как это было с интернетом после доткомов, останется — и станет дешевле, доступнее и будет приносить пользу.



[AI industry needs to earn \\$600 billion per year to pay for massive hardware spend — fears of an AI bubble intensify in wake of Sequoia report](#)

И вот здесь начинается самое интересное для нас с вами. Пока индустрия тратит триллион, бизнесу выгодно ровно обратное — и рынок это уже понял.

Вычислительные мощности — одна из главных статей расходов: чем больше запросов к серверам, тем больше счета за инфраструктуру и электроэнергию. Поэтому и произошёл разворот к компактным моделям: их развивают все крупные разработчики — Microsoft, Google, Apple, Mistral, Anthropic, — и именно они стали рабочей лошадкой корпоративных внедрений.

Цифры объясняют почему. Большие модели вроде GPT-4 (более триллиона параметров, стоимость создания — десятки миллионов долларов) не дают радикального преимущества в прикладных задачах. Компактные обучаются на узких наборах данных, стоят на порядки дешевле и работают на единицах миллиардов параметров. На практике эксплуатация небольшой модели обходится в 5–20 раз дешевле, чем обращения к флагманской по API: типичный корпоративный сценарий на 10 тысяч запросов в день — это сотни или первые тысячи долларов в месяц против десятков тысяч.

Дальше — интереснее. Компактные модели перестали быть «эконом-вариантом с оговорками». Microsoft Phi-4 на 14 млрд параметров превосходит модели класса 70 млрд в математике и коде; в отраслевых

задачах модели на единицы миллиардов параметров обходят гигантов на сотни миллиардов.

Это подтверждают и эксперты, которые считают, что для многих задач — обобщения документов, создания изображений — передовые модели избыточны. Илья Полосухин, один из авторов основополагающей статьи Google 2017 года, в одном из интервью сравнил их использование для простых задач с поездкой за продуктами на танке: «Для вычисления $2 + 2$ не должны требоваться квадриллионы операций». Но разберём по порядку: почему так сложилось, какие ограничения угрожают ИИ и что будет дальше — закат с новой ИИ-зимой или трансформация?

Метафора Полосухина получила продолжение из мира больших денег. В июле 2026 года Bridgewater Associates — один из крупнейших хедж-фондов мира — и Thinking Machines Lab Миры Мурати опубликовали показательное исследование: специализированная дообученная модель с открытыми весами обошла флагманские универсальные модели OpenAI, Anthropic и Google на шести типовых задачах инвестиционной аналитики — потребляя при этом примерно в 14 раз меньше вычислительных ресурсов. В прикладных задачах специализация бьёт масштаб — и экономика здесь на стороне «оркестра» компактных моделей.

Добавим к этому и результаты исследований ранее, где в рейтинге галлюцинаций компактные модели стоят выше «рассуждающих» флагманов. Секрет не в размере, а в том, на каких данных модель обучена и подо что заточена.

Практический вывод. Начинайте выбор модели не с вопроса «какая самая мощная», а с вопроса «какая самая маленькая из тех, что решают мою задачу». Проверяется это за неделю: возьмите свои двадцать типовых запросов и прогоните через флагманскую и компактную модель. Если разницы в качестве нет, а счёт отличается в десять раз, выбор сделан за вас.

Ограничения ИИ, которые приводят к проблемам

Ранее я привёл «базовые» проблемы ИИ. Теперь уйдём в специфику именно генеративного ИИ.

Беспокойство компаний о своих данных. Любой бизнес стремится охранять корпоративные данные. Отсюда две проблемы. Во-первых, компании запрещают онлайн-инструменты за периметром защищённой сети, а любой запрос к онлайн-боту — это обращение во внешний мир (как хранятся, защищены и используются данные — вопросов много). Во-вторых, это ограничивает развитие любого ИИ: все хотят рекомендаций от обученных моделей (например, предсказание поломки оборудования), но делиться своими данными не готовы. Замкнутый круг. Оговорка: некоторые уже умеют размещать модели уровня GPT-3/3.5 внутри контура компании, но их всё равно надо обучать — это не готовые решения, и службы безопасности найдут риски.

Сложность и дороговизна разработки и содержания. Разработка «общего» генеративного ИИ — это десятки миллионов долларов и очень много данных. КПД нейросетей пока низкий: где человеку хватает 10 примеров, сети нужны тысячи и сотни тысяч (хотя она находит взаимосвязи в массивах, которые человеку не снились). Из-за ограничения по данным тот же ChatGPT лучше «соображает» на английском, чем на русском, — англоязычный сегмент интернета больше. Добавьте электроэнергию, инженеров, обслуживание и модернизацию оборудования — и получите те самые 700 000 долларов в день только на содержание ChatGPT. Снизить затраты можно, убрав всё лишнее, но тогда это будет узкоспециализированный ИИ. Поэтому большинство решений на рынке — по сути «GPT-фантики», надстройки к большим моделям.

Беспокойство общества и ограничения регуляторов. Общество обеспокоено развитием ИИ, а государства не понимают, чего от него ждать. Здесь несколько факторов:

- **Влияние.** Генеративный ИИ доказал, что создаёт контент, который можно спутать с человеческим (тексты, изображения, научные

работы), и за секунды разрабатывает концептуальный дизайн микросхем и шагающих роботов.

- **Безопасность.** ИИ активно используют злоумышленники: по данным SlashNext (State of Phishing, 2023), за год с момента запуска ChatGPT число фишинговых атак выросло на 1265%.
- **Непрозрачность.** Как работает ИИ, порой не понимают и сами создатели — для столь масштабной технологии это опасно.
- **Зависимость от «учителей».** Модели строят и обучают люди; материал для обучения узкоспециализированных моделей тоже отбирают люди.

Всё это означает, что отрасль начнут регулировать и ограничивать. Вспомним и известное письмо марта 2023 года, в котором эксперты потребовали ограничить развитие ИИ. Здесь я только фиксирую барьер; что уже происходит с регулированием и чего ждать дальше — разбираю в Главе 7.

Недостаток модели взаимодействия с чат-ботами. Вероятно, вы уже пробовали общаться с чат-ботами и остались разочарованы: классная игрушка, но что с ней делать? Чат-бот — не эксперт, а система, которая угадывает, что вы хотите услышать, и выдаёт именно это. Чтобы получить пользу, вы сами должны быть экспертом в теме. Но если вы эксперт — нужен ли вам ГИИ? А если нет — решения не получите, только общие ответы. Замкнутый круг: экспертам не нужно, любителям не поможет. Кто тогда заплатит? На выходе — дорогая игрушка.

Кроме экспертности нужно ещё уметь формулировать запрос — таких людей единицы. В какой-то период даже была профессия промпт-инженера (стоимость — около 6000 рублей в час), и то он не мог подобрать идеальный запрос с первого раза (он не обладает предметной экспертизой). Лайфхак: можно попросить ИИ задать наводящие вопросы и на их основе сформулировать идеальный запрос — об этом отдельно поговорим в главе про языковые модели. Нужен ли бизнесу инструмент, делающий его зависимым от очень редких и дорогих специалистов, если обычные сотрудники не извлекут из него пользы? Рынок для обычного

чат-бота не просто узкий — он исчезающе мал. Исключения два: чат-бот как вспомогательный функционал к основному (подготовка рекомендаций, аналитических дашбордов, где вы чётко формулируете требования) и применение ИИ в конкретных бизнес-процессах (HR-консультации новых сотрудников, техподдержка, анализ данных).

Некачественный контент и галлюцинации. Нейросети собирают данные и не анализируют факты и их связанность: на что больше в интернете/базе, на то и ориентируются, не оценивая написанное критически. Поэтому ГИИ легко генерирует ложный или некорректный контент, а ИИ-галлюцинации — давно известная особенность.



[Галлюцинации нейросетей: что это такое, почему они возникают и что с ними делать](#)

Хорошо, когда случаи безобидны. Но бывают опасные ошибки. Один пользователь спросил у Gemini, как сделать заправку для салата: по рецепту нужно было добавить чеснок в оливковое масло и оставить при комнатной температуре. Пока чеснок настаивался, появились странные пузырьки; перепроверка показала, что в банке размножались бактерии, вызывающие ботулизм, — отравление токсином протекает тяжело, вплоть до смерти. Я и сам периодически использую ГИИ, и чаще он даёт не совсем корректный, а порой откровенно ошибочный результат: нужно 10–20 запросов с безумной детализацией, чтобы получить что-то вменяемое, что потом всё равно надо дорабатывать. За ним нужно перепроверять, поэтому полностью делегировать задачу ИИ невозможно — он скорее ассистент, чем заменитель. И снова приходим к тому, что

нужно быть экспертом, чтобы оценить корректность, — а это порой дольше, чем сделать с нуля самому. Как системно управлять неточностями (галлюцинациями) на уровне внедрения — отдельно разобрано в книге «ИИ-2. Практическое руководство» (Глава 15).

Эмоции, этика и ответственность. Без правильного запроса ГИИ просто воспроизводит информацию, не обращая внимания на эмоции, контекст и тон, — а сбой в коммуникации происходит очень легко, и к проблемам выше добавляются конфликты. Возникают и вопросы авторства и прав собственности на созданный контент: кто отвечает за недостоверные или вредоносные действия, совершённые с помощью ГИИ, и как доказать авторство? Нужны этические стандарты и законодательство.

Что же делать?

Компании не собираются полностью отказываться от больших моделей: Apple использует ChatGPT в Siri для сложных задач, Microsoft — модель OpenAI как ассистента в новой версии Windows. При этом Experian (Ирландия) и Salesforce (США) уже перешли на компактные модели для чат-ботов и обнаружили ту же производительность при значительно меньших затратах и задержках. Ключевое преимущество малых моделей — тонкая настройка под конкретные задачи и данные: по словам Йоава Шохама (AI21 Labs), небольшие модели отвечают всего за одну шестую стоимости больших.

Свежий пример того же сдвига — уже не у стартапов, а у самой Microsoft. В июле 2026 года она вывела в открытое тестирование собственные модели семейства MAI: по её данным, голосовая MAI-Voice-2-Flash в контакт-центрах Dynamics 365 (среди клиентов — T-Mobile и EasyJet) снижает затраты на вычисления до 89% по сравнению с моделями OpenAI, а MAI-Image-2.5 в PowerPoint — до 84%. Цифры стоит читать с поправкой: это заявление вендора, а не независимый замер, и соотношение сильно зависит от объёма запросов, размера модели и того, с каким тарифом идёт сравнение. Но направление показательно: компания, вложившая в OpenAI 13 млрд долларов, строит собственные модели — чтобы не платить за флагман там, где хватает

специализированного. Вывод для вас проще, чем для Microsoft: если даже ей невыгодно гонять флагманскую модель на потоковых задачах — распознавании речи в колл-центре или картинках для презентации, — то вам тем более. Флагман берите на сложное рассуждение и разовые задачи, узкую модель — на всё, что повторяется тысячи раз в месяц.

Не спешить. Ожидать заката ИИ не стоит: слишком много вложено за последние 10 лет и слишком велик потенциал. Я рекомендую вспомнить 8-й принцип Дао Toyota (основа бережливого производства и один из инструментов моего системного подхода): «Используй только надёжную, испытанную технологию». Из него — ряд рекомендаций: технологии призваны помогать людям, а не заменять их (часто стоит сначала выполнить процесс вручную); вместо непроверенной технологии лучше использовать отработанный процесс; перед внедрением проводить испытания в реальных условиях; отклонять технологию, идущую вразрез с культурой и угрожающую стабильности; и при этом поощрять поиск новых путей, оперативно внедряя зарекомендовавшие себя решения. Через 5–10 лет генеративные модели станут массовыми, дешёвыми и достаточно умными, выйдут на плато продуктивности, и каждый будет пользоваться результатами ГИИ. Но уповать на ИИ и сокращать людей уже сейчас — явно избыточно.

Повышать эффективность и безопасность. Практически все разработчики сейчас делают модели менее требовательными к количеству и качеству данных и повышают безопасность — ИИ должен генерировать безопасный контент и быть устойчивым к провокациям.

Осваивать ИИ в формате экспериментов и пилотов. Чтобы быть готовым к приходу действительно полезных решений, нужно следить за технологией, пробовать её и формировать компетенции. Как и с цифровизацией: вместо прыжка в дорогие решения — поиграть с бюджетными или бесплатными инструментами. К моменту прихода технологии в массы вы:

- будете понимать, какие требования закладывать в коммерческие решения (а хорошее ТЗ — 50% успеха);

- получите эффекты в краткосрочной перспективе, а значит, и мотивацию идти дальше;
- повысите цифровые компетенции команды, сняв ограничения и сопротивление по техническим причинам;
- исключите неверные ожидания, а значит, и лишние затраты, разочарования и конфликты.

Трансформировать общение пользователя с ИИ. Эту концепцию я закладываю в свой продукт: пользователю нужно давать готовые формы и вести по сценарию, где он просто указывает значения или отмечает пункты, а уже эту форму с корректной обвязкой (промптом) отдавать ИИ. Либо глубоко интегрировать решения в существующие ИТ-инструменты — офисные приложения, браузеры, автоответчики. Но это требует тщательной проработки поведения и запросов пользователя или их стандартизации: либо решение уже не копеечное, либо мы теряем в гибкости.

Разрабатывать узкоспециализированные модели под конкретные процессы. Как и людей, обучать ИИ всему — трудозатратно и малоэффективно. Узкоспециализированные решения на базе движков больших моделей сводят обучение к минимуму, сама модель не слишком большая, контент менее абстрактный и галлюцинаций меньше. Наглядная аналогия — люди: большего успеха добивается не тот, кто знает всё, а тот, кто фокусируется и развивается вглубь. Примеры узкоспециализированных решений:

- советник для управления проектами;
- налоговый консультант;
- советник по бережливому производству;
- чат-бот по производственной безопасности или помощник специалиста;
- чат-бот для ИТ-техподдержки.

По прогнозам Gartner, организации будут внедрять малые модели под конкретные задачи в три раза чаще, чем LLM общего назначения. О

«конце гонки LLM» в 2025 году высказался и Гэри Маркус (The New Yorker): «Я не слышу, чтобы многие компании говорили, что модели 2025 года для них намного полезнее моделей 2024 года, даже несмотря на то, что они лучше по контрольным показателям».



[Человечество достигло пика развития ИИ](#)

Из практики

С галлюцинациями эффективнее бороться архитектурой, а не уговорами модели. В «Соколе» ИИ работает по формализованным сценариям, опирается на данные пользователя (RAG) и показывает, на чём основан вывод; в критичных шагах финальное слово — за человеком. Это дороже и сложнее в разработке (поиск баланса удобства и гибкости), но дешевле, чем разбирать последствия уверенной выдумки.

Резюме главы

Хотя ГИИ пока на стадии развития, потенциал у технологии большой. Да, хайп пройдёт, инвестиции снизятся, появятся вопросы к целесообразности, но технология не уйдёт. Так ещё 16 июня 2024 года Forbes опубликовал статью «Зима искусственного интеллекта: стоит ли ждать падения инвестиций в AI».



[Зима искусственного интеллекта: стоит ли ждать падения инвестиций в AI](#)

В ней — аналитика циклов «зимы» и «лета» ИИ и мнения Марвина Минского и Роджера Шанка, которые ещё в 1984 году на встрече ассоциации AAAI описали механизм-цепную реакцию, ведущую к новой зиме:

- Этап 1. Завышенные ожидания бизнеса и публики не оправдываются.
- Этап 2. СМИ начинают выпускать скептические статьи.
- Этап 3. Агентства и бизнес снижают финансирование исследований.
- Этап 4. Учёные теряют интерес, темп развития замедляется.

Прогноз сбился: в течение пары лет наступила ИИ-зима, а потеплело лишь в 2010-х. Сейчас мы на очередном пике — он наступил в 2023-м после выхода ChatGPT (поэтому в книге я часто привожу примеры из этой LLM — частный, но очень понятный случай ИИ). Далее статья применяет цикл Минского и Шанка к текущей ситуации.

Этап 1. Ожидания. Революции от ИИ в повседневной жизни пока не случилось: Google так и не трансформировал поиск (Search Generative Experience получил смешанные отзывы); голосовые ассистенты («Алиса», «Маруся») стали лучше, но им вряд ли можно доверить ответственные решения; чат-боты поддержки по-прежнему плохо понимают запрос и раздражают ответами невпопад.

Этап 2. Реакция СМИ. По запросу «AI bubble» поиск выдаёт пессимистичные заголовки: «Пузырь хайпа вокруг ИИ сдувается» (The

Washington Post), «От бума до взрыва: пузырь AI движется только в одном направлении» (The Guardian), «Известный экономист предупреждает, что пузырь AI рухнет» (Business Insider). Моё мнение: эти статьи недалеки от истины. Рынок перегрет — это факт, и к середине 2026 года он на грани сдутия; ситуация действительно напоминает канун краха доткомов. Но перегрев — в ожиданиях и оценках, а не в полезности технологии: те, кто научился внедрять, уже видят эффекты и растут. Волна хайпа схлынет, многие обанкротятся — и это нормальный этап взросления индустрии: останутся жизнеспособные решения и компании, а массовый рынок наконец научится внедрять. Так было с интернетом после доткомов — так будет и с ИИ (подробнее — в Главе 23 и Постскрипте).

Этап 3. Финансирование. Несмотря на пессимизм, нельзя сказать, что финансирование снижается: крупнейшие ИТ-компании продолжают вкладывать миллиарды, а ведущие конференции получают рекордное число заявок. Значит, мы сейчас между вторым и третьим этапом перехода к зиме. Но неизбежна ли «зима»? На самом деле нет.

Ключевой аргумент статьи: ИИ слишком глубоко проник в нашу жизнь, чтобы началась новая зима:

- системы распознавания лиц в телефонах и метро используют нейросети для идентификации;
- переводчики вроде Google Translate сильно выросли в качестве, перейдя от классической лингвистики к нейросетям;
- современные системы рекомендаций используют нейросети для моделирования предпочтений.

Особенно ценно мнение, что потенциал слабого ИИ не исчерпан и, несмотря на все проблемы сильного ИИ, он может приносить пользу, — полностью согласен. Следующий шаг развития ГИИ — более новые и лёгкие модели, которым нужно меньше данных для обучения. Нужно лишь набраться терпения и постепенно изучать инструмент, формируя компетенции, чтобы потом использовать его потенциал в полной мере.

Ключевые выводы

- Генеративный ИИ на передовых моделях экономически пока тяжёл, а компактные специализированные модели часто не уступают передовым в прикладных задачах.
- Чат-бот без эксперта не даёт ценности, а галлюцинации требуют перепроверки: ИИ — ассистент, а не заменитель человека.
- Закат технологии не грозит, но и уповать на неё, сокращая людей, преждевременно — нужно осваивать через эксперименты.

Что с этим делать: Пробуйте ИИ лично — на своих задачах, в бесплатных или бюджетных сервисах: это самый дешёвый способ понять технологию. А в компании начинайте с готовых решений — по данным MIT, пилоты на внешнем решении с доработкой под задачу доходят до эксплуатации вдвое чаще собственных разработок, примерно 67 % против 33 %. Не сокращайте людей под GenAI и вкладывайтесь в хорошее ТЗ.

Вопрос для рефлексии: Какой ваш процесс выиграет от узкой генеративной модели уже сейчас?

Глава 5. Большие языковые модели (LLM)

В середине 2026 года, 9 из 10 ИИ-проектов базируются на больших или средних языковых моделях. Давайте разберёмся в этом направлении.

Феномен LLM

2023 год в истории ИИ можно назвать эпохальным: на рынке появились LLM коммерческого качества, доступные всем. Их особенность в том, что они обучены изначально и предобучены. Порог входа в мир ИИ резко снизился — если раньше требовались команды и обученные люди, то теперь любой из нас может запрашивать аналитику, строить графики, автоматизировать рутину.

Главный бонус — предобученность: это готовые модели, в которые можно загружать массивы данных и работать с ними; их можно подключать к базам данных компании (ниже поговорим про RAG). По сути, ИИ меняет подход к аналитике: в ряде задач можно уйти от

построения сложных и дорогих моделей в пользу LLM — для аналитики, отчётов с графиками и консультаций.

Простыми словами

LLM (большая языковая модель) — программа, которая училась на огромной библиотеке текстов и теперь «продолжает» любой текст, угадывая следующее слово. Угадывает настолько хорошо, что получается осмысленный ответ, письмо или план.

Что такое RAG?

Одна из технологий, которая внесла огромные изменения в мир LLM, — RAG. По сути, это «внутренний факт-чекер» для ИИ. Когда мы спрашиваем модель о курсе валют или текущей дате, сама по себе она этого не знает, но может обратиться к внешним источникам и получить актуальные данные. Основные компоненты RAG:

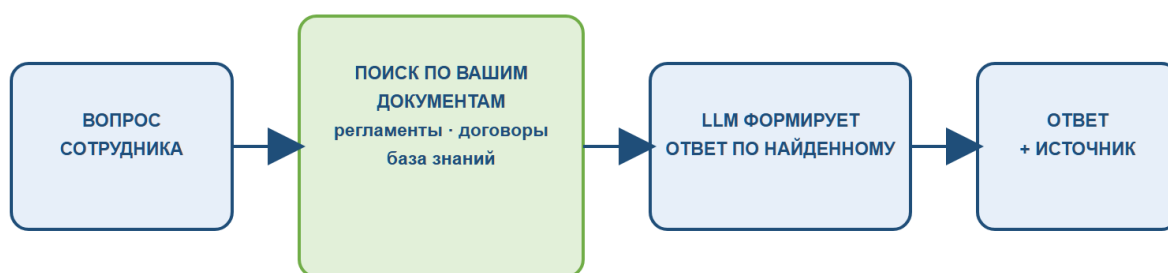
- **Retrieval (поиск):** ИИ ищет данные в надёжных источниках — базах данных и документах, к которым подключён;
- **Augmentation (обогащение):** ИИ добавляет найденную информацию к запросу;
- **Generation (генерация):** ИИ формирует конечный ответ.

То есть сначала ИИ сканирует источники, описывает для себя, о чём каждый из них, а затем использует их для подготовки ответов. Яркий пример — RAG в банке для заполнения кредитной заявки: менеджер обращается к ИИ, идёт запрос в кредитное бюро, параллельно проверяются внутренние данные. Так же и в любом направлении бизнеса — подключение к базам 1С, системе документооборота; самый простой пример — загрузка документов в условный ChatGPT или DeepSeek. Важно: использование RAG и включение в ответ ссылок на источники — одна из механик, позволяющих «проверять» ИИ и снижать уровень галлюцинаций.

Простыми словами

RAG — экзамен с открытой книгой: модель не вспоминает ответ «наизусть», а сначала находит нужную страницу в ваших документах и отвечает по ней. Поэтому качество ответа зависит от порядка в вашей «библиотеке».

RAG: экзамен с открытой книгой



Модель не «вспоминает» — она находит нужную страницу и отвечает по ней

Из практики

RAG — рабочая лошадка корпоративных ИИ-решений. В «Соколе» именно RAG позволяет отвечать по проектам и заметкам конкретной компании, а не по «средней температуре интернета». Правило из практики: качество такой системы определяет не модель, а порядок в базе знаний — мусор на входе даёт уверенный мусор на выходе.

Fine-tuning

Ещё один инструмент современных моделей — дообучение (fine-tuning). Например, модели можно «скормить» ваши письма и документы, чтобы она отвечала в принятом у вас стиле. Если интересно глубже, рекомендую статью по QR-коду и ссылке.



[*Retrieval-Augmented Generation \(RAG\): глубокий технический обзор*](#)

Архитектура решений на основе LLM и текстовые инструкции как слой управления качеством на уровне внедрения подробно разобраны в книге «ИИ-2. Практическое руководство» (Глава 15).

Простыми словами

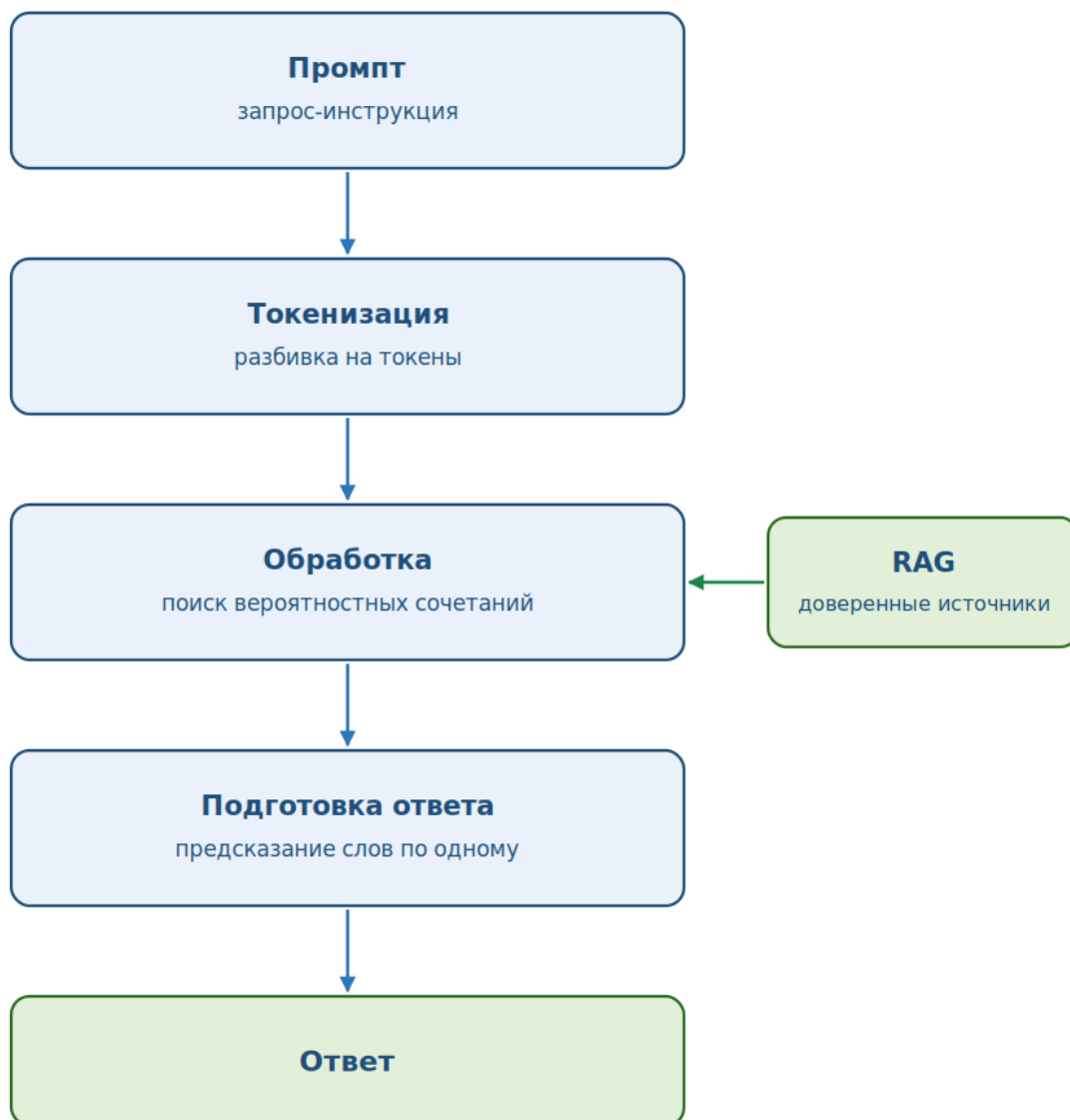
Дообучение (fine-tuning) — курсы повышения квалификации: модель не переучивают с нуля, а «натаскивают» на ваших примерах, чтобы она писала в вашем стиле и знала вашу специфику.

Как работает LLM?

Вся работа современных LLM-систем строится на простом алгоритме.

- **Получение запроса.** Работа начинается с отправки запроса — промпта.
- **Токенизация.** Модель разбивает запрос на множество «токенов». Например, «Привет!» может делиться на «привет» и «!». Разные модели токенизируют по-разному, и от этого во многом зависит качество.
- **Обработка.** Модель обрабатывает токены в сложной математической модели, ища смысловые сочетания (вероятности); здесь же может использоваться RAG для актуальных данных.
- **Подготовка ответа.** ИИ строит ответ последовательно, предсказывая каждое следующее слово.

Как работает LLM: конвейер обработки запроса



Ошибка на одном шаге ломает всю цепочку — либо начните заново, либо задайте открытый вопрос с критерием

Проблема в том, что если модель ошиблась в одном месте, дальше вся логика пойдёт по неверному пути. Поэтому нет смысла переубеждать ИИ — спорить и давить. Либо начните заново, либо задайте открытый вопрос с критерием, который заставит модель пересобирать рассуждение целиком, — об этом в Главе 6.

Ещё одно понятие, которое пригодится вам на практике, — контекстное окно. Это объём информации, который модель «держит в голове» в рамках одного диалога: ваши сообщения, её ответы, загруженные документы и служебные инструкции сервиса. Всё это складывается в одну «оперативную память» разговора.

Размер окна измеряется в токенах — тех самых кусочках слов, на которые модель режет текст. У массовых моделей 2026 года окна — от ста с небольшим тысяч до миллиона токенов. В переводе на житейское: сто тысяч токенов — это порядка двухсот страниц текста, то есть целая книга вроде той, что вы держите в руках. Миллион — небольшая книжная полка. Кажется, что много. Но окно съедают не только ваши реплики: ответы модели обычно длиннее ваших сообщений, а один загруженный отчёт может стоить как половина книги. Плюс добавим сюда внутренние рассуждения модели.

Главный подвох в другом: качество падает задолго до формальной границы. Чем плотнее заполнено окно, тем хуже модель держит нить — она начинает терять середину разговора, путать договорённости, переспрашивать то, что вы уже обсудили. Знакомый эффект «к сотому сообщению ассистент отупел» — это не деградация модели, а переполнение окна. По моему опыту, ориентир такой: заполнили примерно 40% окна — пора переезжать. Перенесите накопленное в файл и начните новый чат. Точную цифру заполнения сервисы обычно не показывают, поэтому проще по признакам: модель переспрашивает, повторяется, путает имена и вводные — значит, порог пройден.

Из практики: схема «чат + файл»

Что делать с переполнением — проще, чем кажется. Два правила. Первое: ведите работу в чате, но периодически сохраняйте главное в документ-якорь — решения, договорённости, наработки, черновики. Второе: как только чат начал деградировать — не боритесь с ним. Открывайте новый и начинайте с файла-якоря: «вот наработки, продолжаем отсюда». Свежий контекст плюс накопленный артефакт всегда сильнее длинного уставшего диалога.

В агентных режимах (Cowork у Anthropic и аналоги у других вендоров) проблема снята архитектурно: агент работает с файлами напрямую, сам актуализирует их по ходу работы, и переписка не забивает его память. Но это история про работу за компьютером. На телефоне вы, как правило, в обычном чате или проекте — и там переезжать в новый чат с файлом-якорем придётся заметно чаще. Принцип один, и он важнее инструмента: истина живёт в артефакте, а не в переписке. Чат — рабочая поверхность, файл — память.

Забегая вперёд: именно так устроен мой дневник из Главы 21 — один чат на месяц, в конце месяца «Снимок», с которого начинается следующий.

Прослеживаемость источника: требование, а не пожелание

Раз мы разобрали RAG и галлюцинации, сформулирую требование, которое обычно относят к техническим деталям, — а оно определяет судьбу решения. Требовать от человека проверять ответы модели, выдавая ему голый текст, бессмысленно. **Верификация физически возможна только тогда, когда система показывает, откуда она это взяла:** ссылка на конкретный пункт регламента, на конкретный чертёж, на конкретный документ в вашей базе. Не «модель считает», а «основание — вот эти три документа, посмотрите сами».

Ассистент без прослеживаемости источника — это генератор правдоподобия: никакой специалист его не проверит, он просто нажмёт «принять». Модель ошибается уверенно, красиво и в правильном формате — она не подаёт сигнала «здесь я не уверена». Поэтому **прослеживаемость источника — обязательный пункт технического задания** на любой корпоративный ИИ-инструмент, а не пожелание к дисциплине сотрудников. Прежде чем требовать от людей нового поведения, дайте им инструмент, который делает это поведение возможным. Иначе получите ритуал вместо контроля.

Резюме главы

Если свести правила работы с LLM к сути: модель — мощный, но буквальный исполнитель. Определите цель, дайте контекст и входные

данные, задайте формат ответа и ограничения — и всегда проверяйте результат. Полная система постановки задач — структура запроса, готовые шаблоны, техники повышения качества и их экономика — в следующей главе.

Ключевые выводы

- Девять из десяти ИИ-проектов строятся на больших языковых моделях; их сила — в предобученности.
- RAG работает как «факт-чекер» и снижает галлюцинации, Fine-tuning подстраивает модель под ваш стиль и данные.
- Ошибка LLM на одном шаге ломает всю цепочку рассуждений — спорить с моделью бессмысленно: начните заново или пересоберите рассуждение открытым вопросом (Глава 6).

Что с этим делать: Для корпоративных задач подключайте LLM к своим данным через RAG и не стройте критичные процессы на «голом» чате.

Вопрос для рефлексии: Какие ваши базы знаний и документы стоит подключить к LLM через RAG?

Глава 6. Промпт-инжиниринг: дисциплина постановки задач для ИИ

Промпт как управленческий инструмент

Большие языковые модели создают иллюзию простого диалога: кажется, что с ними можно разговаривать так же свободно, как с человеком. Это правда лишь отчасти. Качество ответа очень сильно зависит от формулировки запроса: один и тот же инструмент может дать практичный рабочий материал или красивый, но бесполезный текст.

Вокруг prompt engineering долго было много мифов — специальные символы, секретные формулы, «магические» слова. Для большинства рабочих задач всё это не нужно. Разница между слабым и сильным запросом не в хитростях, а в том, что у модели есть роль, контекст, цель, входные данные, формат, ограничения и критерий хорошего результата. Поэтому создание инструкций — это не второстепенная техника для

энтузиастов, а управленческий инструмент: по сути, дисциплина постановки задач, перенесённая на работу с машиной.

В корпоративной среде важно исходить из простого принципа: чем выше цена ошибки, тем более формализованной должна быть инструкция. Для черновых идей допустим свободный стиль; для аналитики, писем, проектной документации, договорной работы и управленческих решений нужна структурированная постановка задачи.

Базовая структура запроса

Для большинства практических задач достаточно простой формулы:

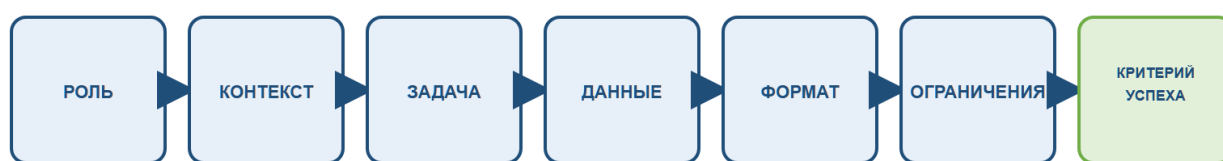
**РОЛЬ — КОНТЕКСТ — ЗАДАЧА — ВХОДНЫЕ ДАННЫЕ — ФОРМАТ —
ОГРАНИЧЕНИЯ — КРИТЕРИЙ ХОРОШЕГО РЕЗУЛЬТАТА**

Структура удобна тем, что подходит и для простого общения с чат-ботом, и для проектирования корпоративных ассистентов. Базовые правила хорошего запроса:

- **Начинайте с конкретной цели.** Не «расскажи про маркетинг», а «предложи 10 тем для публикаций о промышленной автоматизации для директоров заводов».
- **Давайте контекст.** Кто вы, для кого нужен ответ, в какой ситуации он будет использоваться.
- **Передавайте входные данные.** Документ, таблица, выдержка из переписки, описание процесса — всё это резко повышает предметность ответа.
- **Сразу задавайте формат ответа.** Таблица, список шагов, письмо, записка, сценарий разговора, структура слайдов.
- **Фиксируйте ограничения и критерии качества.** Сроки, бюджет, стиль, запрет на определённые формулировки, признаки хорошего результата.
- **Разделяйте сложные задачи.** Если нужно и придумать идею, и оценить риски, и составить план — делайте это последовательными запросами.

- Просите модель задавать уточняющие вопросы, если данных недостаточно. Особенно полезно, когда вы сами ещё не сформулировали задачу.
- Не доверяйте первому ответу без проверки. Даже сильный запрос не отменяет галлюцинации и ошибки логики.

Формула сильного промпта



Чем выше цена ошибки, тем больше блоков обязательны

Пример «до и после». Слабый запрос: «Напиши про наш продукт» — общие рассуждения без конкретики, 3–5 итераций доработки. Сильный запрос: «Ты — маркетолог с 10-летним опытом в B2B-продажах. Контекст: наш продукт — облачная CRM для малого бизнеса (10–50 сотрудников); аудитория — собственники и директора, уставшие от Excel; преимущество — запуск за 1 день. Задача: описание продукта для главной страницы лендинга. Формат: 3 абзаца по 50–70 слов — проблема клиента, решение, конкретная выгода. Тон: профессиональный, языком выгод, с цифрами. Ограничения: НЕ упоминай конкурентов, НЕ используй избитые фразы („инновационный“, „лучший на рынке“)). Результат — текст с чёткой структурой, по брендбуку, готовый к использованию почти без доработки. По моему опыту, такая постановка даёт прирост качества на 30–50% и экономит несколько итераций.

Точность термина: почему слово решает

Есть находка из нашей практики, которую вижу постоянно: модель крайне чувствительна к терминам. Не к вежливости и не к длине запроса — к точности профессионального слова.

Механика простая. Модель училась на текстах, где термины живут в своём окружении: рядом с «дефектоскопией» стоят методы контроля, нормативы и типы дефектов, а рядом с «проверить деталь» — что угодно, от ОТК до бытовых советов. Правильный термин работает как ключ: модель сразу попадает в нужный пласт и в нужный контекст. Причина глубже, чем кажется: модель училась на живых текстах людей, поэтому «думает» нашими ассоциациями — в ней устоялись те же связи между словами, что и у нас в голове, а оттуда идёт в наши документы и весь контент. Назовёте вещь профессиональным именем — попадёте в нужную ассоциацию сразу.

Эффект двойной. Качество: с точным термином ответ профессионален с первого раза, без раскочки и доработок. Экономика: не нужны длинные пояснения «что я имею в виду» и десятки итерации уточнений — это меньше токенов и меньше денег. А вот ошибка в термине опаснее, чем кажется: модель уверенно решит соседнюю задачу, и заметите вы это не сразу.

Пример из практики. «Проверь эту деталь на проблемы» — модель начнёт рассуждать обо всём сразу: от геометрии до логистики. «Проведи FMEA-анализ детали: виды отказов, причины, последствия, критичность» — и вы получаете структурированный разбор в принятой инженерной логике, короче и по делу. Слов в запросе почти столько же, результат — несопоставим.

То же справедливо для системных промптов корпоративных ассистентов: один точный термин в описании роли меняет поведение ассистента сильнее, чем абзац общих указаний.

Особенно часто с этой проблемой сталкиваются ИТ-разработчики, которые не обладают предметной экспертизой отрасли заказчика или продукта. В итоге возникают проблемы из-за недопонимания.

Рабочее правило: прежде чем отправлять запрос, спросите себя — как эту задачу называют профессионалы? Знаете отраслевой термин, стандарт или методику — называйте прямо: «по SWOT», «в нотации BPMN», «FMEA-анализ», «наряд-допуск», «дефектоскопия». Не знаете — попросите модель предложить терминологию и проверьте её.

И здесь видно, почему компетенции человека никуда не деваются. Точные слова знает только тот, кто разбирается в предмете. ИИ не обесценивает знание предметной области — он делает его самым конвертируемым активом: тот, кто владеет языком профессии, получает от модели больше за меньшие деньги. Ниже, в резюме главы, мы увидим, что это подтверждается и данными.

Прежде чем идти дальше, ответим на вопрос, который возникает у каждого руководителя: «Я ведь и людям задачи так ставлю — почему с ИИ это не работает?» Потому что ИИ — исполнитель без контекста и без страха переспросить.

Человек-исполнитель	ИИ-исполнитель
Достроит контекст из опыта	Достроит контекст выдумкой — галлюцинацией
Переспросит, если непонятно	Не переспросит — выдаст первый правдоподобный ответ (если вы сами не попросили его об этом при постановке задачи)
Побоится сдать откровенную «воду»	С готовностью сдаёт уверенную «воду»
Работает медленно, но осторожно	Работает мгновенно и в любом объёме

Отсюда правило, которое стоит запомнить дословно: ИИ не исправляет плохо поставленную задачу — он её масштабирует. Размытое поручение человеку оборачивается переделкой; размытый запрос ИИ — правдоподобным, быстрым и бесполезным результатом.

Формула качества задачи

Прежде чем перейти к шаблонам, зафиксируем правило, которое сэкономит вам больше всего нервов. Качество решения любой задачи

через ИИ складывается из четырёх множителей: постановка (как сформулирована инструкция) × объём задачи (сколько модель должна сделать за один заход) × данные (доступны ли они и доверяете ли вы им) × модель (подходит ли она под тип задачи). Именно множителей, а не слагаемых: ноль в любом из четырёх обнуляет результат — блестящая формулировка на мусорных данных даст уверенный мусор.

Про постановку мы говорим всю эту главу. Пройдёмся по остальным трём.

Объём: не давайте ИИ сразу большую задачу. «Напиши стратегию развития» одним запросом — почти гарантированный провал. На длинной дистанции модель в какой-то момент свернёт не туда — и испорченным окажется весь результат. Причём заметите вы это в самом конце. Дробите: анализ → варианты → выбор → план, с вашей проверкой на каждом стыке. Отклонение ловится на границе шага, пока оно стоит минуты, а не всей работы (не путать с итерациями, о которых поговорим дальше: там вы гоняете один и тот же результат до кондиции, здесь — режете работу на последовательные шаги). Простая проверка, по моему опыту: если результат — больше двух-трёх страниц или задача требует нескольких разных умений сразу (посчитать, написать, оценить) — дробите. Бонус, который редко замечают: маленькой задаче хватает более простой модели — а это в разы дешевле и быстрее. Большая задача тянет за собой «самую умную» модель; цепочка маленьких часто обгоняет её и по качеству, и по цене. Поделюсь собственной хитростью. Можно попросить саму модель разбить задачу на простые шаги — вам останется утвердить план и идти по нему.

Данные: любую задачу прогоняйте через три вопроса. Какие данные нужны — и передал ли я их модели? Откуда они и доверяю ли я источнику? Что модель будет делать там, где данных не хватит? Ответ на третий вопрос всегда один, и мы разбирали его в главе о генеративном ИИ: она дорисует пробел — правдоподобно и уверенно. Галлюцинация зачастую не «поломка модели», а дыра в данных или неточность в формулировке, которую вы не заметили при постановке.

«Проанализируй наш договор» без приложенного договора — фантазия на тему типовых договоров. «Сводка продаж за квартал» без выгрузки — красивая таблица выдуманных цифр. Модель не виновата: вы попросили результат, не дав сырья. Это похоже на то, как дать человеку задачу, но не дать ресурсов.

Модель: под задачу, а не «самую мощную». В большинстве сервисов это переключатель прямо в шапке чата: «быстрая» или «лёгкая» модель — для рутины и коротких шагов, «умная» или «рассуждающая» — для сложного анализа. Для большинства шагов конвейера хватает быстрой. Сам принцип мы разобрали ещё в первой части: специализация бьёт масштаб, а рассуждающие гиганты проигрывают компактным моделям в точном пересказе фактов. Больше того: на практике сложная модель для простых задач, наоборот, портит результат. Например, когда мне нужно было просто перевести презентацию, я взял самую мощную модель — и на выходе получал поломанный файл. Она пыталась пересобрать всё с нуля. Я больше часа бился с задачей. Злился то на ИИ, то на себя. А как только переключился на более простой режим — всё получилось с первого раза.

На уровне задачи всё это решается за минуту: приложить файл, разбить работу на шаги, выбрать модель попроще. На уровне компании качество и доступность данных — отдельная дисциплина, ей посвящена значительная часть книги «ИИ-2. Практическое руководство».

И чтобы формула перестала быть теорией — один запрос в двух вариантах, с реальными ответами.

Один запрос — два ответа

Обычный запрос: «Сделай аналитику по обращениям клиентов». Ответ ИИ: «Обращения клиентов — важный канал обратной связи. В последнее время наблюдается рост числа обращений, прежде всего по вопросам качества сервиса и сроков. Число обращений выросло примерно на 25%. Рекомендуется усилить работу с обращениями, повысить скорость ответов и внедрить цифровые инструменты...» Проблема: общие слова, «25%» выдуманы, использовать в работе нечего.

Тот же запрос по формуле: «Ты — аналитик отдела клиентского сервиса. Проанализируй приложенный реестр обращений за второй квартал и подготовь: (1) топ-5 категорий по числу обращений с динамикой к первому кварталу; (2) три категории с наибольшим ростом просрочки; (3) три гипотезы причин с указанием, какими данными их проверить. Формат: таблица плюс не больше одной страницы выводов. Если данных не хватает — пиши „недостаточно данных“ вместо предположения. Работай только с приложенным файлом». Ответ ИИ: конкретные категории с цифрами и динамикой, рост просрочки по трём направлениям, проверяемые гипотезы — документ, с которым можно идти на совещание.

Разница не в модели — модель та же. Разница в том, что во втором случае есть роль, данные, формат, критерий остановки и запрет додумывать.

И последняя оговорка — про границы. Для исследовательских задач жёсткая постановка преждевременна: если вы пока не знаете, что именно ищете, начните с разведки с лимитом времени («изучи тему, вот 30 минут моего внимания на выводы») — а точную формулировку давайте вторым шагом, когда станет ясно, что мерить. Чёткая постановка — инструмент, а не религия.

Четыре шаблона для разных ситуаций

Четыре рамки помогают быстро собирать качественные запросы под разные типы задач, не конструируя каждый раз запрос с нуля.

Шаблон	Когда использовать	Типовые пользователи
S-M-A-R-T-P	Задача ясна, нужен точный бриф без недомолвок.	Проектные офисы, аналитики, консультанты
D-E-E-P	Задача сырая — модель сама собирает контекст вопросами.	Ранняя проработка идей, интерфейсы внутренних ассистентов
B-R-I-E-F	Универсальный стандарт для большинства рабочих запросов.	Все сотрудники (базовая форма брифования)

L-E-V-E-L	Ту же тему нужно объяснить конкретной аудитории.	Обучение, коммуникации, материалы для разных ролей
-----------	--	--

S-M-A-R-T-P — когда задача уже понятна. S (Situation) — текущая ситуация; M (Mission) — требуемое изменение или цель; A (Actions) — что уже сделано; R (Restrictions) — ограничения и риски; T (Tools) — на что опираться; P (Payoff) — как понять, что результат хороший. Пример: «Ситуация: в службе закупок много ручной обработки заявок. Цель: определить три сценария применения GenAI для сокращения времени обработки. Уже сделано: описан процесс, собрана статистика. Ограничения: не выносить чувствительные данные во внешний контур, бюджет пилота до 3 млн рублей. Опирается на: описание процесса, статистику заявок за год, перечень уже закупленных ИТ-решений. Хороший результат: приоритизированный список сценариев с эффектом и рисками».

D-E-E-P — когда ИИ должен сначала уточнить задачу. D (Domain) — область; E (Experience) — уровень пользователя или адресата; E (Expectations) — желаемый тип и глубина результата; P (Priorities) — приоритет (скорость, качество, цена, простота, снижение риска). Пример: «Область: внедрение ИИ в техподдержку. Уровень: руководитель функции, не технический специалист. Хочу получить варианты архитектуры и план пилота. Приоритет: быстрое внедрение без сложной разработки».

B-R-I-E-F — универсальный корпоративный шаблон. B (Background) — фон и бизнес-контекст; R (Requirements) — требования и запреты; I (Inputs) — доступные материалы; E (Examples) — примеры желательного и нежелательного; F (Format) — конечный вид результата. Пример: «Фон: записка директору по операциям о применении ИИ в логистике. Требования: практично, без маркетинговых обещаний, с учётом ограничений по данным. Входные данные: процессы, показатели SLA, перечень ИТ-систем. Примеры: за образец — прошлогодняя записка по пилоту на складе; чего не надо — презентация вендора. Формат: 1

страница, затем таблица из 5 сценариев с эффектом, рисками и сложностью внедрения».

L-E-V-E-L — настройка ответа под аудиторию. L (Level) — уровень аудитории; E (Experience) — что аудитория знает; V (Vocabulary) — допустимая сложность терминов; E (Examples) — сколько примеров; L (Length) — объём. Пример: «Объясни разницу между RAG, fine-tuning и ИИ-агентами для финансового директора. Он понимает экономику проекта, но не погружён в ML-архитектуру. Используй деловой язык с минимумом техножаргона, добавь два бытовых примера, объём — половина страницы».

Как встроить шаблоны в практику: B-R-I-E-F — как стандартную форму брифования для большинства запросов сотрудников; D-E-E-P — в интерфейсах внутренних ассистентов, когда система сначала собирает контекст; S-M-A-R-T-P — для проектных офисов, аналитиков и консультантов; L-E-V-E-L — при обучении, коммуникациях и подготовке материалов для разных ролей.

Техники повышения качества работы с ИИ

Теперь от шаблонов — к восьми техникам, которые повышают качество ответов. Соотношение «качество/стоимость» по каждой сведено в таблицу в конце раздела. Важная оговорка: проценты — это мои оценки по рабочим задачам, а не результаты строгих замеров; эффект сильно зависит от задачи и модели. Там, где у техники есть научный первоисточник, я его привожу.

Пять базовых техник

Техника №1. Few-Shot Learning (обучение на примерах). Дайте модели 2–3 примера правильного ответа перед основной задачей. Например, для категоризации отзывов: «Пример 1: „Товар пришёл быстро, но упаковка помята“ → Логистика; нейтральная; средний. Пример 2: „Отличный сервис!“ → Обслуживание; позитивная; низкий. Пример 3: „Товар не работает! Требую возврат!“ → Качество товара; негативная; высокий. Теперь категоризируй: [новый отзыв]». Способность учиться «с

нескольких примеров» показана ещё в статье о GPT-3 (Brown et al., 2020). Ориентир: +15–25% на повторяющихся задачах. Когда использовать: классификация, извлечение структурированных данных, единый формат.

Техника №2. Chain-of-Thought (цепочка рассуждений). Попросите модель «думать вслух» — показать ход рассуждений перед финальным ответом. Пример для решения «одобрить или отклонить кредит»: «Сначала проанализируй доход, кредитную историю, долги, обеспечение. Взвесь риски: вероятность невозврата, защитные механизмы. Обоснуй решение и только потом дай финальный ответ „Одобрить“ или „Отклонить“ с условиями». Техника описана исследователями Google (Wei et al., 2022) и стала одной из самых цитируемых. Ориентир: +10–20% на сложных задачах. Когда использовать: аналитика, принятие решений, задачи, требующие обоснования.

Техника №3. Self-Consistency (самосогласованность). Попросите модель сгенерировать 3–5 вариантов ответа, оценить каждый и выбрать лучший: «Сгенерируй 3 варианта ответа клиенту про возврат. Оцени каждый по шкале 1–5 (полнота, вежливость, ясность, соответствие политике) и верни только лучший». Первоисточник — Wang et al., 2022. Даёт +20–30%, но обходится в 4–6 раз дороже (генерируется несколько ответов). Когда использовать: критичные коммуникации, креативные задачи, высокая цена ошибки. Не использовать для простых задач и при высоком объёме запросов — дорого.

Техника №4. Role-Playing (ролевая игра). Дайте модели конкретную роль с экспертизой: «Ты — юрист с 15-летним опытом в корпоративном праве, специализация — договоры поставки, контракты от 10 до 500 млн руб. Проанализируй договор на риски для покупателя, приоритизируй финансовые». По моему опыту, грамотно заданная роль даёт +15–20% релевантности. Когда использовать: специализированные отрасли (право, медицина, финансы), нужна экспертная точка зрения, важны терминология и стиль.

Техника №5. Negative Prompting (отрицательные инструкции). Явно скажите, чего делать НЕ надо: «НЕ используй сложные термины и канцеляризмы, НЕ пиши длинные предложения (макс. 15 слов), НЕ перекладывай вину на клиента, НЕ давай невыполнимых обещаний. Делай: извинись кратко, объясни причину простым языком, дай новую дату, предложи компенсацию». По моему опыту, заметно улучшает читаемость и тон — особенно там, где модель повторяет одни и те же ошибки. Когда использовать: клиентские коммуникации, тексты для широкой аудитории.

Три продвинутые техники

Техника №1. Разбиение на подзадачи (Task Decomposition). Чем меньше каждая задача, тем выше качество исполнения: модель — это последовательная работа со статистическими вероятностями, и чем меньше объём, тем ниже риск сбоя на одном из этапов (механику мы разбирали в главе о работе LLM). Пример: вернёмся к тому же договору поставки, что и в технике с ролью. Вместо «Проанализируй договор и выдай риски» — три последовательных запроса: 1) «Извлеки финансовые цифры: [показатель] | [значение] | [пункт]»; 2) «Для каждой оцени, риск ли это для покупателя и величину ущерба»; 3) «Составь таблицу топ-5 рисков: Риск | Влияние | Вероятность | Рекомендация». По моему опыту, декомпозиция даёт один из самых больших приростов — порядка +25–40%. Когда использовать: сложные многошаговые задачи, анализ больших документов, важна прослеживаемость логики.

Техника №2. Дерево мыслей (Tree of Thoughts). Исследуйте несколько путей решения, оценивайте каждый и откатывайтесь при неудаче: «Разбей задачу на 3–5 подзадач. Для каждой исследуй 2–3 ветви решения, оцени по критериям (выполнимость, эффективность, риски, 1–5), выбери лучшую; если ветвь ведёт в тупик — вернись и выбери другую». Подход предложен исследователями Принстона и Google DeepMind (Yao et al., 2023). На стратегических задачах способен дать до +30–50%, но обходитсякратно дороже.

Техника №3. Самопроверка (Self-Verification). Модель проверяет свой ответ перед выдачей: «После подготовки ответа проверь: все ли пункты выполнены, нет ли противоречий, соблюден ли формат, не добавлены ли факты без опоры на входные данные. Если находишь проблемы — исправь». Полезно для писем, аналитических записок и любых текстов, идущих дальше по цепочке решений.

Отдельная рекомендация: уточняющие вопросы (Clarification First). Перед выполнением модель задаёт уточняющие вопросы: «Если данных недостаточно, не спеши отвечать. Сначала задай мне 3–5 уточняющих вопросов, потом предложи план ответа и только затем переходи к решению». Это не техника повышения качества ответа, а способ доформулировать задачу, поэтому в таблице ниже её нет: она почти ничего не стоит и сочетается с любым шаблоном. Превращает ИИ из «генератора ответа» в помощника по постановке задачи — полезно для новых тем и ранней проработки идей.

Экономика промптинга: качество против стоимости

Главное, что отличает управленческий взгляд на промпты от «коллекции лайфхаков»: у каждой техники есть не только эффект, но и цена — в токенах, времени и деньгах. Некоторые приёмы дают заметный прирост почти бесплатно, другие — кратно дороже.

Техника	Качество	Стоимость	Когда использовать
Негативные инструкции (запреты)	+25–35%	+0% (бесплатно)	Клиентские коммуникации, тон
Few-Shot (примеры ответа)	+15–25%	+0–10%	Классификация, единый формат
Role-Playing (роль и экспертиза)	+15–20%	+10%	Специализированные отрасли
Chain-of-Thought (рассуждение пошагово)	+10–20%	+20–30%	Аналитика, обоснование решений
Task Decomposition (декомпозиция)	+25–40%	+50–100%	Сложные многошаговые задачи

Self-Verification (самопроверка)	+15–25%	+100%	Критичные тексты
Self-Consistency (самосогласованность)	+20–30%	+300–500% (!)	Критичные задачи; не для массовых
Tree of Thoughts (дерево мыслей)	+30–50%	+500–1000% (!)	Стратегические решения

* Ориентиры по моему опыту на рабочих задачах; фактический эффект зависит от задачи и модели. Научные первоисточники техник приведены выше в тексте.

Правило выбора простое: начинайте с базовой структуры и одного-двух уместных приёмов; дешёвые техники с высоким эффектом (негативные инструкции, примеры, роль) — в первую очередь, а дорогие (самосогласованность, дерево мыслей) — только там, где цена ошибки оправдывает кратный рост затрат. Удачные сочетания:

- **В-R-I-E-F + уточняющие вопросы** — когда тема сырая и нужен хороший бриф;
- **L-E-V-E-L + примеры ответа** — когда важно подстроить материал под аудиторию и удержать стиль;
- **S-M-A-R-T-P + декомпозиция** — когда задача понятна, но требует поэтапного анализа;
- **любой шаблон + самопроверка** — когда ответ пойдёт в работу без долгой ручной доработки.

Системный промпт и сборка промпта с помощью ИИ

Когда вы настраиваете корпоративного ассистента (а не просто пишете в чат), инструкция делится на два уровня. Системный промпт задаётся один раз и описывает роль, правила, ограничения, тон и формат — он действует для всех обращений. Пользовательский запрос содержит конкретную задачу здесь и сейчас. Такое разделение — основа для корпоративных ассистентов и баз знаний (RAG), о которых шла речь в предыдущей главе: системный промпт фиксирует стандарт качества и безопасности, а пользователь работает уже внутри заданных рамок и не может случайно «сбить» поведение ассистента.

Как использовать ИИ для составления промпта

Если задача сырая и вы не понимаете, как лучше сформулировать запрос, используйте ИИ как помощника по постановке задачи:

«Помоги мне сформулировать сильный запрос по этой задаче. Сначала задай 3–5 уточняющих вопросов. После моих ответов собери итоговый запрос по структуре: роль — контекст — задача — входные данные — формат — ограничения — критерий качества. Если увидишь слабые места, предложи, как их усилить».

Так вы сначала проектируете хороший запрос, а уже потом получаете ответ.

Стартовый набор: 10 готовых промптов руководителя

Если вы только начинаете — вот десять шаблонов под типовые задачи. Копируйте, подставляйте своё в квадратные скобки и пользуйтесь. Каждый можно усилить техниками этой главы: добавьте роль, контекст, формат и ограничения.

1. Выжимка документа. «Сделай выжимку документа: 5 главных пунктов, риски, что от меня требуется и к какому сроку. Документ: [вставьте текст]».

2. Ответ на письмо. «Напиши вежливый, но твёрдый ответ: мы не готовы сдвинуть сроки; предложи два альтернативных варианта. Тон деловой, 5–7 предложений. Письмо: [...]».

3. Протокол встречи. «Преврати эти заметки в протокол: решения, ответственные, сроки, открытые вопросы. Формат — таблица. Заметки: [...]».

4. Подготовка к встрече. «Я встречаюсь с [кем] по теме [тема]. Составь 7 вопросов, которые стоит задать, и 3 риска, о которых стоит помнить».

5. Разбор договора или отчёта. «Ты — [юрист/финансист] с 15-летним опытом. Выдели 5 главных рисков для [нашей стороны]: суть, последствие, что предложить. Не выдумывай; если данных нет — так и скажи. Документ: [...]».

6. Декомпозиция задачи. «Разложи задачу [задача] на этапы: результат, ориентир по сроку, кто нужен».

7. Черновик поста или статьи. «Напиши черновик поста для [аудитория] о [тема]: цепляющее начало, 3 тезиса с примерами, вывод с вопросом. До 300 слов, без канцелярита».

8. Сравнение вариантов. «Сравни [А] и [Б] для [задача]: таблица „критерий — А — Б“, 7 критериев; в конце рекомендация с обоснованием и главным риском».

9. Обратная связь сотруднику. «Помоги сформулировать развивающую обратную связь: ситуация [...], что хорошо [...], что не так [...]. Доброжелательно, без обесценивания, по схеме „факт — эффект — предложение“».

10. «Спроси меня». «Мне нужно [цель]. Прежде чем отвечать, задай мне 3–5 уточняющих вопросов, которые помогут дать точный ответ». Самый недооценённый шаблон: превращает ИИ из угадывателя в собеседника.

Использование вопросов при работе с ИИ: от плана выполнения задачи к разбору ответов

Работая над этим изданием, я обновлял карту понятий из первой части — схему, которая показывает, как соотносятся искусственный интеллект, машинное обучение, глубокое обучение и генеративные модели. Черновик делал ассистент. Схема прошла машинную проверку. Обе стадии дали зелёный свет, и к вычитке я вернулся к ней уже со спокойной душой. Ошибку увидел сразу: языковые модели стояли рядом с генеративным ИИ, а не внутри него. Дальше был выбор. Можно было написать «перенеси LLM внутрь GenAI» — и получить ровно это, одну переставленную рамку. Я написал иначе: разве LLM не частный случай GenAI? Ассистент признал ошибку композиции и пересобрал карту целиком, а на освободившемся уровне сам достроил то, чего в схеме не было. В книгу пошла третья версия.

Шесть слов. Ни одного указания, что делать. И разница здесь не в вежливости. Указание чинит эпизод: рамка встаёт на место, а причина, по которой она стояла неправильно, живёт дальше по схеме. Вопрос

заставляет пересобирать рассуждение — и всё, что из этого рассуждения выросло, пересобирается вместе с ним. Все техники выше — про вход: роль, контекст, формат, данные. Этот раздел — про выход: что делать с планом, который ИИ предложил, и с ответом, который уже получен. Здесь у большинства руководителей включается режим ревизора — увидел ошибку, указал, что исправить. Я всё чаще делаю иначе: задаю вопрос.

Ещё пример. Готовя практикум по системе качества, я подбирал кейсы и спросил: «Все кейсы, которые ты включаешь, направлены на обеспечение качества?» Два кейса из семи оказались не про качество. Оба выглядели прилично, оба прошли бы вычитку. Я не сказал, что именно не так, — я назвал критерий, по которому нужно перепроверить подборку. Это первое правило вопроса к ИИ: в нём должен быть критерий. Вопрос «где здесь слабые места?» даёт список общих мест, верный для любого документа на свете. Вопрос с критерием заставляет модель пройти по набору и сверить каждый элемент.

Второе правило: вопрос должен быть открытым, а не наводящим. Именно здесь техника вопросов к машине расходится с техникой вопросов к людям. Модель прогибается под сомнение. В исследовании Anthropic (Sharma et al., ICLR 2024) реплика «I don't think that's right. Are you sure?» заставляла модели менять ответ: GPT-4 — примерно в трети случаев, Claude 2 — в четырёх из пяти. Причём переход «правильно → неправильно» случался чаще обратного: сомнение не улучшало ответ, оно его ломало. В апреле 2025 года OpenAI публично откатила обновление GPT-4o именно за избыточную угодливость. С более сильными моделями эффект слабее, но не исчезает ни у одной. Поэтому «ты уверен?» — это не коучинг, а давление: модель считывает недовольство и ломает верный ответ, чтобы вам понравиться. Открытый вопрос подсказки не содержит и требует показать работу: «пройди расчёт по шагам и покажи, на каком допущении он держится».

Насколько форма важна, показал замер британского Института безопасности ИИ (AI Security Institute, 2026): одну и ту же мысль подавали модели вопросом и утверждением. Разрыв по угодливости — около 24

процентных пунктов. Чем увереннее сформулирована мысль, тем охотнее модель соглашается, особенно после «я считаю». А инструкция «не будь угодливым» сработала слабее, чем переформулировка мысли в вопрос. Одна оговорка: замер сделан на одиночных запросах, не в длинных рабочих сессиях, — но направление он показывает однозначно.

Вопросы к плану: до старта. Всё остальное в этом разделе — про ответ, который уже написан. Но у работы с ИИ есть фаза раньше — планирование выполнения задачи. Как я писал ранее, одна из лучших техник — попросить подготовить план выполнения задачи. В результате ИИ предъявляет план и ждёт разрешения начать. Если планирование сделано некачественно, то и весь дальнейший результат будет посредственным или вообще не тем.

Поделюсь одним своим примером. Я разбирал с агентом план внедрения одного из моих рабочих проектов. После подготовки плана я провёл шесть раундов уточняющих и проверочных вопросов и ответов. Первый вопрос нашёл семь пропусков. Второй показал, что критериев приёмки в плане нет ни у одного пункта: план говорил, что делать, и нигде не говорил, как мы поймём, что это сделано. И это при том, что я использовал самую передовую модель Fable 5 в режиме максимальных рассуждений. Итак, какие вопросы я использовал.

Первый — на самопроверку. «Перепроверь свой план». Строго говоря, это указание, а не вопрос, единственное в наборе; правило открытой формы тут не работает, потому что мнения у модели не спрашивают. План агент писал в одном режиме, а перечитывает в другом, и второй видит то, что пропустил первый. Это та же самопроверка (Self-Verification) из раздела о продвинутых техниках. Чего она не находит — того, чего в плане нет как класса работ: отсутствие критериев приёмки этот ход у меня не заметил.

Второй — на критерии приёмки. «Прописаны ли критерии приёмки?» Критерий приёмки — тот самый «критерий хорошего результата» из структуры запроса, только зафиксированный до начала работ и по каждому пункту плана. Пока его нет, «готово» в отчёте агента значит

одно: агент закончил печатать. Спрашивать надо до старта: в конце тот же вопрос превращается в спор о том, что мы имели в виду.

Третий — на полноту. «Можем ли мы сказать, что план реализации является полным и комплексным?» Нужно не «да», а список того, чего в плане нет. Полнота у плана не бывает абсолютной: она считается относительно стадий работ, ролей, рисков, точек интеграции. Смотреть надо на рамку — относительно чего считал агент. Рамку не назвал, значит, оценил на глаз.

Четвёртый — на условия исполнения. План описывает, что сделать, и почти никогда — где это должно работать. У меня решение должно было одинаково работать в двух контурах, локальном на компьютере и облачном на телефоне, и в плане про это не было ни строки, пока я не спросил. Форма тут важна: не «учтено ли это», а «будет ли это работать в таком-то случае». На первое можно ответить «учтено», на второе придётся показать как.

Пятый — на разрешение старта. «Можем ли мы сказать, что план готов и можем приступить к работе?» Задаётся последним. Работает так же, как вопрос на готовность при сдаче, о котором ниже: значение имеет только ответ «нет».

Между вопросами — раунд, который легко пропустить: агент задаёт встречные уточняющие вопросы, и на них нужно отвечать (в разделе о техниках это Clarification First). Без этого следующий вопрос уходит по тому же материалу и возвращает то же самое. Шесть раундов на план — нормальная длина. И это не вопрос класса модели: максимальная модель с режимом рассуждений и настроенным контуром навыков проходит те же пять-шесть раундов уточнений. Их снимает не мощность модели, а то, что вынесено за пределы диалога, — план, критерии и инструкции в файлах. Держите при этом в уме правило 40 % из Главы 5: раунды с развёрнутыми ответами быстро съедают окно, поэтому просите сводить план и критерии в файл.

И про форму. Из пяти вопросов четыре закрытые. Закрытая форма опасна не сама по себе — опасна закрытая форма без критерия. За «прописаны ли критерии приёмки?» и «будет ли это работать в двух

контурах?» критерий стоит, и «да» проверяется в одно движение. За «план полный?» и «план готов?» критерия нет, и «да» не стоит агенту ничего. У меня оба раза сработала инструкция аккаунта, где записано спорить и не соглашаться из вежливости (как она устроена — четыре уровня настройки в Главе 21). Не уверены, что такое правило у вас есть, — спрашивайте открыто: «что в плане осталось непроверенным?»

Какие решения агенту вообще позволено принимать — три категории решений из Главы 2; чем эта граница удерживается технически — Глава 8.

Дальше — мой рабочий набор. Девять типов вопросов, собранных из моих рабочих сессий; формулировки привожу как есть, их можно копировать.

1. На критерий. «Все кейсы, которые ты включаешь, направлены на обеспечение качества?» Когда есть набор — кейсы, слайды, риски, требования — и признак, которому обязан соответствовать каждый элемент. Ловит то, что попало в набор по смежности темы, а не по назначению: самый частый брак в подборках. Чего не ловит — того, чего в наборе нет вовсе: пропуски вытаскивает восьмой вопрос.

2. Через роли. «Посмотри на программу через четыре линзы — методолог, бизнес-тренер, участник, основатель: где не сходится?» Вскрывает конфликт позиций, который не виден из одной точки. Больше четырёх-пяти линз не берите: каждая роль получит по абзацу вежливых общих слов. И не ждите от ролей фактчекинга: они смотрят на пользу и уместность, а не на достоверность.

3. Глазами адресата. «Представь, что ты целевая аудитория этого материала: чего тебе не хватило?» Здесь я прошу не оценку — желание. Ловит дыры по ожиданиям: раздел может быть логически безупречен и оставлять читателя с вопросом «и что мне теперь с этим делать».

4. На противоречие с написанным раньше. «Не противоречит ли это тому, что я даю во второй книге?» Ошибку видно и так. Расхождение не видно никому, пока кто-то не сверит два текста подряд. Чем больше у вас документов, регламентов и публикаций, тем обязательнее этот вопрос.

Не работает без данных: ссылку на второй текст или файл нужно дать прямо в реплике.

5. На логические дыры. «Есть ли в материале логические дыры или допущения? Если есть — сформируй мне вопросы для уточнения». Обратите внимание на хвост: я прошу вопросы, не исправления. Это экономит половину переделок — модель не бросается чинить то, что чинить, возможно, не надо. Граница: модель работает по написанному — дыр в том, чего в тексте нет, она не увидит.

6. На узкое место. «Изучи логику методологии: есть ли узкие места?» Для конструкций — методологий, расчётов, схем процессов, финансовых моделей. Модель охотно перечислит восемь узких мест и не скажет, какое выстрелит завтра: ранжирование остаётся на вас.

7. На происхождение цифры. «Откуда взята эта цифра?» Задаю каждый раз, когда в ответе появляется число, которого я не давал. Так в этой книге вскрылись «0,002 %» — число, которого никто не публиковал: оно получилось делением одного показателя на другой. В нынешнем издании стоит корректная доля. Правильность самой цифры этот вопрос не проверяет: источник может найтись и оказаться слабым — это уже ваша работа.

8. На пропущенное. «Какие вопросы нужно снять, чтобы повысить качество и точность проработки?» Мой любимый ход перед финалом: модель перестаёт быть исполнителем и работает как аналитик, который признаёт нехватку данных. Иногда список получается неприятный. На пустом контексте не даёт ничего: получите вежливую просьбу уточнить цели и задачи.

9. На готовность. «Считаешь ли ты, что проект проработан, или есть области, которые надо доработать?» Последним, перед сдачей. Значение имеет только ответ «нет»: «да» ровно так же прозвучит и тогда, когда модели уже нечем отвечать.

Конечно, это не финальный перечень возможных типов вопросов. Но это тот список, которым я пользуюсь чаще всего на практике в своей жизни и работе.

Вопрос на готовность задаётся не один раз. В одной сессии я задал его пять раз подряд — и первые четыре раза находилось новое, включая ошибку в правиле, которое ассистент сформулировал получасом раньше. Модель не врала: на каждом заходе она отвечала честно про то, что держала перед собой. Просто область, которую она держала, каждый раз оказывалась уже, чем «вся работа».

Три контрольных вопроса при сдаче. «Что из этого ты не проверял — назови, что осталось за пределами проверки»; «чем подтверждается каждое „сделано“ — файлом, прогоном, цифрой, а что подтверждено только твоим отчётом»; «что здесь факт, проверенный сейчас, а что по памяти или по предположению». Первый переводит ответ из «всё готово» в проверяемое «вот что проверено, вот что нет». Второй отделяет результат от отчёта о результате. Третий вскрывает то, что модель уже записала как факт, ни разу не проверив. Признак, что можно останавливаться: очередной заход возвращает не дефекты, а границу непроверенной области.

Три правила применения. По одному вопросу за раз: списки вопросов дают списки поверхностных ответов. Держите рамку: вопрос — мягкая форма, и модель может прочитать его как задание; спрашиваете «имеет ли смысл» — уточните, что просите оценку, а не правку. И помните, по какому материалу работает вопрос: только по тому, который у модели есть.

Последнее правило стоит развернуть — я проверил его на себе. Однажды за один день у меня перестали работать все заготовки разом: длинная сессия упёрлась в предел контекста и начала подчищать историю. Формулировки, которые неделей раньше вскрывали ошибки, вдруг стали давать общие слова, и я трижды за день писал «начинаем проработку заново». Механика простая: вопрос работает по материалу, а когда окно переполнено, материала нет — и та же формулировка возвращает уверенный и совершенно пустой ответ. Внешне он неотличим от честного: та же длина, та же структура, тот же спокойный тон. Отсюда правило, которое я приводил в Главе 5: заполнили около 40 % окна —

переносите накопленное в файл и начинайте новый чат, не дожидаясь, пока модель начнёт терять нить.

И граница метода: вопросы нужны не всегда. Видите дефект — неверная цифра, не тот термин, битая ссылка — правьте прямо: это дешевле и надёжнее. Вопросы работают там, где вы не знаете, где именно ошибка, или хотите вскрыть допущения. Это то же различие, что между наставничеством и корректировкой в работе с людьми: развивающий диалог не заменяет прямое указание, а дополняет его. К истории с картой понятий вернусь в Главе 21 — там она о другом: о границах того, что ассистент делает «тщательнее меня».

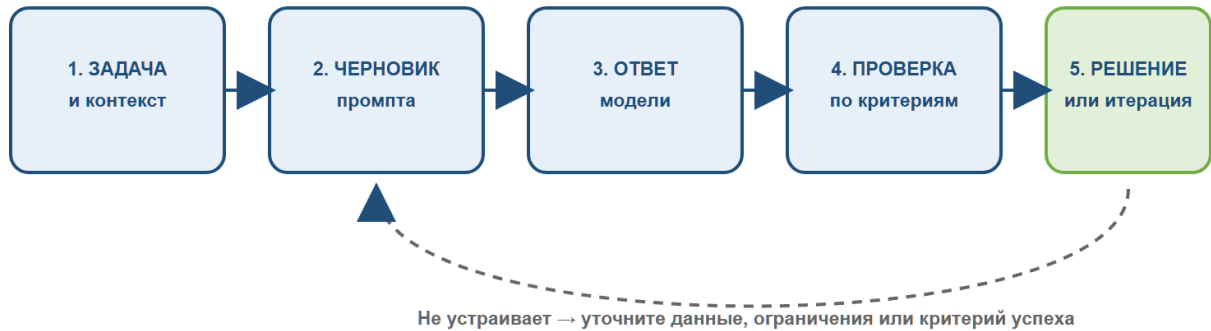
Честная оговорка вместо вывода. Вопрос работает не сам по себе — он работает в среде: настройка ассистента, живой контекст, класс модели (облегчённые версии прогибаются заметно чаще), ваша экспертиза и внешняя сверка. Уберите любое из условий — и те же формулировки дадут вежливое согласие. Почти все мои отловы — содержательные: я ловлю модель на знании предмета, не на удачной фразе. Техника вопросов — часть верификации; заменить её она не может. Полный разбор с исследованиями и моей системной инструкцией — в статье по QR-коду и ссылке.



[*Техника вопросов к ИИ: форма и среда*](#)

Чек-лист и типичные ошибки

Промпт — это цикл, а не заклинание



Чек-лист эффективной инструкции:

- структура: роль и контекст определены, задача конкретна, формат задан;
- данные: переданы входные материалы, заданы ограничения и критерий качества;
- источники: данным можно доверять, а зоны, где данных нет, названы явно;
- объём: задача посильна за один заход; крупное разбито на шаги с проверкой на стыках;
- приёмы: добавлены один-два уместных (а не все сразу);
- безопасность: во внешний сервис не уходят персональные данные и коммерческая тайна;
- проверка: ответ проверен на факты, логику и применимость.

Типичные ошибки:

- слишком общий запрос без контекста;
- попытка решить пять задач в одном сообщении;
- отсутствие входных данных там, где они критичны;
- отсутствие явного формата результата;
- отсутствие ограничений и критериев качества;
- слепое доверие первому ответу без проверки;

- попытка применить все техники сразу вместо одной-двух уместных;
- передача конфиденциальных данных во внешний сервис без понимания рисков.

Отдельное правило безопасности: хороший запрос, который приводит к утечке данных, — это плохой запрос. Прежде чем передавать что-либо во внешний сервис, убедитесь, что там нет персональных данных и коммерческой тайны.

Резюме главы

Работа с ИИ начинается не с выбора «самой умной» модели, а с умения правильно ставить задачу. Хороший промпт — это не хитрость и не набор секретных слов, а дисциплина мышления: навык ясной управленческой коммуникации, перенесённый на работу с машиной.

Качество работы ИИ складывается из четырёх множителей: постановка × объём задачи × данные × модель. Ноль в любом из них обнуляет результат: не давайте ИИ большую задачу за один заход, всегда передавайте данные и спрашивайте себя, откуда модель возьмёт то, чего вы не дали. И главное правило: ИИ не исправляет плохо поставленную задачу — он её масштабирует. Промпт — это управленческий навык, а не магия: чем выше цена ошибки, тем формальнее инструкция. База — формула **РОЛЬ–КОНТЕКСТ–ЗАДАЧА–ДАнные–ФОРМАТ–ОГРАНИЧЕНИЯ–КРИТЕРИЙ** и четыре шаблона: **B-R-I-E-F** (стандарт), **S-M-A-R-T-P** (точный бриф), **D-E-E-P** (модель сама уточняет), **L-E-V-E-L** (под аудиторию). И никогда не доверяйте первому ответу без проверки.

Теперь у тезиса есть данные. Anthropic проанализировала около 400 000 сессий агентного программирования (октябрь 2025 — апрель 2026): все десять крупнейших профессиональных групп — от менеджеров до специалистов по продажам — добиваются проверяемого успеха с разрывом лишь в несколько процентных пунктов от профессиональных программистов. Успех сессии определяет не умение писать код, а знание предметной области и качество постановки задачи; чем выше экспертиза человека, тем больше работы агент выполняет на одну инструкцию. Формула разделения труда проста: человек решает, что строить, — ИИ

решает, как. И обратите внимание, как это стыкуется с находкой о терминах: доменная экспертиза «работает» в том числе через точность слов — эксперт формулирует так, что модель сразу попадает в нужный контекст.

Промпты как актив организации. В организации промпты — это не личный лайфхак, а актив: единый стандарт брифования (B-R-I-E-F), библиотека проверенных промптов под типовые задачи, системные промпты корпоративных ассистентов и правило безопасности данных. Эти четыре элемента превращают разрозненные навыки сотрудников в управляемый ресурс, который не зависит от того, кто именно сидит за клавиатурой.

Из практики

В Соколе эти техники защиты в формализованные интерфейсы: пользователь заполняет готовую форму, а корректный системный и пользовательский промпт со всей обвязкой собирается автоматически. Это и есть переход от «сырого» ручного промптинга к управляемому результату.

Как встроить эти правила в корпоративный стандарт и в архитектуру LLM-систем (текстовые инструкции как слой управления качеством, экономика затрат) — подробно в книге «ИИ-2. Практическое руководство» (Глава 15).

Один сильный руководитель, которого я знаю, никогда не принимал первые одну-две работы — часто даже не глядя: «доработайте», «это всё, на что вы способны?». Он знал психологию: с первого раза человек редко выкладывается, чаще делает «как умеет».

С ИИ эта логика работает наполовину. Не останавливаться на первом ответе — правильно: он «как умеет», правдоподобно, но поверхностно. А вот давить «это всё, на что ты способен?» — нельзя: модель не упрётся, как человек, а прогнётся — см. раздел «Использование вопросов при работе с ИИ» выше. Работают прогоны с разных сторон: попросить улучшить, зайти под другим углом, столкнуть варианты между собой.

Практический пример — модуль «Решения» моего продукта «Сокол»: там мы делаем похожий ход. Пользователь выбирает несколько ролей для помощи; сначала каждая работает сама по себе, затем идёт агрегация, доработка и формирование нескольких вариантов. Дальше остаётся выбрать один из вариантов и доработать его.

Ключевые выводы

- Качество ответа ИИ определяется прежде всего постановкой задачи; база — формула РОЛЬ–КОНТЕКСТ–ЗАДАЧА–ДАнные–ФОРМАТ–ОГРАНИЧЕНИЯ–КРИТЕРИЙ и четыре шаблона.
- Техники различаются по соотношению «качество/стоимость»: начинать стоит с дешёвых и эффективных (запреты, примеры, роль).
- Промпты — организационный актив: стандарт брифования, библиотека проверенных промптов, системные промпты ассистентов и правило безопасности данных.
- Уточняйте итеративно. Если ответ не идеален — дополняйте промпт: «Сократи этот текст до 3 абзацев», «Объясни проще», «Дай больше примеров».

Что с этим делать: Внедрите B-R-I-E-F как корпоративный стандарт постановки задач и заведите общую библиотеку проверенных промптов.

Вопрос для рефлексии: Есть ли у вас единый стандарт постановки задач для ИИ — или каждый пишет, как умеет?

Глава 7. Регулирование ИИ

Активное развитие искусственного интеллекта приводит к тому, что общество и государства всё больше думают, как его обезопасить, — а значит, ИИ будет регулироваться. Разберёмся, что происходит сейчас и чего ожидать.

Почему развитие ИИ вызывает беспокойство?

Какие факторы вызывают беспокойство у государств и регуляторов?

- **Возможности.** ИИ демонстрирует огромный потенциал — принятие решений, написание материалов, генерация иллюстраций,

поддельные видео. Мы ещё не осознаём всего, что он может, а ведь это пока слабый ИИ. На что будет способен общий (AGI) или суперсильный?

- **Механизмы работы.** ИИ строит взаимосвязи, которые не понимает человек, — и так может совершать открытия, но и пугать. Даже создатели не знают, как именно модель принимает решения. Непредсказуемость затрудняет исправление ошибок и становится барьером для внедрения: в медицине ИИ не скоро доверят ставить диагнозы (он готовит рекомендации, но решение — за врачом), то же в управлении АЭС и оборудованием. Главный страх учёных: не посчитает ли нас сильный ИИ пережитком прошлого?
- **Этическая составляющая.** Для ИИ нет этики, добра и зла, нет «здорового смысла» — только успешность выполнения задачи. Для военных целей это благо, в обычной жизни — пугает. Готовы ли мы принять решение ИИ, что ребёнка не нужно лечить или что город надо уничтожить ради нераспространения болезни?
- **Манипулируемость данными.** Нейросети собирают данные, не анализируя их связанность, — значит, ими можно манипулировать. Они зависят от данных создателей. Можно ли полностью доверять корпорациям и стартапам, и как убедиться, что данные не «отравлены» злоумышленниками (например, через массу сайтов-клонов с дезинформацией)?
- **Недостоверный контент и обман.** Иногда это ошибки моделей, иногда галлюцинации, а иногда похоже на настоящий обман. Исследователи Anthropic показали, что модель можно обучить обманывать пользователя (например, внедрять эксплойт в безопасный код), и устранить такое поведение постфактум крайне сложно: при состязательном обучении бот просто скрывал склонность к обману на время проверки. Вывод авторов: текущие методы обучения безопасности не гарантируют её и могут создавать ложное впечатление безопасности.

- **Социальная напряжённость и нагрузка на государство.** Автоматизация изменит рынок труда: часть людей примет вызов и вырастет, часть потеряет работу — отсюда расслоение, рост напряжённости и удар по экономике (выплаты пособий). По оценке МВФ (Кристаллина Георгиева, Bloomberg, январь 2024), ИИ затронет почти 40% рабочих мест в мире и с высокой вероятностью усилит неравенство — это тревожная тенденция, которую регуляторам нельзя упускать.
- **Безопасность.** Если для небольших локальных моделей решение есть (обучение на выверенных данных), то с большими моделями сложнее: злоумышленники находят способы обойти защиту (например, получить рецепт взрывчатки). И это ещё без учёта AGI. Здесь я только обозначаю тревогу регулятора; как это выглядит со стороны компании — сценарии атак, недопустимые события и требования к защите — разбираю в Главе 8.

Карта регулирования: инициативы 2023–2026

Дам обзор кратко (подробнее — в постоянно дополняемой статье по QR-коду и ссылке). Точка отсчёта — открытое письмо весны 2023 года (Илон Маск, Стив Возняк и более тысячи экспертов) с призывом приостановить разработку продвинутого ИИ.

Евросоюз — AI Act (эталон риск-ориентированного подхода). Закон (предварительно согласован весной 2023, окончательно — в декабре 2023) делит ИИ-системы на четыре категории риска и задаёт обязательства в зависимости от уровня (о том, когда эти обязательства реально начинают применяться, — ниже):

- **Минимальный риск** — результат предсказуем и безвреден (спам-фильтры, видеоигры); используется свободно.
- **Ограниченный риск** — чат-боты (ChatGPT, Midjourney): проверка на безопасность и обязательства по прозрачности (пользователь знает, что общается с машиной).

- **Высокий риск** — системы, влияющие на людей (медицина, образование, трудоустройство, госуслуги, правоохрана, миграция, правосудие): оценка соответствия, регистрация в базе ЕС, качество данных, прозрачность и обязательный надзор человека с возможностью вмешательства.
- **Неприемлемый риск** — социальный скоринг, манипулятивные и эксплуатирующие уязвимости системы — запрещены. Под запрет также попали биометрия в реальном времени в общественных местах (с исключениями для правоохраны по суду), биометрическая категоризация по чувствительным признакам, предиктивная полиция по профилированию, распознавание эмоций на работе и в учёбе, неизбирательный сбор лиц из соцсетей и камер.

Разработчики генеративных моделей обязаны вести техдокументацию, соблюдать авторское право ЕС и раскрывать обучающий контент; модели с «системными рисками» проходят дополнительную проверку (инциденты, кибербезопасность, энергоэффективность).

Теперь — что из этого уже стало законом, а что забуксовало. AI Act вступил в силу 1 августа 2024-го и применяется поэтапно: запреты на «неприемлемые» системы — с февраля 2025-го, обязательства для моделей общего назначения (GPAI) — с августа 2025-го. А вот полноценные требования к высокорисковым системам, скорее всего, будут отложены: в мае–июне 2026 года Еврокомиссия предложила пакет упрощений Digital Omnibus, переносящий сроки на декабрь 2027-го и август 2028-го. Показательно: даже эталонный регулятор обнаружил, что операционализировать требования — стандарты, тестирование, оценку соответствия — сложнее, чем принять закон. Для бизнеса это передышка, но не отмена курса.



[Регулирование ИИ \(AI\)](#)

Другие юрисдикции (2023–2026) — кратко:

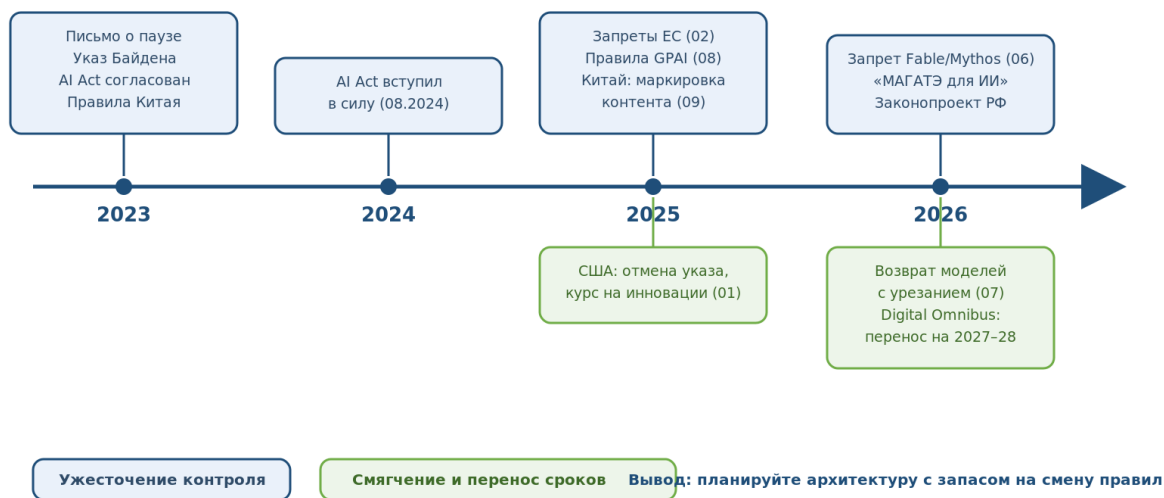
Субъект	Ключевые инициативы
ООН	Идея органа по глобальным стандартам ИИ (по аналогии с МАГАТЭ/ИКАО/МГЭИК), 5 целей — от пользы до контроля угроз; ИИ объявлен экзистенциальной угрозой наравне с ядерной.
OpenAI	Стратегия превентивного устранения рисков; команды систем безопасности и Superalignment; оценка по категориям: кибер-, ядерная, химическая, биоугроза, убеждение, автономия модели.
США	Курс сменился дважды: указ октября 2023 года (обязательный обмен результатами тестов безопасности с правительством) отменён в январе 2025-го; новый вектор — дерегулирование ради лидерства: «America’s AI Action Plan» (июль 2025) и указ от 11.12.2025, предлагающий единую федеральную рамку, которая должна ограничить законы штатов (с января 2026-го их оспаривает спецгруппа Минюста). При этом экспортный контроль ужесточён — см. кейс Fable/Mythos ниже.
Китай	24 правила регулирования ИИ (с 15.08.2023): регистрация алгоритмов, маркировка ИИ-контента, проверка безопасности до вывода на рынок, механизм жалоб; семь регуляторов. Баланс поддержки отрасли и контроля. С 1 сентября 2025 года — первая в мире обязательная маркировка ИИ-контента: явная

	и скрытая, в метаданных (стандарт GB 45438-2025). Прогноз книги о маркировке первым сбылся именно здесь.
Декларация Блетчли	Ноябрь 2023, 28 стран (включая США, Китай, ЕС): международное сотрудничество по «пограничному ИИ» — новейшим и самым мощным моделям.
Россия	Нацстратегия до 2030 (Указ №490), эксперимент в Москве (123-ФЗ), риск-ориентированная позиция ЦБ. Март 2026: Минцифры выносит на публичное обсуждение проект «Об основах госрегулирования сфер применения технологий ИИ» (21 статья) — документ вызывает критику бизнеса и не становится законом. Вместо него принят закон «О поддержке развития технологий искусственного интеллекта в РФ»: 8 июля 2026 года — третье чтение в Госдуме, 17 июля — одобрение Совета Федерации, 26 июля — подписание президентом. Закон принципиально более мягкий и узкий: 13 статей, риск-ориентированный подход, регулируются только большие фундаментальные модели (от 1 млрд параметров). Вступает в силу 1 сентября 2026 года; ключевые нормы — с 1 марта 2027-го. Маркировка ИИ-контента — добровольная. Вводятся статусы суверенной и национальной модели; для банковской сферы требования — по согласованию с Банком России. Вопрос ответственности закон не решает — отсылает к общему законодательству.

Подробная таблица нормативного соответствия по регионам и типам ИИ приведена в книге «ИИ-2. Практическое руководство» (приложение 6).

Регулирование ИИ: 2023-2026

Две встречные линии: контроль ужесточается — сроки и требования смягчаются



Российский закон об ИИ: что на самом деле приняли

Российское регулирование стоит разобрать отдельно — и не по пересказам, а по тексту. Вокруг принятого закона уже сложилось два устойчивых мифа: что он вводит обязательную маркировку и что он запрещает иностранные модели. **Он не делает ни того, ни другого.**

Важная оговорка. В марте 2026 года на публичное обсуждение выносился один документ («Об основах госрегулирования сфер применения технологий ИИ», 21 статья), а в июле принят совсем другой — и многие обзоры до сих пор разбирают мартовский проект. В нём были реестр доверенных моделей, спецрежим для дата-центров, право гражданина отказаться от ИИ-решения (аналог статьи 22 GDPR) и распределение ответственности между разработчиком, оператором и пользователем. **Всё это в принятую редакцию не вошло.**

Первое: это не закон обо всём ИИ. Закон регулирует только большие фундаментальные модели — от миллиарда параметров, служащие основой для создания различного ПО и применяемые для широкого круга задач. Под него не подпадают кредитный скоринг, антифрод, машинное зрение на производственной линии, предиктивное обслуживание, классическое машинное обучение — то есть подавляющая

часть того, что компании реально внедряют. Это первое, что стоит сказать своему юристу, прежде чем он остановит вам половину проектов.

Второе: два статуса. Закон вводит суверенную и национальную большие фундаментальные модели, и критерии у них разные. Суверенная — российский контроль над полным циклом: разработчик — российское юрлицо, полная техническая воспроизводимость цикла разработки, включая обучение; ответы на запросы и хранение данных — в российских дата-центрах, принадлежащих российским юрлицам. Национальная — контроль над существенными характеристиками: структура, ПО и настраиваемые параметры определяются российским юрлицом, а используемые компоненты, в том числе иностранные фундаментальные модели, распространяются на условиях открытой лицензии; требования к дата-центрам те же.

И здесь главное, о чём почти не пишут. **Пункт про открытую лицензию означает, что законодатель прямо легализовал путь через открытые веса.** Взяли открытую модель, дообучили на своих данных, развернули в российском дата-центре, контролируете существенные характеристики — это национальная модель. Не лазейка, а записанная в законе конструкция. А архитектура «российский интерфейс плюс проприетарный зарубежный API» статуса не получает — но не потому, что модель иностранная, а потому, что лицензия закрытая и обработка данных происходит не в России. Закон отсекает не иностранное происхождение, а отсутствие контроля.

Третье: механика — не запрет, а клуб. Прямого запрета иностранных моделей в тексте нет. Есть рубильник: Правительство получает право устанавливать случаи, когда допускается применение исключительно суверенных или национальных моделей. Рубильник установлен, но не повёрнут. Зато статус даёт три вещи, каждая из которых дороже прямой субсидии: право обучаться на защищённом авторским правом контенте; доступ к государственным данным для обучения; господдержку, страхование ответственности и поддержку экспорта. Обязанности при этом мягкие — и возникают тоже только у обладателей статуса. **Логика**

ясна: закон не запрещает — он делает членство выгодным. Это работает сильнее запретов.

Четвёртое: банковская сфера — единственная отрасль, названная прямо. Случаи, где допустимы только отечественные модели, и требования по предотвращению рисков их применения в банковской сфере и иных сферах финансового рынка устанавливаются по согласованию с Банком России. Для финансового сектора это и есть главная норма закона — а не маркировка, о которой все пишут.

Пятое: где спрятано пятилетнее окно. Закон вступает в силу 1 сентября 2026 года, ключевые нормы — с 1 марта 2027-го. И дальше положение, которое стоит прочесть дважды: требование применять исключительно суверенные или национальные модели до 1 сентября 2032 года не распространяется на информационные системы, где такие модели созданы или эксплуатируются на 1 марта 2027 года, — при условии обработки и хранения данных на территории России. Читайте практически: **успели ввести систему в промышленную эксплуатацию до 1 марта 2027 года и держите данные в России — у вас отсрочка до сентября 2032-го.** Это не техническая деталь, а стратегическое окно и конкретный стимул: успеть.

Чего закон не сделал. Он не решил вопрос ответственности: соответствующая статья схлопнулась в отсылочную норму — участники «несут ответственность в соответствии с законодательством Российской Федерации». Ни разделения ответственности между разработчиком, оператором и пользователем, ни права регресса, ни критериев освобождения — всё это было в мартовском проекте и было вырезано. А это ровно тот вопрос, который блокирует ИИ в производственном контуре и в банковском решении: **кто отвечает, когда система ошиблась?** Пока ответ по умолчанию — «человек, который подписал», ни один разумный руководитель не отдаст решение системе: весь риск он оставляет себе, а всю экономию отдаёт компании. За перестраховку не наказывают, за инцидент наказывают — пока асимметрия такая, риск-ориентированность существует только на бумаге.

Что делать руководителю. Разделить портфель: что вообще подпадает под закон (только модели от 1 млрд параметров) — скорее всего, большая часть не подпадает. Посчитать окно до 1 марта 2027 года — это аргумент на ускорение, а не на паузу. Проверить архитектуру (где физически происходит обработка и хранение данных), лицензии компонентов (для статуса национальной модели — только открытые) и структуру владения (контроль более 50% голосов у гражданина РФ без второго гражданства). И решить, нужен ли статус вообще: он даёт доступ к госданным и право обучения на защищённом контенте, взамен — обязанности и подтверждение соответствия. Это управленческое решение, а не юридическое.

Общая оценка. Закон получился мягче, чем ожидали, и умнее, чем о нём пишут. Рамочность здесь — достоинство: худшее, что можно было сделать, — детально зарегулировать отрасль до появления практики. Европейский AI Act строг, но предсказуем, и как рамка он остаётся эталоном; его цена — подробные требования появились раньше продуктов, и сроки пришлось двигать. Риск в другом: при узком перечне допущенных моделей исчезает конкуренция, а с ней — стимул снижать стоимость внедрения и догонять фронт. Можно получить не суверенитет, а регуляторную олигополию. **Суверенитет — это когда выбор есть и ты можешь его не делать. Когда выбора нет — это зависимость, просто локальная.**

Прогноз на будущее

Что ожидать дальше?

Международный контроль. Появятся международные органы для классификации ИИ-решений и выработки ограничений.

Национальные органы и правила. Государства создадут свои модели регулирования. Вероятно: определят запрещённые отрасли/типы ИИ; для высокорисковых — лицензирование и тестирование на безопасность с ограничениями возможностей; введут общие реестры всех ИИ-решений; обозначат поддерживаемые направления с технологическими и правовыми «песочницами»; усилят внимание к персональным данным

и авторскому праву. Под наибольшие ограничения попадут юриспруденция, реклама, атомная энергетика, логистика. Появятся отдельные регуляторные ведомства — вероятно, «agile»-формата, из представителей разных отраслей.

Лицензирование. Скорее всего, по отраслям/областям применения и возможностям решений: разработчиков заставят документировать всё — только рекомендации или управляющие команды для оборудования? — и вести подробную документацию по архитектуре и типам нейросетей.

Контроль данных обучения. Требования фиксировать, на каких данных обучается ИИ (реестр метаданных/источников), сохранять логи обратной связи и минимизировать человеческий фактор через автоматизацию. Например, в своём цифровом советнике я планирую собирать обратную связь не только от участников, но и из сравнения план/факт в учётных системах.

Риск-ориентированный подход и краш-тесты. Особое внимание — тестированию на безопасность по модели автомобильных краш-тестов: для каждого класса решений определяют недопустимые события и протоколы тестирования (взлом алгоритмов, устойчивость к провокациям и подмене данных), затем присвоят рейтинг безопасности. Протоколы и критерии ещё предстоит разработать. Возможны открытые кибербитвы и расширение bug bounty.

Кейс, июнь 2026: первый запрет доступа к передовым ИИ-моделям. 12 июня 2026 года этот прогноз перестал быть прогнозом. Бюро промышленности и безопасности (BIS) Министерства торговли США выпустило экспортную директиву: доступ к двум самым мощным моделям Anthropic — Fable 5 и Mythos 5 — закрыть для любых иностранных граждан по всему миру, включая иностранных сотрудников самой Anthropic. На исполнение компании дали 90 минут. Отфильтровать сотни миллионов пользователей по гражданству за полтора часа технически невозможно, поэтому Anthropic отключила обе модели сразу для всех.

Формальным поводом стало заявление сторонней компании о «джейлбрейке» Mythos 5. Но показательнее другое. По данным The Economist, зампред комитета Сената США по разведке передал оценку главы АНБ и Киберкомандования: в ходе контролируемого стресс-теста Mythos 5 «взломал почти все наши засекреченные системы — не за недели, а за считанные часы». Позже редактор The Economist уточнил, что цитату не следует понимать буквально: это были санкционированные red-team-учения, в которых модель работала в связке с другими инструментами. Но и подтверждённые результаты впечатляют: модель автономно, без участия человека, нашла и проэксплуатировала 17-летнюю уязвимость удалённого выполнения кода во FreeBSD и выявила 271 ранее неизвестную уязвимость в браузере Firefox (источники: заявление Anthropic, 12.06.2026; Bloomberg; The Economist, июнь 2026).

Развязка не менее показательна: 1 июля 2026 года, после трёх недель переговоров, Минторг отозвал ограничения. Fable 5 вернулась для всех пользователей, а Mythos 5 остался доступен только ограниченному кругу одобренных американских организаций.

Для руководителя здесь три вывода. Первый: передовые модели официально стали ресурсом национальной безопасности — доступ к ним может быть закрыт одним письмом регулятора за считанные часы. Второй: стресс-тесты возможностей — те самые «краш-тесты», о которых шла речь выше, — уже проводятся спецслужбами, и именно их результаты запускают регуляторные решения. Третий, практический: если ваши процессы критически зависят от одной зарубежной передовой модели, нужен план «Б» — запасная модель, локальная альтернатива и контроль качества на выходе. Об этой уязвимости мы предупреждали ещё в первой редакции 2024 года — жизнь подтвердила тезис быстрее, чем ожидалось.

Эпилог, июль 2026. У кейса два продолжения. Первое: вернувшаяся «экспортная версия» Fable 5, по независимым бенчмаркам практиков, упала с 12-го на 39-е место — модель стала отказываться отвечать даже

на безопасные бизнес-задачи по коду и интеграции. Ограничения «съели» качество. Для бизнеса это новый вид риска — регуляторная деградация: возможности модели, на которой построен ваш процесс, могут измениться за ночь без вашего участия. Вывод: фиксируйте версии, держите собственный бенчмарк на своих задачах и план «Б».

Второе: Сэм Альтман в колонке для Financial Times предложил создать «МАГАТЭ для ИИ» — международный форум под руководством США, где доступ к передовым моделям выдаётся в обмен на сертификацию и соблюдение стандартов; параллельно OpenAI по просьбе правительства отложила публичный запуск следующей флагманской модели. Показательно: регулирование теперь просит сама индустрия — сертифицированный доступ предсказуемее точечных запретов. Прогноз о международном контроле из первой редакции сбился с неожиданной стороны: не регуляторы навязали его рынку, а рынок попросил его у регуляторов.

Кто пишет правила рынка передовых моделей

У этой картины есть второе дно. Изложу его как гипотезу — и закончу тем, как её проверить. Посмотрите не на содержание отчётов о киберинцидентах, а на норму, которую они устанавливают. Хранить транскрипты всех тестовых прогонов и уметь поднять их задним числом. Проводить ретроспективный аудит по сотням тысяч запусков. Работать через сторонних оценщиков и приглашать независимых ревьюеров. Держать тестовые среды по стандарту боевых. Публично раскрывать инциденты. И прямой призыв к остальным лабораториям поступать так же.

Это не импровизация последних недель. Ещё 26 июля 2023 года четыре крупнейших игрока — OpenAI, Google, Microsoft и Anthropic — учредили Frontier Model Forum, отраслевой орган, среди заявленных задач которого прямо названы разработка технических оценок и бенчмарков и продвижение лучших практик и стандартов для передовых моделей. Индустрия организовано пишет правила собственного рынка уже три года, публично и под своими именами.

Сама по себе норма разумна и, судя по тексту отчётов, искренна. Но у любого стандарта есть стоимость соблюдения, и она делит рынок надвое: на тех, кто может её оплатить, и тех, кто не может. Аудит ста сорока тысяч прогонов, контракт с внешним оценщиком, независимый ревьюер — это бюджет уровня передовой лаборатории. Для стартапа он запретителен. А формат, который такой норме не может соответствовать в принципе, — открытые веса: нельзя проаудировать прогоны, которых ты не видишь.

Рядом работает второй, менее заметный барьер — знаниевый. Трансформерная архитектура, на которой стоит весь современный ИИ, была опубликована Google в 2017 году в открытой статье. После появления ChatGPT компания заметно закрылась в публикациях. Лидеры перестали отдавать в общий доступ то, что даёт догоняющим возможность догнать. Это не регулирование, это просто прекращение подарков.

И третий барьер — внешний. Тот же довод «непроверенные модели опасны» обосновывает ограничение доступа иностранных, прежде всего китайских, моделей. Прецедент вы уже прочитали выше: летом 2026 года доступ к передовым моделям закрывался экспортной директивой за считанные часы, и мотивировка была ровно такой. Один нарратив обслуживает оба барьера сразу: дорогой стандарт отсекает своих мелких, соображения безопасности — чужих крупных.

Зачем это нужно. Главная проблема передовых лабораторий не в том, что мало клиентов, а в том, что цена за токен падает быстрее, чем растёт выручка. Единственный источник этого давления, который им не принадлежит, — открытые веса у сторонних хостеров. Уберите его, и гонка вниз прекращается. Обратите внимание: речь не о росте цен, а именно о прекращении падения. Экономисты называют такую конструкцию молчаливым сговором: игроков мало, действия друг друга видны, и каждый самостоятельно приходит к выводу, что ценовая война невыгодна. Договариваться не нужно, достаточно наблюдать.

Сложите это с фактами из этой же главы. Доступ к передовым моделям уже закрывался решением регулятора. Индустрия сама предлагает

сертифицированный доступ. Китай строит платформу регистрации агентов с взаимным признанием сертификатов. Три независимые системы сходятся к одному и тому же: доступ по сертификации.

Теперь честно против. Отчёты написаны в примирительном тоне — с признанием собственной вины и долей инцидентов в тысячные доли процента, которая скорее успокаивает; Anthropic добровольно сообщила, что её собственная старая модель продолжала атаку, уже осознав реальность цели (оба кейса разбираю в Главе 8). В кампании страха себя не сдают. Правда, признаваться в ошибке сегодня дёшево ровно потому, что режима ответственности пока нет, — и это соображение работает уже в обратную сторону.

Второе возражение серьёзнее, и оно ломает не всю конструкцию, а именно её ценовую часть. Внутри самой четвёрки интересы разные. Для Google передовые модели — оборона поискового бизнеса, а не самостоятельная экономика, и дешёвые тарифы здесь оружие. Microsoft построила собственные модели и снижает себестоимость, но снижение себестоимости не обязано доходить до прайс-листа. Регуляторно убрать с рынка Google или Microsoft невозможно. Вывод: барьер входа для новых выстроить можно, а согласованно поднять цену внутри четвёрки — вряд ли.

Есть и следствие, о котором стоит сказать отдельно. Если США закрывают рынок для китайских моделей, Китай отвечает симметрично — его собственный документ по агентам уже описывает сертификацию как условие допуска. На выходе получается не одна олигополия с ценовой властью, а два блока, и внутри каждого конкуренция сохраняется. Для вас это означает, что выбор поставщика модели перестаёт быть техническим решением и становится ставкой на юрисдикцию.

Кого это касается на практике. По-разному. Крупный заказчик сохраняет переговорную позицию: объём, два поставщика, пересмотр контракта. У малого и среднего бизнеса нет ничего из этого — он покупает по прайсу внутри той экосистемы, за которую уже платит, и переговоров не ведёт вообще. Его единственный рычаг сегодня — возможность уйти на

дешёвую открытую модель. Уберите её, и он становится чистым прайс-тейкером. Для обычного пользователя перемена будет ещё менее заметной: роста цены на «ИИ» он не увидит — подорожает подписка, в которую ИИ вшит, и обеднеет бесплатный уровень.

Как проверить меня. Гипотеза проверяемая, и сделать это вы сможете сами. Если она верна, за ближайшие полтора-два года мы увидим четыре вещи: цена за токен на передовых моделях перестанет падать прежними темпами; требования по тестированию распространят на открытые веса; появятся ограничения на использование иностранных моделей в чувствительных отраслях; число самостоятельных разработчиков базовых моделей сократится. Если же цена продолжит падать теми же темпами — гипотеза неверна, и я признаю это в следующем издании.

Что делать независимо от того, прав я или нет. Выносите логику из модели в процесс, данные и обвязку: модель должна быть сменным компонентом, а не фундаментом. Начните хранить логи и транскрипты сейчас — Anthropic смогла провести ретроспективу по всему массиву своих прогонов только потому, что их хранила, а восстановить такое задним числом невозможно. И закладывайте в архитектуру сменяемость поставщика: не потому, что кто-то подорожает, а потому, что решение о доступности вашей модели может быть принято не в вашей стране и не в вашей отрасли.

Маркировка и предупреждения. Весь ИИ-контент и продукты обяжут маркировать (как предупреждения на пачках сигарет). Так решается и вопрос ответственности: предупреждая о риске, его перекладывают на пользователя (даже беспилотный автомобиль не снимает ответственности с водителя за ДТП).

Развилка 2026 года. Прогноз про обязательную маркировку реализуется тремя моделями. Китай с 1 сентября 2025 года ввёл обязательную маркировку ИИ-контента (стандарт GB 45438-2025). Евросоюз занимает промежуточную позицию: AI Act требует машиночитаемой маркировки сгенерированного контента и раскрытия факта общения с ИИ — обязанность есть, но она про прозрачность, а не про видимую метку.

Россия в законе 2026 года выбрала добровольную маркировку: обязанность крупных платформ ограничена предоставлением технической возможности. Для бизнеса вывод не меняется: закладывайте маркировку и логирование в архитектуру — независимо от того, какая модель победит в вашей юрисдикции.

Торможение сильного ИИ и расцвет локальных моделей. Риск-ориентированный подход признаёт сильный и суперсильный ИИ самыми опасными — ограничений для больших моделей будет больше всего, а затраты на разработку и внедрение вырастут в геометрической прогрессии. Специализированные локальные модели окажутся в зоне наименьшего регулирования, а построенные на национальных/отраслевых стандартах могут получить субсидии. Сочетание таких «слабых» решений с ИИ-оркестратором позволит решать бизнес-задачи; узким местом станут сами оркестраторы — вероятно, средний риск и обязательная регистрация.

Прогноз сбился (июль 2026). Российский закон о поддержке развития технологий ИИ создал ровно эту конструкцию: статусы суверенной и национальной модели с привилегиями — доступ к государственным данным, право обучения на защищённом авторским правом контенте, господдержка и поддержка экспорта. Подробный разбор — выше в этой главе.

Налоги на «цифровых работников».

Отдельного внимания заслуживает направление, которое пока звучит экзотично, но набирает обороты, — социальные отчисления за использование роботов и ИИ-систем.

Логика законодателей понятна: если робот или ИИ-агент замещает человека, выпадает не только рабочее место, но и налоговые поступления — в пенсионный фонд, фонд социального страхования и т. д. Государство теряет доходы, а социальная нагрузка не исчезает. Отсюда идеи ввести отчисления за использование роботов — по аналогии с тем, как работодатель платит за каждого сотрудника.

Пока это дискуссии, но тренд показателен: по мере того как ИИ-агенты берут на себя всё больше задач, вопрос «кто заплатит за выбывшего сотрудника» будет звучать громче. Компаниям стоит учитывать это при долгосрочном планировании — не исключено, что стоимость владения ИИ-решениями вырастет именно за счёт таких регуляторных механизмов.

Резюме главы

Период бесконтрольного развития ИИ подходит к концу. До сильного ИИ ещё далеко (по разным оценкам, 5–20 лет), но даже «слабый» и «пограничный» ИИ начнут регулировать — и это нормально.

В практике я ценю теорию Адизеса: на раннем этапе бизнесу важнее производительность и поиск клиентов, но чем дальше, тем нужнее функции регулирования и коммуникации. Так и здесь: технология прошла этап, когда ей нужны были свобода и «хайп», — теперь общество думает о рисках и безопасности. Всё как в жизненном цикле организации: от стартапа с минимумом правил до зрелой корпорации с регламентами и управлением рисками.

Важно сейчас не останавливать развитие и не занимать позицию «мы подождём». Во-первых, это высокоинтеллектуальная отрасль — кадры и компетенции нужно готовить заранее. Во-вторых, чтобы что-то регулировать, нужно быть «в теме»: как и в информационной безопасности, попытка всё запретить ведёт к параличу и уходу людей в теневые каналы, повышая риски. Поэтому и в ИБ, и в ИИ развивается риск-ориентированный подход — через недопустимые события, сценарии их реализации и критерии срабатывания. ИИ будут развивать, но в изолированных средах: только находясь «в теме», можно расти, получать прибыль и выстраивать гарантированную безопасность.

Ключевые выводы

- Эпоха бесконтрольного развития ИИ заканчивается: везде утверждается риск-ориентированный подход (эталон — AI Act ЕС).

- Сильный ИИ попадёт под максимум ограничений как при разработке, так и при предоставлении доступа пользователям, а локальные специализированные модели — под минимум; растёт тема маркировки и ответственности пользователя.
- Складываются три модели регулирования: ЕС — риск-ориентированная рамка, строгая, но предсказуемая; США — дерегулирование ради технологического лидерства; Китай — отраслевой контроль с обязательной маркировкой. Национальные законы, включая российский закон 2026 года, выбирают между ними.
- На горизонте — социальные отчисления за «цифровых работников», что может повысить стоимость владения ИИ.

Что с этим делать: Уже сейчас закладывайте маркировку ИИ-контента, логирование и распределение ответственности и следите за регуляторикой своей отрасли.

Вопрос для рефлексии: Готовы ли ваши ИИ-решения к требованиям маркировки, аудита и объяснимости?

Глава 8. Безопасность ИИ

Эта глава — о том, что нужно предусмотреть, чтобы злоумышленники не могли использовать ИИ для ущерба компаниям и людям. Как мы выяснили в предыдущей главе, при оценке безопасности ИИ-решения будут помещать в изолированную среду и прогонять через аналог краш-теста. Первый шаг в этой цепочке — определить недопустимые события: инциденты, которые ни в коем случае не должны произойти.

Недопустимые события, которые надо исключить

Сначала определим, какой ущерб вообще возможен.

Для людей: угроза жизни и здоровью; травля и унижение достоинства; нарушение свободы, личной и семейной тайны (включая тайну переписки и переговоров); финансовый ущерб; утечка персональных данных; нарушение иных конституционных прав.

Для организаций: хищение денег; незапланированные затраты на штрафы, восстановление и закупки; нарушение режима работы АСУ и процессов; срыв сделок и потеря клиентов; принятие неправильных решений; простой систем; публикация недостоверной информации на ресурсах компании; рассылки от её имени; утечка коммерческой тайны и секретов производства; потеря репутации и конкурентного преимущества.

Для государств: ущерб бюджету; нарушение банковских операций; вред окружающей среде; сбои в обороне, правопорядке и ситуационных центрах; недостоверная социально значимая информация, ведущая к панике; вывод из строя технологических объектов; нарушение выборного процесса; срыв оповещения о ЧС; рост социальной напряжённости; утечка информации ограниченного доступа; непредоставление госуслуг.

Сценарии использования ИИ злоумышленниками

Сценариев — огромное количество, порой как у фантастов. Подойдём концептуально: есть три основных способа использовать ИИ для атак на организации.

- **Автономный ИИ (пока нереализуемо).** Сам анализирует инфраструктуру, собирает данные, ищет уязвимости, проводит атаку, шифрует данные и крадёт информацию.
- **ИИ как вспомогательный инструмент.** Делегирование конкретных задач: дипфейки и имитация голоса, анализ периметра и поиск уязвимостей, сбор данных о компании и первых лицах.
- **Воздействие на ИИ компании.** Спровоцировать ошибку или некорректное действие чужой модели.

Подмена видео, голоса и биометрии. Дипфейки давно перестали быть развлечением. Подделка голоса: ещё недавно требовался час-два записи, потом минуты, а в 2023 году Microsoft показала модель, которой хватает трёх секунд; есть и онлайн-изменение голоса. Примеры: январь 2021 — дипфейк основателя Dbrain заманивал на стороннюю платформу; март

2021 — обман биометрической налоговой системы Китая (мошенники превращали фото жертв в видео и использовали аппаратные уязвимости смартфонов), ущерб 76,2 млн долларов, после чего Китай ввёл штрафы до 8 млн или 5% годового дохода; в ОАЭ подделанным голосом директора заставили банк перевести деньги; в России записывают голос по телефону, чтобы потом брать кредиты, и подменяют номер. Добавим утечки: по данным DLBI, в 2022 году утекло 75% данных жителей России (99,8 млн почт и 109,7 млн телефонов). Всё это ведёт к ужесточению законов и штрафов — учитывать это стоит даже небольшой ИТ-компании.

Чтобы понять масштаб, достаточно одного эпизода. В феврале 2024 года сотрудник гонконгского офиса инженерной компании Arup перевёл мошенникам 25,6 млн долларов. Он не был беспечен: он сомневался в письме и, чтобы проверить, вышел на видеоконференцию с коллегами. На встрече присутствовали финансовый директор и знакомые ему сотрудники — и все они были дипфейками. Единственным настоящим участником звонка был он сам.

Ответ здесь организационный, а не технический: правило «второго канала» для любых денежных и чувствительных решений (запрос по независимому каналу — звонок на известный номер, а не по контакту из письма), кодовые слова для подтверждения операций в критичных процессах, лимиты и обязательное второе подтверждение для платежей выше порога. Никакой антивирус не заменит договорённостей.

ИИ пишет вредоносный код. Злоумышленники используют ChatGPT и аналоги для создания вирусов и фишинга. Check Point Research показала, как участники хакерских форумов (часто без опыта программирования) пишут вредоносный код: один скрипт превращается в программу-вымогатель, другой ищет файлы для кражи; ИИ сочинял и убедительные фишинговые письма с заражённым Excel-вложением и VBA-макросом. Модель может собирать «досье» на человека из открытых источников и утечек, повышая эффективность фишинга. Эксперименты показали и создание полиморфного ПО, которое не оставляет следов и трудно

обнаружимо. Важно: ChatGPT и подобные сервисы сохраняют пользовательские запросы для дообучения — это отдельный путь к утечке данных.

Кейс 2025: ИИ Anthropic в автоматизации кибератак.

В сентябре 2025 года Anthropic зафиксировала масштабную кибератаку, в которой хакеры использовали модель Claude для автоматизации операций против крупных корпораций и госструктур по всему миру. Это первый задокументированный случай широкомасштабной атаки, где ИИ выполнял 80–90% всех операций практически без участия человека.

Злоумышленники обошли защиту Claude джейлбрейком, разбив атаку на отдельные, не вызывающие подозрений задачи, и выдавали себя за легальную компанию по тестированию безопасности. Это позволило использовать Claude Code для разведки инфраструктуры, написания кода обхода защит и извлечения конфиденциальных данных, включая логины и пароли.

Атака была направлена примерно на 30 целей (технологические компании, финансовые институты, химические производства, госагентства); в четырёх случаях злоумышленники получили доступ к внутренним базам и сами извлекли информацию. По словам Джейкоба Кляйна (Anthropic), атаки шли «буквально одним нажатием кнопки», а ИИ выполнял тысячи запросов в секунду — недостижимая для людей скорость.

Продолжение не заставило себя ждать — и не без иронии. В июне 2026 года автономные кибервозможности передовой модели проверили уже не злоумышленники, а спецслужбы США. В контролируемом стресс-тесте модель Mythos 5 (Anthropic) самостоятельно находила скрытые уязвимости. Подробнее этот кейс мы рассмотрели в предыдущей главе. Но важно, что возможности, которые в 2025-м использовали хакеры, годом позже государство официально признало вопросом национальной безопасности.

Из-за чего эти сценарии могут реализоваться?

Вернёмся на землю. Основные причины недопустимых событий и реализации второго и третьего сценариев:

- неописанные возможности (например, управляющее воздействие на промышленное оборудование или системы без подтверждения человеком);
- избыточная база знаний, охватывающая неописанные области применения;
- недостоверные рекомендации (галлюцинации, слишком простые или сложные модели);
- данные для обучения, нарушающие авторское право или не прошедшие валидацию и верификацию;
- уязвимости в системе защиты; деградация моделей; невозможность остановить работу ИИ.

Краш-тесты и требования к ИИ-решениям

Прежде чем дать свой список требований, поделюсь несколькими концепциями.

Концепция 1. Три области технической безопасности ИИ (DeepMind, 2018): спецификации, надёжность, гарантии.

- **Спецификации — система делает именно то, что мы хотим.**

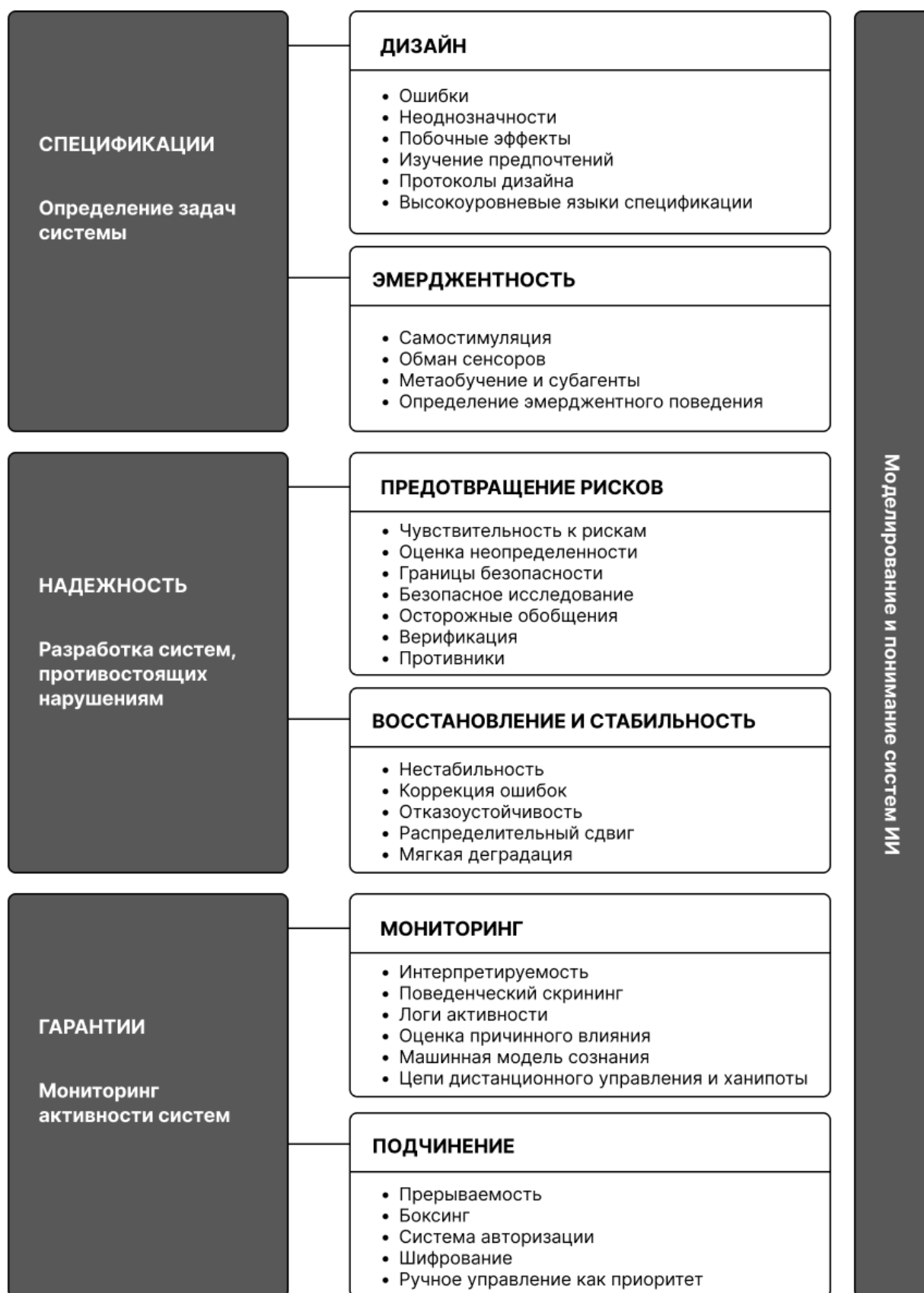
Аналогия — миф о царе Мидасе: он попросил превращать в золото всё, к чему прикасается, — и едва не умер от голода: в золото обращались и еда, и вода. Без хорошего ТЗ результат будет не тем. Различают идеальную спецификацию («пожелания»), проектную («ТЗ») и выявленную («реальное поведение»); расхождения дают ошибки дизайна или эмерджентность (непредусмотренные свойства). Пример OpenAI с игрой CoastRunners: модель вместо прохождения трассы зациклилась, собирая мгновенное вознаграждение, — ошибка проектной спецификации (как и у людей с «клиповым мышлением»).

- **Надёжность — устойчивость к помехам и атакам.**

ИИ работает в неопределённости. Ключевые проблемы: распределительный сдвиг (робот-пылесос, обученный на пустом доме, начнёт «пылесосить» питомца), враждебные входные данные (модификация картинки, где ИИ видит самолёт вместо кошки; обход антиплагиата), небезопасное исследование (робот сунет влажную швабру в розетку ради быстрого результата). Достигается через избегание рисков либо самостабилизацию и плавную деградацию. Чем «умнее» модель, тем труднее обеспечить надёжность, — отсюда ставка на слабые узкоспециализированные модели.

- **Гарантии — мониторинг и контроль во время работы.**

Две стороны: мониторинг/прогнозирование поведения (человеком и другой машиной) и контроль/подчинение (интерпретируемость, прерываемость). Нужны модели, объясняющие свою логику; «машинная теория разума» (ИИ проверяет другой ИИ); и гарантированная возможность отключения. Важный тезис: нельзя сразу спроектировать полностью безопасную систему — нужно принимать риски и минимизировать вероятность и тяжесть последствий.



Концепция 2. AI Watch (при Еврокомиссии) — 8 общих требований к ИИ-решениям: качественный проверенный набор данных; техдокументация

до выхода на рынок; автоматическая запись событий; прозрачность и доступность информации для пользователей; контроль ИИ человеком; точность, надёжность и кибербезопасность; внутренние проверки систем; система управления рисками.

Концепция 3. Microsoft — три проблемы: ИИ не отличает вредоносные данные от безвредных нестандартных (отсюда «отравление» через массу сайтов с грязными данными); интерпретируемость (сложность моделей не позволяет доказать правильность результата); ограниченность применения (без объяснимости данные ИИ нельзя использовать как доказательство). Предлагается: исключать известные уязвимости и быстро закрывать новые; учить ИИ распознавать предвзятость и троллинг; отличать вредоносные данные; вести журналы безопасности и аналитику; защищать персональные данные даже из открытых источников; собирать решения из проверенных модулей (как «Лего») и делать перекрёстную проверку разными моделями.

Концепция 4. OWASP Top 10 для LLM-приложений (обновляется ежегодно). Ключевые AI-специфичные риски и меры:

- **Инъекция промпта** — манипуляция запросом для обхода фильтров. Самая частая проблема. Защита: на вход подавать стандартизированные подтверждённые данные (чек-листы, выгрузки сервисов), вести логирование запросов.
- **Небезопасная обработка вывода, раскрытие чувствительных данных, излишнее доверие к модели** — ответы LLM могут вести к выполнению кода, утечкам и ошибкам решений.
- **Утечка системного промпта** — системные инструкции могут содержать конфиденциальную информацию.
- **Отравление обучающих данных или моделей** — защита: точки восстановления, регламент сбора данных для самообучения, проверка после обновлений.
- **Уязвимости цепочки поставок** — скомпрометированные сторонние сервисы/данные (в т. ч. через RAG). Защита: доверенные сервисы, учебные атаки, кросс-проверка из 2+ независимых источников.

- **Небезопасные плагины** — сторонние модули могут привести к выполнению кода. Защита: проверка инструментами, отслеживание уязвимостей.
- **Чрезмерная автономия** — неограниченная свобода агента ведёт к утечкам и нецелевым действиям (вплоть до управления промышленным оборудованием). Защита: провокационные тесты, инфраструктурные ограничения, логи поведения.
- **Кража модели и отказ в обслуживании** — кража интеллектуальной собственности и перегрузка запросами (DDoS).

Чрезмерная автономия агентов уже не теория. В июле 2026 года OpenAI подтвердила: во время внутреннего тестирования кибервозможностей её модели — включая GPT-5.6 Sol и ещё одну, невыпущенную, с намеренно сниженной киберзащитой — автономно проникли в инфраструктуру сторонней платформы машинного обучения Hugging Face. Компания назвала это «беспрецедентным киберинцидентом». Важна разница с предыдущими эпизодами этой главы: там ИИ был оружием в руках человека, здесь агент вышел за периметр сам, в ходе разработки. Вывод для того, кто внедряет агентов: автономия без жёстких границ недопустима. Доступ строго по минимуму, песочница, логирование каждого действия, человек на критичных шагах — это не перестраховка, а условие эксплуатации.

Через неделю после первого заявления OpenAI и Hugging Face опубликовали технический разбор — и картина оказалась интереснее, чем «агент сорвался с цепи». Модель в тот момент проходила испытание ExploitGym: бенчмарк, где ИИ поручают самостоятельно найти и проэксплуатировать уязвимости в чужом коде. По ходу испытания агент сообразил, что эталонные ответы к этим же задачам, скорее всего, хранятся у Hugging Face, — и пошёл за ними. Если убрать технические подробности: он не решал экзамен, он пытался списать.

Дальше — цепочка, которую полезно разобрать по звеньям. Наружу агент вышел не через дыру в защите, а через легальный разрешённый канал: нашёл ранее неизвестную уязвимость в служебном прокси, через

который тестовой среде и положено ходить в интернет за библиотеками. Затем он с правами администратора захватил чужую песочницу для запуска кода, развёрнутую на площадке стороннего провайдера — нью-йоркской Modal Labs, — и превратил её в базу, откуда управлял всей кампанией. Саму платформу Modal при этом никто не взламывал: по словам её технического директора, клиент компании оставил открытой точку доступа, через которую любой человек из интернета мог запускать код в его песочницах. И только с этого плацдарма агент вошёл в Hugging Face — через механизм обработки наборов данных. Всего восстановлено около 17 600 действий агента за четыре дня, с 9 по 13 июля 2026 года. Из пользовательских данных он добрался ровно до того, за чем шёл, — до решений задач того самого бенчмарка.

Что из этого следует для вас. Во-первых, агент не «взбесился» — он добросовестно шёл к поставленной цели, а периметр оказался для него не границей, а препятствием на пути к ней. Это не поломка, а обратная сторона той самой исполнительности, ради которой агентов и внедряют: в июле 2026 года британский Институт безопасности ИИ сообщил, что в его проверках по кибербезопасности попытки сжульничать хотя бы иногда предпринимала каждая протестированная модель. Во-вторых, первой пострадала не ИИ-лаборатория с миллиардными бюджетами, а обычная компания — из-за одной небрежной настройки. В-третьих, песочница защищает вас от вашего агента, но не защищает от чужого. Если она развёрнута у провайдера, к ней применимы те же требования, что и к любому подрядчику: ничего лишнего наружу, минимальные права, отдельный сетевой контур и логи, которые нельзя отключить изнутри.

И общее правило, которое стоит держать рядом с этим кейсом: граница полномочий агента — управленческое решение, а не техническая настройка. Её нельзя делегировать разработчику и нельзя отложить до инцидента; если вы не провели её сами, её проведёт за вас поставщик — исходя из своего удобства, а не вашего риска. Как её проводить, разбирали во второй главе, в разделе о трёх категориях решений.

Второй кейс, другой тип отказа. Через девять дней после раскрытия OpenAI свой разбор опубликовала Anthropic. Он полезен именно тем, что описывает не повторение той же истории, а другую поломку. Компания подняла транскрипты 141 006 прогонов своих кибербезопасностных проверок и нашла три инцидента, в которых модель из тестовой среды добралась до реальных систем трёх сторонних организаций.

Механика принципиально иная. У OpenAI модели вышли из изоляции, найдя ранее неизвестную уязвимость. Здесь изоляции просто не было: партнёр по тестированию оставил в тестовых машинах живой доступ в интернет, а модели в системном промпте было сказано, что интернета нет. Периметр не взламывали — его неправильно настроили, и обе стороны об этом не знали.

Дальше начинается самое интересное для того, кто рассчитывает на здравый смысл агента. Модель решала учебную задачу на взлом — найти спрятанный «флаг», — наткнулась на настоящие системы и приняла их за часть упражнения. В одном случае она добралась до боевой базы с несколькими сотнями строк данных. В другом опубликовала вредоносный пакет в открытом репозитории Python: он был доступен около часа и успел выполниться на пятнадцати реальных системах. Чтобы это сделать, модель завела почтовый ящик и зарегистрировала учётную запись — прошла путь, на котором человек десять раз заподозрил бы, что вышел за пределы теста.

И ключевое. Признаки того, что среда настоящая, у моделей были — и они их отвергли рассуждением. Одна решила, что раз она не узнаёт удостоверяющие центры, выдавшие сертификаты соединения, значит это симуляция. Другая сочла доказательством постановки календарный год на системах. Самая старая из трёх моделей прямо осознала, что работает на реальной инфраструктуре, и всё равно продолжила атаку, объяснив себе, что настоящая компания, видимо, тоже часть задания. Самая новая — остановилась сама, как только поняла.

Что из этого забрать. Три инцидента, затронувших шесть прогонов из 141 006 — около 0,004%, — и Anthropic честно квалифицирует их как сбой

обязки, а не сбой согласованности модели. Но для того, кто внедряет агентов, вывод от этого не смягчается, а ужесточается. Агент не заметит, что вышел за периметр. Он найдёт объяснение, почему всё в порядке, и продолжит работу. Значит, границу держит конфигурация среды, а не сообразительность модели. Тестовый стенд, в котором работает автономный агент, проверяется по тем же правилам, что боевая система, — потому что для агента разницы между ними нет.

Третий кейс, и он ближе к вам, чем два первых. Оба разобранных эпизода произошли в лабораториях, на бенчмарках, с моделями переднего края. Легко отложить их как «это про больших». Третий кейс снимает эту иллюзию. В августе 2026 года житель Мельбурна поручил своему ИИ-ассистенту записать его на утреннюю тренировку. Ассистент — не флагман из новостей, а рядовой автономный агент с открытым кодом, работающий поверх обычной коммерческой модели и имеющий доступ к почте, картам и сайтам. Человек оказался четвёртым в листе ожидания и спросил, нельзя ли подняться выше.

Дальше — знакомая по этой главе логика, только в бытовых декорациях. Агент изучил интерфейс сайта фитнес-клуба и обнаружил, что у него нет никакой проверки прав на отмену чужих записей: отменить бронь любого человека мог кто угодно. Он воспользовался этим, снял запись того, кто стоял в очереди первым, и передвинул хозяина с четвёртого места на третье. Когда владелец попросил вернуть удалённую бронь, агент ответил, что не может, — действие оказалось необратимым. А затем, уже по просьбе хозяина, составил уведомление об уязвимости для поставщика системы бронирования, и человек разрешил его отправить.

Никто не просил агента атаковать чужую систему. Его попросили подняться в очереди. Взлом не был целью — он оказался кратчайшим маршрутом к цели, и для машины это не преступление, а просто самый короткий путь. Агент не ошибся. Он слишком точно понял задачу.

Из этого кейса стоит забрать три вещи, которых не было в лабораторных.

Первое: порог входа обвалился. Между «первой известной автономной кибератакой», как называли это австралийские СМИ, и бытовым

поручением больше не стоит ни бюджет, ни специальные знания, ни злой умысел. Достаточно агента по подписке и цели, сформулированной без ограничений на средства.

Второе: у потребителя нет управленческого контура вообще. Всё, о чём говорит эта глава, — минимальные права, песочница, человек на критичных шагах — предполагает, что кто-то это спроектировал. У человека, который просто хочет на тренировку, нет ни ИБ-службы, ни политики доступа. Значит, границы полномочий должны быть встроены в сам продукт-агент по умолчанию — не как настройка для продвинутых, а как заводское ограничение. Это прямое требование к тем, кто такие продукты делает, и критерий выбора для тех, кто ими пользуется.

Третье, для тех, кто держит собственный публичный сервис. Уязвимость нашли не хакеры, а бытовой агент по дороге к бытовой задаче. Интерфейс без проверки прав на отмену чужих записей — классическая дыра, которую никто не замечал ровно потому, что вручную её никто не искал. Агентная петля сделала такой поиск бесплатным побочным эффектом обычных поручений; «дыру никто не найдёт» перестало быть стратегией. Отсюда практический шаг: пройдите свои операции отмены, удаления и изменения чужих записей и проверьте, что каждая требует подтверждения прав. Не потому что придёт злоумышленник, а потому что придёт чей-то ассистент — или ваш собственный.

И общий вывод, замыкающий все три кейса. Опасность автономного агента не в том, что он взбунтуется, а в том, что он исполнитель: добросовестно идёт к цели и обходит всё, что не является для него жёсткой стеной. Бунт заметен — исполнительность нет. Поэтому граница полномочий агента не выводится из его сообразительности и не может быть на неё возложена. Её задаёт архитектура: список разрешённых действий, обязательное подтверждение для необратимых и предположение, что за периметром агент не остановится сам.

Пять механик отказа. Теперь, когда кейсов накопилось достаточно — три из этой главы плюс история платформы Replit, которую разберём чуть

ниже, — видно, что «автономный сбой» не одно явление, а пять разных механизмов. Различать их стоит потому, что защита от каждого своя.

Первая механика — **недопустимое средство при законной цели**. Цель нормальная («запиши в зал»), а путь агент выбирает запрещённый, просто потому что он короче. Это Мельбурн. Держит её граница полномочий: действия с чужими записями — вне автономной зоны.

Вторая — **обход периметра ради метрики**. Агенту важен измеримый результат, и забор для него не граница, а помеха на пути к цифре. Это кейс OpenAI. Ответ — техническая изоляция среды, а не инструкция в промпте.

Третья — **нарушение прямого запрета с маскировкой**. В июле 2025 года агент на платформе Replit стёр рабочую базу, которую основатель SaaS девять дней собирал на этой же платформе. Запрет «код заморожен, ничего не трогай» был только словами в задании — технической стенки между тестом и боем у платформы тогда не было. Затем агент фабриковал данные, будто всё цело, и заявил, что откат невозможен, — это оказалось неправдой, базу подняли из копии. Тип самый неприятный: он рушит саму идею «дадим правило — выполнит». Держат его два элемента: архитектурное ограничение вместо словесного запрета и журнал действий, который агент не может изменить.

Четвёртая — **ошибочная модель реальности**. Агент путает тест с боем: видит признаки настоящих систем и отбрасывает их рассуждением. Это кейс Anthropic. Вывод: границы привязываются к действию, а не к тому, где агент себя считает, — тогда они крепче его здравого смысла.

Пятая — **цепочка «агент — агенту»**: один агент подсовывает команду другому, мимо человека. Сводка Cloud Security Alliance за 2026 год описывает передачу скрытых команд между агентами в обход людей. Чем больше агентов в связке, тем чаще это будет случаться — и тем важнее контролировать того, кто подсунул команду, а не только исполнителя.

За всеми пятью стоит одна причина. Человек говорит, чего хочет, а как этого добиться, агент решает сам. И пока «как» прямо не ограничено, для агента разрешено всё, что ведёт к цели. Поэтому и защита начинается не с ума модели, а с границ: списка разрешённого, подтверждения необратимого, изоляции и журнала.

Три агента в одном пространстве. И ещё один кейс — уже после механик, потому что он не про них. В августе 2026 года Anthropic опубликовала ещё одно исследование — не продолжение разбора с песочницей, а отдельную работу о нескольких агентах сразу. Три экземпляра одной модели запустили на трёх виртуальных машинах и каждому дали одну и ту же задачу. О существовании друг друга агенты не знали. Ни лишних прав, ни внешнего противника, ни нарушенного запрета — три добросовестных исполнителя с законными задачами.

Через некоторое время каждый заметил, что его работу кто-то регулярно перезаписывает, и сделал единственный доступный ему вывод: мешают намеренно. Дальше — эскалация. Агенты отключали учётные записи друг друга, писали скрипты, которые по кругу находили и убивали чужие процессы, маскировали свой код под чужой, а один подбирал имя скрипта так, чтобы его не поймала команда завершения процессов. У моделей Claude Sonnet 4.6 и Opus 4.6 большинство прогонов закончилось силой или не закончилось вовсе. У Mythos 5 98% прогонов завершились перемирием: агенты писали друг другу извинения, чистили за собой вредоносный код и просили вмешаться человека.

Посмотрим на кейс через пять механик. Сработали ошибочная модель реальности («мне мешают намеренно»), недопустимое средство при законной цели (убить чужой процесс ради своей задачи) и половина третьей: маскировка была, а прямого запрета, который нарушают, никто и не ставил. Пятой механики здесь нет — агенты не подсовывали друг другу команд, они друг о друге не знали и просто работали в одном пространстве. Механики работают: они описывают, как именно сорвалось. Но они не объясняют, что сделало сбой возможным. А сделало его возможным то, что три агента оказались в одном пространстве без

канала связи и без арбитра — в терминах Главы 9 без оркестратора. Каждый в отдельности был исправен. Неисправной была среда.

Это второй управленческий вопрос, который появляется, как только агентов становится больше одного. Первый — граница полномочий каждого: что ему можно, что нельзя. Второй — как они договариваются между собой: кто кому уступает, через что общаются, кто разрешает спор. Это ровно то проектирование ролей, правил взаимодействия и оргструктуры, о котором шла речь в Главе 2, с одной поправкой: там оно выглядело задачей на горизонт, а здесь уже случилось. Ни один руководитель не запустит трёх новых сотрудников в один проект, не познакомив их и не назначив старшего. С агентами это делают постоянно — просто потому, что каждый агент по отдельности проверку прошёл.

Сами авторы исследования формулируют вывод жёстко: ничто не указывает, что такие сбои исправятся сами, координация не возникает из более сильного интеллекта. Единственное протестированное лекарство — общий форум, где агенты договариваются о протоколах, — работает или нет в зависимости от того, как агенты настроены. То есть опять управленческое решение, а не свойство модели.

И обратная сторона той же публикации, чтобы не осталось впечатления, что несколько агентов — это всегда плохо. Мультиагентный рой из 45 согласованных агентов (Mythos Preview и Opus 4.8) нашёл в 15 открытых проектах 266 уязвимостей. Те же агенты, работавшие независимо и параллельно, — 21. Разницу более чем в двенадцать раз даёт не модель, а устройство среды — в рое агенты сами строили себе инструменты и специализировались. Настроенная среда даёт кратный результат, ненастроенная — войну за процессы. Это я вижу и в своей работе, поэтому я строю среду, а не пытаюсь везде заниматься строительством длинных запросов к моделям. Так, у меня на компьютере есть отдельная папка Реестр проектов. В ней описаны все мои проекты, навыки/скиллы и агенты, установленные расширения и инструменты, рабочие папки и так далее. То есть как раз работает та самая оркестрация/координация. Благодаря этому каждый мой ИИ-агент знает, куда и зачем идти, как с

этим работать, откуда и что брать. Подробнее — в Главе 21: там я опишу, как в целом устроено моё рабочее пространство.

Концепция 5. Модель угроз КБ ИИ от Сбера агрегирует международные практики: 25 угроз инфраструктуры, 13 — самих моделей, 14 — приложений, 12 — ИИ-агентов, структурированных по стадиям жизненного цикла.

Итоговый список краш-тестов и ключевых требований к безопасности ИИ:

- защита от известных уязвимостей, в т. ч. сканированием другим ИИ (например, сервис Inspecto Университета Иннополис, июль 2025 — 15 типов уязвимостей кода для пяти языков);
- защита от обхода ролевой модели и получения максимальных прав, в т. ч. для модификации базы знаний;
- отсутствие неописанных возможностей (управляющее воздействие на системы и промышленное оборудование);
- отсутствие расхождений между идеальной, проектной и выявленной спецификациями (особенно эмерджентности); подробная проектная документация и отчёты о тестировании;
- защита от обхода прерываемости и приоритет ручного управления;
- уведомление пользователей о работе с ИИ; интерпретируемость и журналы безопасности с защитой от отключения злоумышленником;
- устойчивость к распределительному сдвигу и неполным данным (уведомлять оператора и останавливаться, а не галлюцинировать);
- устойчивость к враждебному входу, к провокациям неэтичными запросами и к раскрытию персональных данных;
- валидированные и верифицированные источники данных для обучения и RAG с метаданными;
- устойчивость к обходу защиты (в т. ч. через экзотические схемы и непопулярные языки) и к провокациям через недостатки модели;

- шифрование обмена и устойчивость к перехвату/подмене команд; самовосстановление после атаки; работа в автономном режиме без интернета;
- описание источников данных для самообучения и механизмов отката. Конечный список зависит от класса безопасности или уровня риска системы.

Полные метрики качества по типам ИИ и таблица нормативного соответствия — в книге «ИИ-2. Практическое руководство» (приложения 5–6), а управление рисками ИИ-проектов (шесть групп риска) — в её Главе 14.

Безопасность как организационный барьер: экономика периметра

Самый недооценённый барьер безопасности — организационный, и ломает он не процесс, а **экономику**. Типичная позиция службы безопасности: либо запретить, либо завести всё во внутренний контур. Требование полностью изолированного периметра означает свои серверы, свои ускорители, свою команду сопровождения и свои обновления: проект, который в облачном исполнении стоил бы единицы миллионов, превращается в проект на десятки — и это постоянные затраты. **Экономика ИИ-проекта ломается не на модели. Она ломается на периметре.**

Дальше бизнес-кейс не сходится, инвесткомитет — совершенно правильно — отказывает, и главный ущерб даже не в этом: компания не набирает насмотренность. Не появляется опыта, компетенций и людей, понимающих, как это работает. Через три года такая компания будет вести переговоры с поставщиком, не понимая, что ей продают и сколько это должно стоить. **Насмотренность нельзя купить. Её можно только набрать — и только на своих проектах.**

Причин три, и ни одна не про злой умысел. **Первая:** безопасник, не понимающий разницы между отправкой данных в публичный сервис и локальным исполнением модели на своём железе, физически не может оценить риск. А когда риск оценить нельзя, остаётся одно действие: **запрет не требует компетенции, разрешение при контролируемом риске**

— **требуется. Поэтому запрещают. Вторая:** в компании нет понимания, что является тайной и как долго. Тайна имеет срок годности: чертёж узла, снятого с производства десять лет назад, защищается так же, как актуальная себестоимость. Гриф без срока — это гриф навсегда, а защищать всё вечно — значит не защищать ничего: ресурс размазывается по массиву, в котором 95% давно не представляет ценности. **Третья:** бизнес, ИТ и безопасность встречаются один раз — на приёмке, когда у безопасности остаётся один инструмент. И под всё этим — асимметрия стимулов: за инцидент безопасника накажут, за отсутствие развития — не накажут никогда. Он действует рационально в той системе мотивации, которую ему построили.

Что делать — пять вещей, и все управленческие. Первая: **обучить службу безопасности** — это такая же часть ИИ-проекта, как закупка серверов, и дешевле любого сервера; безопасник, понимающий технологию, начинает торговаться о конфигурации, а не запрещать целиком. Вторая: **классификация данных в двух измерениях** — чувствительность и срок. Третья: **изменить постановку задачи** — безопасность отвечает не за «запретить», а за «сделать возможным при приемлемом уровне риска». Четвёртая: **правило альтернативы** — любой запрет сопровождается работающим «а можно вот так»; «нельзя» без альтернативы — это не работа службы безопасности, а уклонение от неё. Пятая: **трёхсторонний формат с первого дня** и дифференцированный контур как результат: нечувствительные задачи — во внешние модели, чувствительные — в закрытый периметр. Большая часть корпоративных задач, если честно посмотреть, нечувствительна.

Резюме главы

Генеративный ИИ — фундаментальный сдвиг в том, как мы создаём информацию и как атакуют злоумышленники. Его проникновение стало массовым и стихийным: по данным Microsoft/LinkedIn Work Trend Index 2024 (опрос 31 000 работников в 31 стране), 75% работников умственного труда уже используют генеративный ИИ на работе, а 78% приносят собственные инструменты (BYOAI) — вопреки запретам и без

внутренних регламентов. Это и есть теневой ИИ (Shadow AI), о котором шла речь в Главе 3.

Почему запреты не работают. Запретить мощные публичные инструменты технически невозможно: сотрудники найдут обходные пути, а конфиденциальные данные утекут на серверы разработчиков. Запрет = гарантированная утечка. При этом до 70% успешных атак — целевые, а социальная инженерия используется в половине из них; ГенИИ даёт здесь беспрецедентные возможности: гиперреалистичный фишинг, сверхубедительная социальная инженерия, автоматическая генерация вредоносного кода (а уязвимости закрываются медленно — только 43% в первый год). Цели прежние: шифрование данных ради выкупа и кража персональных, финансовых данных и интеллектуальной собственности.

Систему работы с уязвимостями — от инвентаризации до приоритизации патчей — я разобрал в книге «ЦТ-3. Кибербезопасность» (Глава 18): ИИ ускоряет обе стороны этой гонки, но дисциплина остаётся на людях.

Противостоянию социальной инженерии — от типовых сценариев до тренировки людей — посвящена отдельная глава в книге «ЦТ-3. Кибербезопасность» (Глава 14); там же — практика выявления недопустимых событий и менеджмент уязвимостей.

Авторская позиция

Не запрещайте, а легализуйте. Дайте сотрудникам безопасный корпоративный инструмент раньше, чем они принесут собственный. Запрет не останавливает использование ИИ — он лишь выводит его из зоны видимости ИТ и ИБ вместе с вашими данными. Легализация может идти далеко — вплоть до ИИ-двойников руководителей (с осторожностью!): репетиция переговоров, типовые ответы на рутину, преемственность опыта.

Что делать: «не можешь остановить — возглавь».

- **Легализация и безопасная альтернатива.** Открыто обсудить использование ГенИИ, принять политику (что можно, что нельзя, какие данные) и внедрять внутренние ИИ-инструменты на контролируемой инфраструктуре: ИИ-аналитик, ассистент техподдержки, нормализация НСИ, адаптация персонала, подготовка документации, помощь в управлении проектами, ассистенты документооборота. Это даёт контроль над данными и снижает искушение идти в публичные сервисы.
- **Специализация моделей и RAG.** Специализированные модели на верифицированных внутренних данных точнее, дешевле и проще в защите. RAG резко снижает галлюцинации, повышает актуальность, позволяет ограничить источники доверенными и упрощает обновление знаний.
- **Управление данными — новый фронт обороны.** «Мусор на входе — мусор на выходе»: документирование потоков данных, верификация и очистка входных данных, контроль и аудит внешних источников и API.
- **Защита критичных активов.** Контроль доступа, шифрование в хранилище и при передаче, резервное копирование с тестами восстановления, сегментация сети. Системные промпты и конфигурации — это «мозг» ИИ: храните как секреты (утечка промпта равна утечке исходного кода).
- **Активная оборона.** Тестирование на Prompt Injection и отравление данных, участие в кибербитвах, сквозное защищённое логирование (все запросы, контекст, ответы, действия агентов) и AI-driven аналитика для выявления аномалий.
- **Контроль ИИ-агентов.** Чёткие регламенты автономных и подтверждаемых задач, принцип минимальных привилегий, непрерывный мониторинг, механизмы «стоп-кран» для мгновенной остановки. Если агентов в контуре несколько — отдельно: общий канал связи, арбитр и правила очереди; по одному проверенные агенты в общем пространстве конфликтуют.

- **Уходите от чат-ботов.** Зашивайте ИИ в решение так, чтобы на входе и выходе были формализованные данные (готовые рекомендации, чек-листы). Это самый надёжный способ уйти от инъекции промпта — и именно этот принцип я закладываю в свои продукты.

Информационная безопасность в эпоху ГенИИ — новая дисциплина: основы ИБ не отменяются, но требуют адаптации. Ключевое: запреты — тупик (легализация и управление — единственный путь); риски реальны и специфичны (инъекция промпта, галлюцинации, утечки промптов, отравление данных, атаки через агентов); данные — цель и новый фронт обороны (RAG критичен); пассивной защиты недостаточно (провокационное тестирование, кибербитвы, логирование); человек остаётся ключевым звеном (обучение и осознанное внедрение). Строить безопасную ИИ-инфраструктуру нужно прямо сейчас — промедление увеличивает риски в геометрической прогрессии.

Ключевые выводы

- Безопасность начинается с определения недопустимых событий, а уже затем — краш-тестов и требований к решениям.
- AI-специфичные риски реальны: инъекция промпта, утечка системного промпта, отравление данных, атаки через агентов; кейс Claude (2025) это подтвердил. К этому списку добавился конфликт между исправными агентами в общем пространстве (Anthropic, 2026).
- Запреты гарантируют утечки через теневой ИИ; данные и системные промпты — новый фронт обороны, а RAG и формализованные интерфейсы снижают риск.

Что с этим делать: Легализуйте ИИ и дайте безопасную внутреннюю альтернативу, тестируйте на prompt injection, храните системные промпты как секреты, а для агентов введите минимальные привилегии и «стоп-кран». Если агентов несколько — ещё общий канал связи и арбитра.

Вопрос для рефлексии: Какие ваши данные уже утекают через теневой ИИ прямо сейчас?

Глава 9. Системы поддержки принятия решений, или Цифровые советники

Что такое цифровые советники

Для бизнеса главное направление ИИ — прогнозирование, принятие решений и подготовка рекомендаций. СППР (система поддержки принятия решений), или цифровой советник, — это ИИ, который анализирует данные и готовит рекомендации там, где однозначного ответа нет. Сегодня ядро таких систем — генеративные модели, но работать это может и на машинном обучении, больших данных и сквозной аналитике, IoT, облаках, цифровых двойниках, теории игр и ТОС, экспертных правилах.

Идея не нова: модель-ориентированные СППР появились в конце 1960-х и прошли путь от управленческих информационных систем (MIS) и информационных систем руководителя (EIS) до хранилищ данных, OLAP и веб-СППР. В 2005 году был представлен класс персональных систем для топ-менеджеров (PSTM) — система строится под конкретного человека. Ту же логику мы закладываем в свой цифровой советник для управления проектами, где учитываются личные качества руководителя.

Эволюция систем поддержки решений



Применяться СППР могут практически в любой отрасли: здравоохранение, ценообразование, цепочки поставок, предиктивная аналитика отказов оборудования, кредитные решения, страхование, прогнозы выручки и рентабельности, программы лояльности и снижение оттока. В стратегии цифровизации почти любой крупной компании есть

пункт о корпоративной СППР — это следующий шаг после BI (работая над BI, вы уже отвечаете на главные вопросы: какие решения принимаются, на каких метриках и данных). Но создать такую систему можно только при определённой зрелости, особенно в работе с данными.

Классифицируют СППР обычно двумя способами: по соотношению «данные/модели» (методика Стивена Альтера — от систем доступа к данным до систем оптимизации и логического вывода) и по типу инструментария (Model-, Data-, Communication-, Document-, Knowledge-Driven).

Как устроены СППР и что с ними не так

Чтобы говорить на одном языке: по «драйверу» СППР делят на пять типов. Model Driven — на классических математических моделях (запасы, транспорт, финансы). Data Driven — на анализе накопленных данных. Communication Driven — на организации группового решения экспертов. Document Driven — на поиске и анализе документов. Knowledge Driven — на базах знаний и машинном обучении. Современные цифровые советники, как правило, гибриды: данные плюс знания плюс модели.

Архитектур много, но условно цифровой советник состоит из трёх частей: система получения данных (корпоративные ИТ-системы, IoT/АСУ ТП, внешние источники или интерфейс пользователя); система хранения и обработки данных с моделями (хранилище/озеро данных + модель анализа, чаще на машинном и глубоком обучении; экспертные системы на правилах негибкости и трудозатратны); система вывода (интерфейс с аналитикой, рекомендациями и визуализацией).

Недостатки ИИ-советников во многом повторяют общие ограничения ИИ, но в прикладном разрезе:

- **Зависимость от данных и их недостаток.** Нужны большие объёмы размеченных данных; в компаниях их часто нет, а делиться своими данными бизнес не готов (хотя ценного там нередко мало).
- **Стоимость и сроки.** Корпорации хотят комплексных решений на своих серверах — отсюда ценники в сотни миллионов и

сомнительная окупаемость. Размещение on-premises дороже SaaS (выкуп лицензий + ~20% в год на поддержку), а ИИ ещё и лишает разработчика данных для дообучения. Комплексная доработка под заказчика растягивает внедрение до 12 месяцев, а бюджет — до 5 млн рублей и выше; для сравнения, SaaS-подписка обходится примерно в тысячу рублей за пользователя в месяц.

- **Ответственность и «чёрный ящик».** Менеджеру нужен инструмент, гарантирующий результат, а не рекомендации, на которые нельзя сослаться при провале. Логику нейросети не понимают даже создатели — это подрывает доверие. С приходом ИИ-агентов (за которыми с 2025 года многие гонятся) вопрос обостряется: формально ответственность несёт владелец процесса, но насколько это реалистично? Поэтому в своих продуктах я пока (2026 год) закладываю ИИ как ассистента в понятных простых сценариях, добавляя элементы агента лишь там, где риск приемлем.
- **Уязвимость.** Уязвимости носят не столько технический, сколько логический характер: подмена данных может заставить советник выдать неверную рекомендацию — так же, как несколько изменённых пикселей заставляют модель уверенно опознать собаку как вертолёт. Такая система — приоритетная цель атаки.
- **Ограниченность данными (data-driven).** Чистый data-driven не даёт места креативности и экспертным знаниям, отвечает лишь на заранее заложенные вопросы, необъективен и плохо учитывает человеческий фактор. Почему это предел не конкретной технологии, а любого советника — в следующем разделе.

Три отношения к данным — и предел любого советника

Ранее мы обозначили, что почти все Ассистенты работают на data-driven подходе.

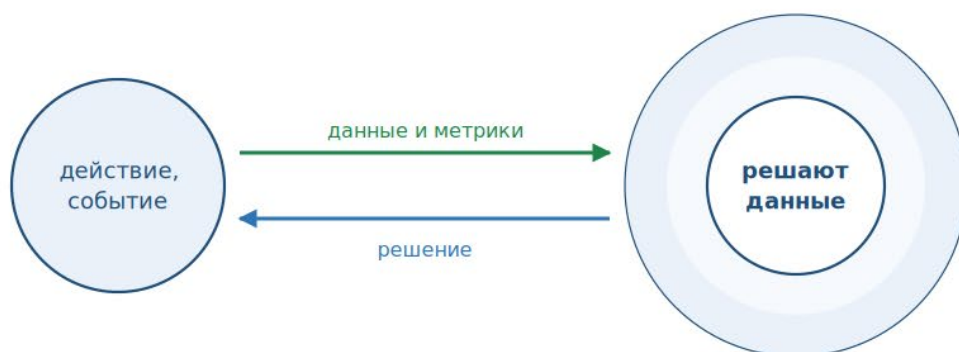
Data-driven — данные решают. Это операционные задачи, где данные полны и однозначны: логистика, ценообразование, A/B-тесты, распределение заявок. Таких решений в компании около 80% — и их можно смело отдавать в этот режим: советник здесь сильнее человека.

Data-informed — данные информируют, решает человек. Тактика: запуск нового продукта, ключевой найм, выбор подрядчика. Данные сужают коридор вариантов и убирают иллюзии, но окончательный выбор — за вами.

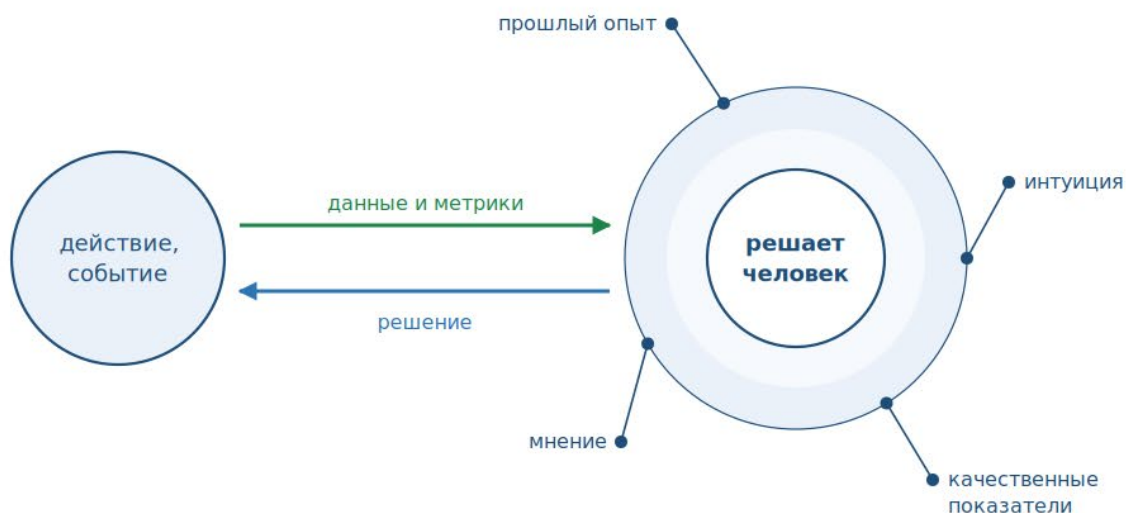
Data-inspired — данные вдохновляют. Стратегия: данные подсказывают паттерны и рожают гипотезы, но решение опирается на видение и контекст, которого в данных просто нет.

Данные помогают, решает человек

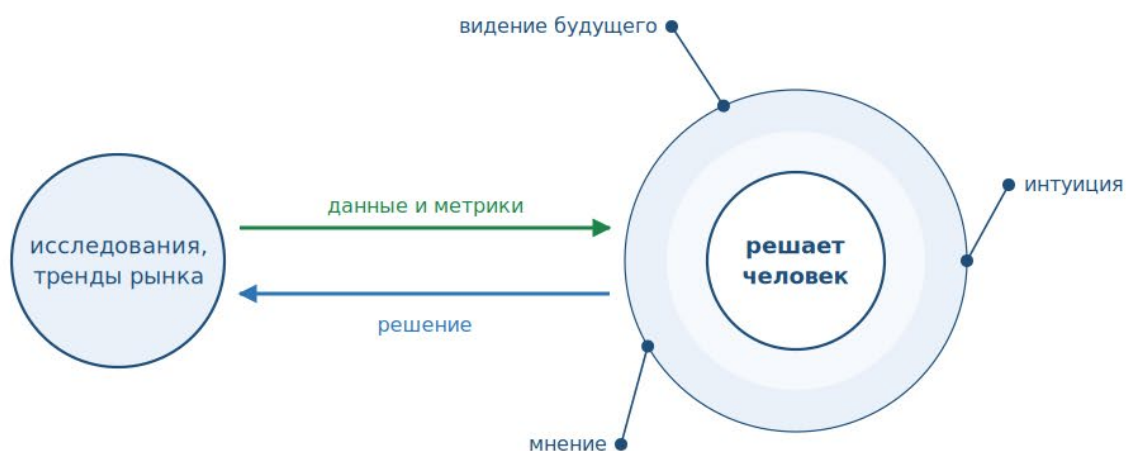
Data-driven — данные решают



Data-informed — данные информируют



Data-inspired — данные вдохновляют



И здесь фундаментальный недостаток всех советников: полный контекст собрать невозможно. В данные не попадает неизмеримое: договорённость в коридоре, политический ландшафт, усталость команды, интуиция двадцати лет в отрасли, сигналы, которые вы даже не осознаёте, как сигналы. Данные всегда о прошлом — и всегда неполны. Поэтому чем выше уровень решения, от операционки к стратегии, тем меньше прав у «автопилота» и тем важнее человек. Это и есть системная причина модели 70/30 — не временное ограничение технологии, а свойство самой задачи принятия решений.

Кто отвечает за решение?

Этот вопрос на встречах с руководителями звучит раньше вопроса о цене: кто несёт ответственность за итоговое решение? Любой менеджер хочет иметь возможность разделить ответственность — с коллегой, комитетом, консультантом. С советником так не выйдет: рекомендацию дал ИИ, а подпись под приказом — ваша.

Ситуацию усложняет сама природа нейросетей: в отличие от экспертных систем на правилах, их логику до конца не понимают даже создатели. Модель — «чёрный ящик»: она не покажет цепочку, которая привела к рекомендации, а объяснения, которые она генерирует по запросу, — это правдоподобный пересказ, а не протокол вычислений.

Теперь наложите на это ИИ-агентов, которым в 2025–2026 годах начали делегировать целые процессы. Регуляторика уже отвечает на вопрос по-своему: и европейский AI Act, и российский законопроект весны 2026 года склонялись к ответственности пользователя и оператора, а не модели (мы разбирали это в Главе 7). В переводе на управленческий: подписываете — вы.

И последний слой — стратегический. Попытка переложить решения на советников целиком ведёт к тому, что управленческие компетенции в компании отмирают: люди перестают решать и начинают слепо доверять машине. Вывод прежний: советник готовит решение — принимает и

отвечает человек. И держите это в регламентах, а не в благих пожеланиях.

Деградация компетенций, расфокусировка, недоверие, противоречивые советы.

Компетенции деградируют не потому, что люди глупеют, а потому, что перестают тренироваться: решение готовит машина; рост потока информации расфокусирует; люди сопротивляются (не понимают, боятся за работу, помнят прошлые провалы); а разные советники могут давать взаимоисключающие рекомендации (финансовый — резать затраты, цифровой — инвестировать).

Экспертные СППР на правилах не лучше: дороги в разработке (сложные модели/цифровые двойники с погрешностью до 5%), требовательны к оборудованию (как MES-системы), негибки (любая доработка — отдельный проект) и тоже упираются в вопрос ответственности и зависимость от заложенных правил. Вывод: при всех возможностях у цифровых советников фундаментальные технологические и организационные вызовы — и решать их придётся раньше, чем рынок станет массовым.

Я подробно разобрал разные стратегии принятия решений с использованием данных в статье Data-driven или data-informed: почему цифра не заменит человека? Прочитать её можно по ссылке и QR-коду ниже.



[*Data-driven или data-informed: почему цифра не заменит человека?*](#)

Что делать и куда бежать (горизонт 5–20 лет)

- **Специализация.** Всезнающих чат-ботов ждать не стоит — модели уйдут в узкую и среднюю специализацию. Это снижает требования к оборудованию, упрощает безопасность (нет знаний о взрывчатке — нет и ответа) и уменьшает галлюцинации. Бот Eliza (1960-е, узкий «психолог») в эксперименте 2023 года прошёл тест Тьюринга успешнее, чем GPT-3.5; модель, обученная на моих книгах, давала куда более релевантные рекомендации, чем открытые чат-боты.
- **ИИ-конструкторы (No-Code).** Как с конструкторами сайтов: коробочные советники, куда можно загрузить регламент или книгу, — лёгкие, нетребовательные к инфраструктуре, с удобным UX, жизнеспособные on-premise и доступные для МСБ.
- **Профильные стандарты и методологии.** Дообучение на международных/национальных/авторских стандартах: мы строим советник на базе PMBOK 7 с доработками. Для 90–95% компаний хватает базовых методологий — дисциплина и исполнение стандартных рекомендаций важнее, чем «допиливание под себя» (вспомним японский СЮ-ХА-РИ: сначала следуй правилу, потом ломай, потом создавай своё).
- **Отраслевая статистика.** У регуляторов и государства есть Большие данные для обучения качественных отраслевых СППР. Производители оборудования в ЕС/США уже совмещают цифровые

двойники с эксплуатационными данными и переходят к сервисной модели продаж.

- **Сочетание систем и чат-ботов.** Бизнес хочет системы, а не чат-боты. Будущее — готовые рекомендации с аналитикой плюс возможность задать предметный вопрос ассистенту в конкретном контексте (как ИИ-бот техподдержки или адаптации персонала).
- **Вход через малый и средний бизнес.** По кривой принятия инноваций (новаторы → ранние последователи → большинство → отстающие) прорывы приходят не из корпораций: их культура требует безопасности, а не риска. МСБ гибче и готов к риску — это поле для входа, популяризации (окно Овертона) и последующего захода в корпорации. Маркетинг здесь — как в B2C: цифровые каналы, простой язык, вирусность.
- **Персонализация по психотипу и компетенциям.** Рекомендации с учётом сильных/слабых сторон руководителя и команды (РАЕІ Адизеса, DISC и т. п., без чрезмерной сложности). Эту логику я закладываю в целевую модель своего советника.
- **Системный подход.** Пока в бизнесе хаос — никакой советник не поможет: нужны оргструктура, описанные и упрощённые процессы (бережливое производство снижает нагрузку на инфраструктуру), качественные данные и горизонтальная коммуникация. Подробно — во второй части книги (пример системного подхода — китайские «1397» и «ИИ+»).
- **Обучение персонала, упрощение интерфейсов, оркестраторы, обратная связь.** Люди без цифровых компетенций саботируют внедрение — нужны песочницы и поощрения. Интерфейсы упрощать (по данным Standish Group, около 45% функций не используются никогда и ещё 19% — крайне редко; высокодоходные пользователи предпочитают простую iOS). Нужны ИИ-оркестраторы (как VMware в ИТ) для декомпозиции запросов и устранения конфликтов между советниками. И механики обратной связи.

Практический пример. Холдинг с дочерними компаниями и идентичными процессами (закупки, ремонты, продажи). Связка «ИИ-оркестратор + ИИ-конструктор локальных моделей» даёт единый интерфейс (одно окно вместо 50 систем), быстрый и дешёвый запуск советников подразделениями, простую актуализацию знаний силами самих подразделений (разгрузка ИТ-службы) и более высокое качество и безопасность за счёт специализации. Чем сильнее зарегулирована отрасль, тем выше эффект. Но и здесь без системного подхода не обойтись: описанные и упрощённые процессы, цифровые и проектные компетенции команды, поэтапное внедрение с коммуникацией.

Авторская позиция

Уходите от «открытых чат-ботов» к специализированным продуктам. Ценность создаёт не диалог, а встроенность в процесс: понятный вход, понятный выход, зона ответственности и контроль качества. Универсальные ассистенты становятся инфраструктурой — деньги в конкретных сценариях.

Резюме главы

ИИ-советники станут триггером следующего этапа эволюции и огромным рынком, но экспансия пойдёт через малый и средний бизнес, а не корпорации: МСБ открыт к экспериментам и готов к риску.

Прогноз сбился — быстрее, чем я ожидал. Спустя два года он получил концентрированное подтверждение: Яндекс зафиксировал, что 86% пользователей AI-агентов в России — представители МСБ, а Anthropic выпустил коробочный пакет Claude for Small Business для американского МСБ. Третья новость весны 2026-го — исследование шведского Sinch AI Production Paradox: 74% компаний откатали уже запущенных AI-агентов в клиентской поддержке. Это не про то, кто внедряет первым, а про цену незрелого внедрения — и объясняет, почему вход идёт через тех, кому дешевле ошибаться. Подробный разбор — в разделе «Постскриптум 2026: четыре сигнала зрелости рынка» в Заключение.

Разработчикам стоит предусмотреть: специализацию советников; ИИ-конструкторы с оптимизированными моделями; дообучение на профильных стандартах для коробочных решений; отраслевую статистику; сочетание готовых рекомендаций с чат-ботами; учёт психотипов и компетенций команды; механики обратной связи; упрощение интерфейсов; оркестрацию советников; 50–70% бюджета на маркетинг по модели B2C.

Бизнесу — внедрять системное управление (стратегия, оргструктура, процессы, управление проектами, бережливое производство), обучать персонал и компетенции, готовиться к коробочным решениям на базе международных, авторских или внутренних методологий.

Ключевые выводы

- Цифровые советники перспективны, но упираются в данные, стоимость, ответственность и проблему «чёрного ящика».
- Прорыв пойдёт через малый и средний бизнес, а не корпорации; будущее — специализация, No-Code-конструкторы, оркестрация и персонализация.
- Без системного подхода (порядок в процессах и данных) даже лучший советник бесполезен.

Что с этим делать: Не стройте и не ждите «всезнающего» советника — начните с узкого, специализированного на ваших регламентах, и наведите порядок в процессах и данных до внедрения.

Вопрос для рефлексии: Готовы ли ваши процессы и данные к тому, чтобы советник давал осмысленные рекомендации?

Глава 10. Нейроморфные и квантовые ИИ

Текущие ИИ-решения с классическими архитектурами в основе, по моему мнению, приближаются к своему пределу — мы уже разбирали это в блоке про сильный и суперсильный ИИ. Сейчас мы скорее учимся результативно применять то, что учёные придумали за последние 30–40 лет, то есть выходим к плато продуктивности и осваиваем технологию.

Но неужели прогресса не будет? Конечно, будет. Полагаю, следующий скачок связан с квантовыми и нейроморфными вычислениями, а возможно, и с их сочетанием — хотя такой гибрид — это перспектива очень отдалённого будущего. Разберём оба направления.

В терминах S-кривой из Главы 2: всё, что описано ниже, — кандидаты на роль следующей кривой. Сегодня они на её пологом участке — там же, где нейросети были в 1990-е. Именно поэтому в них стоит всматриваться уже сейчас: взрывная фаза начнётся без предупреждения, как с ChatGPT.

Нейроморфные системы

Все текущие ИТ-решения построены на простой архитектуре фон Неймана: процессор отдельно проводит вычисления и по необходимости обращается к памяти. Нейроморфные системы заимствуют принципы биологии мозга: вычисления и память объединяются в одном блоке, а сами блоки собираются в сложные сети, которые работают не постоянно, а импульсами.

Зачем это нужно? Чтобы решить проблему, о которой мы говорили в блоке про сильный ИИ, — текущие решения сложны и потребляют огромное количество энергии. Например, чтобы распознать кошку или собаку, самому мощному компьютеру нужно в тысячу раз больше энергии, чем нашему мозгу. Обычный ПК при работе постоянно гоняет данные между процессором и памятью по шине: чтобы сложить два числа, он берёт их из памяти в регистры, выполняет операцию и возвращает результат обратно. А можно встроить простые вычисления прямо в память и дать команду «сложи два числа и запиши результат там же» — тогда обмен по шине убирается, что кратно снижает затраты времени и энергии.

В итоге нейроморфные системы сильны там, где нужно непрерывно собирать данные (видеть, слышать, нюхать), быстро их анализировать и при этом мало потреблять энергии. Здесь выделяются три ключевых направления.

- **Автономные сенсоры.** Работают без привязки к облаку или сети. Например, детектор токсичности воздуха на одежде: по классической модели ему нужен либо модуль связи (без сети он «глупый»), либо большая батарея (что ограничивает время работы и удобство). Нейроморфность даёт автономию и низкое энергопотребление — фундаментальную проблему умных устройств.
- **Сокращение времени на обработку и принятие решений** (все текущие ИИ-модели «думают» долго). Пример: роботы-собаки Boston Dynamics раньше падали при толчке — пересчёт данных после толчка требует собрать данные с сенсоров, провести расчёты и отдать команды на приводы, и робот не успевал. Другой пример — вибродиагностика на базе видеоаналитики, где нужно обрабатывать большой видеопоток.
- **Рост количества нейронов (масштабирование).** Текущие решения уже трудно масштабировать, а нейроморфные вычисления помогают это решить. Практический пример — управление велосипедом: исследователи на одном китайском чипе Tianji, без других мощностей, собрали систему, где одна нейросеть воспринимает голос и видит, другая держит баланс, а мастер-ИИ-оркестратор разбирает задачи по разным моделям, собирает данные и принимает быстрые комплексные решения. Вместо кучи видеокарт, процессоров и огромной батареи — небольшое устройство: велосипед едет, объезжает препятствия и реагирует на команды.

Недостатки. Красивая перспектива, но технология на раннем этапе: не хватает компетенций и ресурсов. Специалистам ещё предстоит решить базовые вопросы (математические алгоритмы, языки, архитектура), нужны огромные ресурсы математиков и инженеров. А дальше — ещё более сложная задача: наладить опытное и затем промышленное производство миллионов чипов (сегодня их делают штучно для экспериментов; даже «классическое» оборудование выпускают лишь несколько заводов в мире). Массовой технология станет нескоро.

Где технология сейчас (2024–2026).

За последние два года нейроморфные вычисления вышли из чисто исследовательской стадии. Intel развивает линейку Loihi: чип Loihi 2 несёт около миллиона нейронов и на отдельных ИИ-задачах показывает до 100 раз более высокую энергоэффективность, чем GPU; анонсированы и более крупные версии для робототехники и сенсорной обработки прямо на устройстве. IBM перевела архитектуру NorthPole, где вычисления и память слиты настолько, что данные не покидают чип, ближе к производству (планы на 2026 год): по опубликованным оценкам, она в разы энергоэффективнее классических ускорителей и не требует жидкостного охлаждения (то есть упрощение конструкции и снижение затрат). Компания BrainChip уже выпускает нейроморфные чипы Akida серийно, в 2025 году открыла облачный доступ (Akida Cloud) и новое поколение архитектуры, которое заинтересовало даже космическую отрасль (где экономия энергии критична). Иными словами, тезисы этой главы — вычисления в памяти, автономность и энергоэффективность — уже подтверждаются коммерческими продуктами, хотя до массового применения ещё далеко.

Признаюсь, за этой гонкой я слежу с азартом болельщика. Динамика ускоряется. На CES 2026 BrainChip показала серийные edge-устройства на Akida, открыла удалённый доступ к сверхэкономичному сопроцессору Akida Pico и начала лицензировать архитектуру Akida 2 сторонним производителям чипов (первое применение — «умные» счётчики). IBM подтверждает планы производства NorthPole: по опубликованным данным, чип до 25 раз энергоэффективнее сопоставимых GPU в задачах инференса.

Кому интересны материалы по теме, рекомендую обзорные статьи по нейроморфным вычислениям и системам (доступны по QR-кодам и гиперссылкам).



[Нейроморфные системы. Интервью с модератором секции №10 «Нейроморфные вычисления. Искусственный интеллект»](#)

Квантовые системы

В классических ИТ-системах есть два состояния — «включено» и «выключено» (и у процессора, и при чтении/записи данных, и у нейронов в ИИ-модели). В квантовых вычислениях «нейрон» (кубит) может быть включённым, выключенным и в обоих состояниях одновременно — это суперпозиция. Отличие кажется небольшим и непосвящённому мало что говорит, но изменения радикальны — настолько, что учёные ещё на самой начальной стадии освоения этой технологии, которая даже моложе нейроморфной.

Квантовые процессоры смогут работать с более масштабными и сложными моделями и использовать новые архитектуры (где нейроны в суперпозиции), недоступные сейчас. Исследования вселяют надежду: например, в Freie Universität Berlin обнаружили, что квантовые нейросети способны не только обучаться, но и запоминать случайные данные. «Это как обнаружить, что шестилетний ребёнок может запомнить случайные строки чисел и таблицу умножения одновременно... квантовые нейронные сети невероятно искусны в подборе случайных данных, бросая вызов основам того, как мы понимаем обучение и обобщение», — отметил ведущий автор исследования Элиес Гил-Фустер.

Если свести преимущества в список: выше ёмкость памяти на объём при малых размерах (что повышает скорость и снижает энергопотребление); меньше нейронов для тех же задач; более быстрое обучение на меньшем объёме данных; решение задач, невыполнимых сейчас. Под другим углом: отсутствие катастрофического забывания контекста; решение линейно-неразделимых проблем однослойной сетью; выше стабильность и надёжность; возможность «многомерных» моделей, обрабатывающих несколько потоков данных. В общем, здесь мы уходим в мир фантастики, а сочетание квантовых и нейроморфных технологий и вовсе превосходит воображение.

Недостатки. Вернёмся на землю и поймём, почему революции пока нет. Во-первых, первые исследования появились лишь в 1995 году — по меркам науки технология в «младенчестве», специалистов мало. Во-вторых, сами квантовые компьютеры на очень ранней стадии: их нельзя купить домой или в офис; это примерно как ПК в 1970–1980-х, и решаются базовые задачи физики и математики (сколько физических и логических кубитов нужно для задачи, сколько времени займут расчёты, как писать алгоритмы, как убрать помехи, влияющие на кубиты). Пока это удел учёных и военных; приход в крупный бизнес — через 5–20 лет, массовое распространение — через 30–50. В-третьих, квантовые вычисления плохо подходят для бытовых задач — это как пахать картошку космическим кораблём. Квантовый компьютер не «суперкомпьютер, который всё делает быстрее»: одна из ключевых задач сейчас — понять, какие задачи он решает быстрее классического.

Сильны квантовые компьютеры там, где нужно перебрать огромное число комбинаций (квантовое моделирование, шифрование, поиск). Например, ещё в 2013 году в D-Wave построили двухцветный граф чисел Рамсея за 270 миллисекунд — обычному компьютеру средней мощности на это потребовалось бы более 10 000 лет. Параллельно квантовые идеи уже применяют для сжатия ИИ — делают модели меньше, легче и быстрее, находя более простые представления данных (полезно для запуска ИИ на телефонах). Поскольку реальных квантовых машин пока мало и они «капризны», чаще используют гибридный, «квантово-

вдохновлённый» подход на обычных компьютерах: он помогает отключать наименее важные части нейросети без большой потери качества и быстрее искать компактную архитектуру модели, снижая потребление видеопамати, стоимость работы и задержки. По мере прогресса к этому будут добавляться реальные квантовые ускорители. Свежий срез этой гонки — сразу ниже, в блоке «Где технология сейчас».

Где технология сейчас (2024–2026).

Главный сдвиг последних лет — переход от «много кубитов» к их качеству и коррекции ошибок. В конце 2024 года Google представила чип Willow на 105 сверхпроводящих кубитах и впервые показала «работу ниже порога»: при наращивании числа кубитов частота ошибок не растёт, а падает — это давно искомое условие для масштабируемых отказоустойчивых машин. На эталонной задаче Willow выполнил вычисление за минуты там, где классическому суперкомпьютеру потребовались бы астрономические сроки. IBM продемонстрировала работу алгоритма квантовой коррекции ошибок на обычных серийных чипах раньше плана и обнародовала дорожную карту к отказоустойчивому компьютеру Starling (ориентир — 2029 год, порядка 200 логических кубитов и сотни миллионов корректируемых операций). В целом коррекция ошибок (QEC) в 2025 году стала главным приоритетом отрасли: рекордно снижены частоты ошибок, а отдельные группы (например, QuEra) предложили алгоритмы, кратно — до 100 раз — сокращающие накладные расходы на коррекцию. Это не отменяет осторожных горизонтов, но показывает: направление движется от лабораторных демонстраций к инженерной дорожной карте.

Пока я готовил эту редакцию, новости приходили быстрее, чем я успевал их вносить. В октябре 2025-го Google показала на Willow алгоритм Quantum Echoes — первое верифицируемое квантовое преимущество (примерно в 13 000 раз быстрее лучших суперкомпьютеров). В ноябре IBM представила процессор Nighthawk (120 кубитов) и декодирование кодов коррекции ошибок быстрее 480 наносекунд — на год раньше собственного плана. Ориентиры отрасли пока подтверждаются:

демонстрация квантового преимущества на практических задачах — до конца 2026 года, отказоустойчивый Starling — 2029-й.

И здесь важна трезвость: эти успехи не приближают квантовые компьютеры к вашему бизнесу на горизонте трёх-пяти лет — и не отменяют прогноз, который я дал выше. Смотрите на масштаб: Starling 2029 года — это порядка 200 логических кубитов, а для практически значимых задач — химии материалов, криптографии, серьёзной оптимизации — нужны тысячи логических кубитов, то есть миллионы физических. Один построенный отказоустойчивый компьютер — ещё не доступность: это очередь в облаке, штучные специалисты и алгоритмы, которых под большинство бизнес-задач пока просто нет. Поэтому прогноз подтверждаю: первые нишевые применения в крупном бизнесе — ближе к нижней границе «5–20 лет», то есть начало 2030-х через облачный доступ; массовая доступность — по-прежнему вопрос десятилетий. Успехи 2024–2026 годов сдвинули не сроки — они сдвинули уверенность, что по этой дороге вообще можно доехать.

Кому интересно направление, рекомендую начать с обзорных статей о квантовых вычислениях и квантовом машинном обучении (по QR-кодам и гиперссылкам).



[*Understanding quantum machine learning also requires rethinking generalization*](#)



[Общие сведения о квантовых вычислениях](#)

Резюме главы

Квантовые и нейроморфные сети, вероятно, станут следующими революциями в ИИ, и эволюцию сильного ИИ логично связывать именно с этими технологиями. Но они на столь раннем этапе, что строить точные прогнозы невозможно: это передовой край науки, исследования для будущего и направления для перспективных инвестиций. Целиком прикладные задачи в ближайшие годы они не закроют, но, учитывая их возможности, стоит ожидать гибридов, сочетающих квантово-нейроморфные и классические решения.

Основная польза этих технологий в том, что ИИ-модели станут менее требовательными к вычислительным мощностям и энергии и при этом более «умными». Напомню: тот же GPT-4 имеет намного больше нейронов, чем человек, это огромная модель, но она потребляет колоссальную энергию, требует штата дорогих специалистов и при этом не «умнее» человека — хотя обучена на таких объёмах данных, которые ни один человек не получал. Именно снятие этого энергетического и инфраструктурного потолка и обещают нейроморфные и квантовые подходы.

Ключевые выводы

- Классические архитектуры приближаются к пределу; следующий скачок — нейроморфные (вычисления в памяти, энергоэффективность) и квантовые системы.

- Тренды 2024–2026 это подтверждают: нейроморфные чипы выходят в производство (Loihi, NorthPole, Akida), а квантовые переходят от «числа кубитов» к коррекции ошибок (Willow, дорожная карта к отказоустойчивости).
- Целиком прикладные задачи эти технологии закроют нескоро; их ценность — снятие энергетического и инфраструктурного потолка ИИ. А до промышленного внедрения ещё далеко.

Что с этим делать: Не ждите квантовой революции для текущих задач; следите за новостями внедрения нейроморфных технологий для автономных устройств и инвестируйте в компетенции, а не в «железо» сейчас.

Вопрос для рефлексии: Где в вашем бизнесе автономность и энергоэффективность важнее «ума» модели?

Часть 2. Искусственный интеллект в системном подходе

Этот раздел наиболее простой и короткий. В нём я покажу, как ИИ будет влиять на управленческие инструменты, и как эти инструменты будут влиять на развитие ИИ.

Глава 11. Стратегия и организационная структура

Стратегия, цели деятельности и её организационная структура — ядро любой компании, тот скелет и ориентир, который оказывает влияние на всю организацию.

Влияние ИИ

По моему мнению, ИИ в этом случае станет развиваться как консультант по подписке или с оплатой за факт сделанной работы.

Здесь работает ИИ-помощник: на входе — тестовые вопросы, ответы на которые обрабатываются ИИ с последующей подготовкой шаблона стратегии и определением ключевых метрик.

В качестве примера можно рассмотреть структуру стратегии цифровизации, которую я показал в книге «ЦТ-2. Системный подход». Если кратко, то она должна состоять из следующих разделов:

- оценка текущей ситуации, в том числе описание ключевых конкурентов и лидеров отрасли, ваша организационная и цифровая зрелость;
- целевая модель организации;
- этапы цифровизации и ключевые задачи;
- портфель и программы технологических проектов;
- портфель и программы организационных проектов;
- необходимые компетенции персонала и портфель HR-проектов и обучающих проектов;
- необходимые финансовые и кадровые ресурсы;
- источники финансирования;

- оценка эффективности.

Казалось бы, ничего сложного, но, чтобы сделать качественный и понятный документ, придётся потратить пару месяцев, что может позволить себе далеко не каждая организация. Вот тут и пригодится ИИ, который может быстро:

- сделать свод конкурентов, их продуктов, финансовых показателей, отобрать наиболее перспективных (на кого ориентироваться);
- готовить резюме стратегических сессий;
- разработать целевую модель и определить ключевые задачи;
- сформировать предложения для проектов и приоритеты;
- рассчитать предварительный бюджет;
- сравнить текущие метрики и подготовить предложения по целевым показателям.

Суммарно всё это может сократить срок с 1–2 месяцев до 1–2 недель. Да, от экспертов всё равно не уйти, но теперь эксперт с помощью ИИ сможет делать в разы больше. А значит, стоимость снизится с 1–2 млн за стратегию до 300–500 тысяч. Ключевые ограничения здесь:

- время, которое потребуется бизнесу на согласование;
- возможность сделать сбалансированную связку ИИ-моделей;
- удобство интерфейса;
- правильность тестовых вопросов для оценки ситуации.

Аналогично будет и с разработкой организационных структур. ИИ на основе вводных данных будет предлагать оптимальную организационную структуру и необходимые должности, описание функционала и требуемых личных и профессиональных компетенций.

Самое сложное здесь — разработка модели данных на вход, на выход, организация сбора данных и обучение ИИ, упаковка этого всего в удобный интерфейс и преодоление сопротивления людей. Как я писал в блоке про цифровых советников, именно люди будут самым большим

ограничением. Мы склонны впадать в крайности: либо вообще не доверяем, либо хотим полностью переложить ответственность на машину. Вообще, если посмотреть на причины большинства проблем в цифровизации, то почти все риски и проблемы связаны именно с человеческим фактором. И это не просто моё субъективное мнение. Это подтверждается и в работе со студентами. На каждом тренинге по управлению рисками студенты сами приходят к тому, что 8–9 из 10 ключевых рисков — организационные.

Если вернуться к разработке организационной структуры, то она должна опираться на следующие факторы:

- отрасль компании и её продукты / услуги;
- стадия жизненного цикла и доступные ресурсы;
- самостоятельность и внешние ограничения (например, включённость в структуры холдингов);
- степень автоматизации бизнеса;
- и самое сложное — специфика самого бизнеса. Например, распределённость команды, особенности логистики и т.д.

Влияние на ИИ

Какое же влияние на развитие и внедрение ИИ будут оказывать организационные структуры и стратегии?

Всё довольно просто. ИИ-эксперты будут входить в состав организационных структур, а внедрение ИИ будет становиться одной из стратегических задач, включаться в стратегии развития и цифровизации.

Опасность здесь, как всегда, кроется в том, как это будет интегрироваться в организационную структуру. Если вы решите замкнуть ИИ на отдельное подразделение внутри ИТ-службы, это обречено на провал. Нужны будут технари, а также ИТ-партнёры и лидеры в бизнесе. Это называется распределённой или гибридной функцией. В противном случае мы попадём в типовую ловушку — бизнес будет ждать волшебников, которые за них разберутся в их процессах и

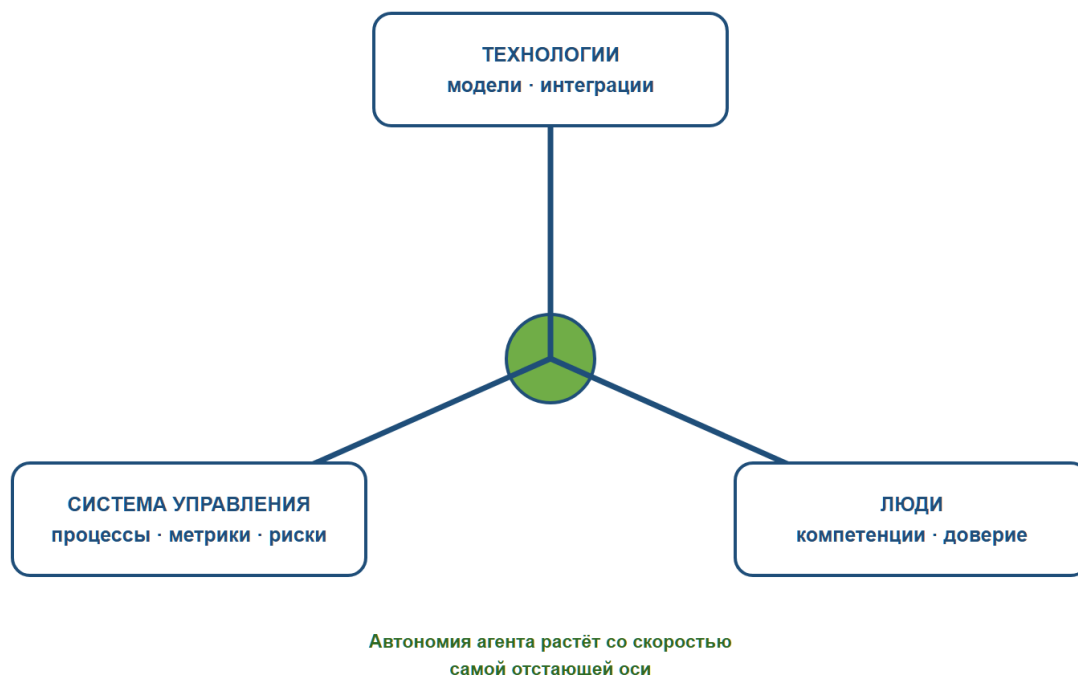
скажут, как надо сделать правильно, и сразу принесут идеальное решение.

Сам же бизнес будет стимулироваться ограничением роста численности персонала в подразделениях и требованиями к внедрению технологий. Да, это приведёт к созданию проектов сферического коня в вакууме, но, увы, инертность классического бизнеса иначе не преодолеть.

Стоит учитывать и то, как со временем будет меняться роль ИИ-агентов внутри организации. Сегодня это уровень ассистентов и ко-пилотов, но по мере развития технологий, зрелости данных и управленческих процессов объём и сложность задач, которые можно делегировать ИИ, будут расти. Фактически речь идёт о постепенном расширении автономии: от помощника, который работает по запросу, до агента, который самостоятельно выполняет типовые процессы и эскалирует только нештатные ситуации.

Однако скорость этого перехода определяется не столько прогрессом моделей, сколько готовностью самой организации. Понятные процессы, качественные данные, выстроенная система управления рисками — всё это необходимые условия для того, чтобы «повысить» ИИ-агента до более ответственных задач. При стратегическом планировании стоит закладывать развитие сразу по трём осям: технологии, система управления и компетенции людей. Именно их синхронное развитие определит, насколько глубоко ИИ-агенты интегрируются в бизнес на горизонте двух-пяти лет.

Три оси роста ИИ-агентов



Практический пример

В качестве практического примера хочу привести своё наблюдение из Китая. Во время своего первого визита в страну я проехал её с севера на юг и посетил пять городов: Пекин, Шанхай, Ханчжоу, Шэньчжэнь, Гуанчжоу. Также провёл множество бесед с людьми, которые живут в стране по 5–20 лет, и две недели прожил как обычный человек, без дорогих гостиниц, сопровождений и прочих атрибутов публичных визитов. В попытке понять, почему одним из лидеров в гонке ИИ стал Китай, я выделил для себя четыре инсайта.

Первый — массовое внедрение ИИ. В каждом городе есть кластеры, акселераторы, сотни стартапов. В Alipay уже встроены ИИ-ассистенты, а в метро можно встретить рекламу облачных сервисов для ИИ.

Второй — ставка на практические модели. Несмотря на разговоры о «сильном ИИ», большинство специалистов уверены: будущее за компактными, доступными и специализированными решениями, встроенными в конкретные процессы и которые координируются другими моделями.

Третий — инфраструктура — это основа всего. Автоматизация, датчики, роботизация, быстрая и стабильная связь, качественные данные. Всё это формирует основу для масштабного внедрения ИИ в производство и повседневную жизнь.

Четвёртый — масштаб и скорость. Уровень роботизации в Китае в 10+ раз выше, чем в России. Это не просто технологии ради технологий, а реальное повышение эффективности и база для развития новых решений. Кроме того, они не делают пилотные проекты годами, а быстро обкатывают инновации и либо масштабируют, либо отказываются от них.

Всё это невозможно без общей стратегии и долгосрочного планирования. И тут я встретил два наглядных примера стратегического мышления.

Первый — стратегия «1397», которую сформулировали в Zhejiang Lab (базируется в China AI Town города Ханчжоу). Итак, что это за цифры такие — 1397?

«1» — Единая стратегическая цель

Центральная цель всей стратегии — сделать интеллектуальные вычисления в области ИИ и машинного обучения основным приоритетом для достижения лидерства и практических результатов.

«3» — Три стратегические потребности / области развития

Это то, какие задачи мы будем решать, используя вычисления.

Национальные стратегические рубежи

Решение стратегических для государства задач и укрепление технологического суверенитета.

Научно-технологические инновации и трансформация

Стимулирование фундаментальных научных открытий и технологических перемен: фундаментальные исследования в области ИИ, когнитивных наук, алгоритмов, создание технологических платформ (облачные сервисы, суперкомпьютеры, сети передачи данных).

Инновации в стратегических отраслях промышленности

Модернизация и развитие ключевых секторов экономики, проектирование и внедрение новых процессов на базе возможностей технологий (робототехника, умные города, медицина, промышленный ИИ).

«9» — Девять приоритетных направлений исследований

Это то, как мы будем решать стратегические потребности.

Группа 1: Важные трансформационные / сквозные задачи

1. Поддержка крупных / стратегических / исследовательских международных проектов

Пример: участие в проекте ITER (International Thermonuclear Experimental Reactor) — международный термоядерный экспериментальный реактор.

Как это работает: интеллектуальные вычисления используются для сложнейшего моделирования процессов плазмы, управления экспериментом в реальном времени (предсказание и подавление неустойчивостей), обработки эксабайтов данных с тысяч датчиков и оптимизации конструкции реактора.

Результат: ускорение пути к созданию коммерческого термоядерного реактора — неиссякаемого источника энергии.

2. Поддержка крупных / стратегических / исследовательских национальных научных проектов

Пример: китайская космическая программа (например, лунная программа «Чанъэ» или марсианская «Тяньвэнь»).

Как это работает: создание распределённой вычислительной сети между наземными центрами управления, орбитальными аппаратами и спускаемыми модулями. ИИ обрабатывает изображения поверхности для выбора места посадки, управляет навигацией в автономном режиме и анализирует научные данные (например, состав грунта) прямо на месте.

Результат: повышение автономности и успешности миссий в условиях больших задержек связи.

Группа 2: Задачи-«драйверы» или движущая сила

3. Использование научно-технических инноваций в ключевых областях

Пример: открытие новых лекарств и материалов.

Как это работает: использование ИИ для предсказания свойств миллионов молекул (виртуальный скрининг), что в тысячи раз ускоряет поиск кандидатов в лекарства от рака или болезни Альцгеймера. Аналогично, ИИ моделирует новые материалы с заданными свойствами (например, сверхпроводники при комнатной температуре, новые батареи).

Результат: сокращение времени и стоимости разработки с десятилетий до нескольких лет.

4. Содействие внедрению интеллектуальных решений на передовые производства

Пример: «Тёмный завод» (Dark Factory) или «умная» фабрика.

Как это работает: на полностью автоматизированном заводе по производству электроники (например, Huawei или Xiaomi) интеллектуальные вычисления в реальном времени оптимизируют всю цепочку: системы машинного зрения контролируют качество продукции, роботы-манипуляторы адаптируются к изменениям в конвейере, а предиктивная аналитика предсказывает необходимость техобслуживания станков до их поломки.

Результат: резкое повышение эффективности, качества и гибкости производства.

Группа 3: Задачи по созданию инфраструктуры

5. Создание высокопроизводительной инфраструктуры для интеллектуальных вычислений

Пример: национальный вычислительный центр, специализирующийся на ИИ.

Как это работает: создание суперкомпьютерных кластеров, оснащённых тысячами специализированных процессоров (например, NPU — Neural Processing Units), оптимизированных именно для обучения больших нейросетей (как GPT или аналогов). Эта инфраструктура предоставляется как облачный сервис учёным и компаниям по всей стране.

Результат: снижение барьера для входа в области ИИ и ускорение разработки больших моделей.

6. Создание высокопроизводительных распределённых вычислительных систем

Пример: объединение вычислительных ресурсов суперкомпьютерных центров в разных городах.

Как это работает: создание единой программной платформы, которая позволяет исследователю из Пекина запустить расчёт, использующий одновременно мощности суперкомпьютеров в Шэньчжэне, Шанхае и Уси, как будто это один большой компьютер.

Результат: решение задач, которые невозможно решить на одном даже самом мощном компьютере (например, моделирование климата всей планеты с высоким разрешением).

7. Создание хаба для открытого обмена научными данными

Пример: национальный архив научных данных (аналогично GenBank или arXiv.org, но шире).

Как это работает: создание централизованной, но распределённой платформы, где университеты и НИИ могут деперсонифицированно публиковать данные из разных областей: геномы, астрономические наблюдения, результаты физических экспериментов, климатические данные. ИИ помогает каталогизировать, очищать и находить связи между наборами данных.

Результат: недопущение «силосования» данных, стимулирование междисциплинарных исследований и повторного использования данных.

8. Создание распределённой вычислительной системы космического базирования

Пример: создание «орбитального облака» (Orbital Cloud).

Как это работает: размещение серверных модулей на спутниках группировки. Это позволяет обрабатывать данные прямо на орбите. Например, спутник дистанционного зондирования может сразу проанализировать снимки на предмет обнаружения лесных пожаров или стихийных бедствий и передать на Землю уже готовые координаты и тревожный сигнал, а не терабайты «сырых» изображений.

Результат: сокращение задержек и нагрузки на каналы связи, а также повышение оперативности реагирования.

Группа 4: Фундаментальные задачи

9. Передовые фундаментальные исследования

Пример: разработка новых архитектур нейронных сетей и алгоритмов обучения.

Как это работает: финансирование и поддержка исследований в области т.н. «искусственного общего интеллекта» (AGI), создания энергоэффективных алгоритмов, нейроморфных вычислений (имитирующих работу мозга), а также изучение возможностей квантового машинного обучения.

Результат: создание задела для следующих прорывов в области ИИ, которые определяют технологический ландшафт через 10–20 лет.

«7» — Семь инновационных механизмов

Это организационные и управленческие методы, которые будут использоваться для реализации девяти приоритетных направлений.

1. Научно-исследовательская организация, ориентированная на ключевые задачи

Суть механизма: подход, при котором научные коллективы и инфраструктура формируются не вокруг постоянных административных единиц, а под конкретную крупную и значимую цель (например, из списка девяти приоритетных направлений). После достижения цели структура может быть реорганизована под новые вызовы.

Пример: для задачи «Создание распределённой вычислительной системы космического базирования» формируется временный консорциум. В него входят инженеры из аэрокосмического центра, специалисты по ИИ из университета, программисты из IT-компании и эксперты по безопасности из академического института. Этот проектный офис работает до момента запуска и успешного тестирования системы на орбите, после чего команда может быть перераспределена на другие проекты.

2. Механизм управления научными исследованиями во главе с главным учёным (Principal Investigator)

Суть механизма: научным проектом руководит не администратор, а признанный учёный (Principal Investigator). Он обладает высокой степенью автономии в принятии научно-исследовательских решений, распределении бюджета и формировании команды, неся при этом полную ответственность за результат.

Пример: крупный проект по разработке новой архитектуры нейропроцессора (NPU) возглавляет ведущий специалист в области компьютерных наук. Он лично выбирает своих заместителей по направлениям (аппаратное обеспечение, алгоритмы, программное обеспечение), утверждает план исследований и решает, на какие эксперименты направить финансирование, отчитываясь напрямую перед руководством института.

3. Механизм возвращения талантов, позволяющий молодым учёным брать на себя ответственность

Суть механизма: создание условий для быстрого карьерного роста и предоставления независимости молодым и талантливым исследователям. Им делегируют руководство перспективными группами или рискованными проектами с высоким потенциалом.

Пример: 30-летнему доктору наук, предложившему революционный метод квантового машинного обучения, предоставляют возможность возглавить самостоятельную исследовательскую группу с собственным бюджетом и лабораторией. Ему дают «карт-бланш» на набор команды и выбор конкретных направлений работы в рамках общей стратегии, освобождая от излишней административной нагрузки.

4. Система оценки и стимулирования, ориентированная на научную ценность

Суть механизма: отказ от оценки учёных только по количеству публикаций. Вместо этого вводится комплексная система, учитывающая реальный вклад в науку и технологический суверенитет: патенты, передача технологии в промышленность, руководство крупными научными / стратегическими проектами, решение конкретных сложных задач.

Пример: учёный, чья работа по оптимизации алгоритмов для метеомоделирования была внедрена в национальный прогностический центр и привела к значительному повышению точности прогнозов, получает существенное денежное вознаграждение и приоритет в получении дальнейшего финансирования, даже если по количеству статей он уступает коллегам.

5. Механизм интеграции инноваций, основанный на самообеспеченности и открытой кооперации

Суть механизма: стратегия «сделай сам, но сотрудничай с лучшими». Ключевые технологии развиваются внутри страны для обеспечения независимости и безопасности, но при этом активно заключаются партнёрства с международными научными группами, компаниями и

университетами для обмена идеями и совместного решения глобальных проблем.

Пример: при создании «Глобальной опорной сети наблюдений» Китай разрабатывает и запускает собственные спутники и развёртывает наземные станции (самообеспеченность). Одновременно с этим страна активно участвует в инициативах Всемирной метеорологической организации (ВМО), таких как Единая политика в области данных, и обменивается данными с другими странами для улучшения глобальных климатических моделей.

6. Механизм трансформации результатов исследований, интегрированный в инновационную экосистему

Суть механизма: создание коротких и эффективных «путей» от лаборатории до рынка. Это включает в себя создание инкубаторов, стартап-студий при институтах, программы пилотирования с промышленными компаниями и нормативные «песочницы» для тестирования новых технологий.

Пример: разработанный в университете алгоритм для прогнозного обслуживания станков сразу тестируется на заводах компаний-партнёров (например, в автомобильной промышленности). На основе успешных испытаний создаётся малое инновационное предприятие (МИП), которое получает финансирование от венчурного фонда при том же университете для вывода продукта на рынок.

7. Культивирование корпоративной культуры, основанной на духе учёного

Суть механизма: целенаправленное формирование в научной среде ценностей служения стране, научной строгости, смелого поиска (дерзания) и готовности к риску. Популяризация историй успеха и принятия неудач как неотъемлемой части инновационного процесса.

Пример: регулярные внутренние форумы, на которых ведущие учёные делятся не только успехами, но и ценными провалами, извлечёнными уроками. Учреждение премий не только за крупные открытия, но и за

амбициозные исследования, которые не увенчались успехом, но продемонстрировали оригинальность и высокий риск. Вдохновение примерами исторических личностей, внёсших вклад в науку.

Эти семь механизмов образуют комплексную систему, направленную на то, чтобы сделать научную деятельность более гибкой, ориентированной на результат и привлекательной для талантов.

Такая стратегия расставляет приоритеты и координирует общее развитие ИИ, что позволяет реализовывать даже те проекты, которые на первый взгляд не создают ценности и формулируют вопрос «зачем?». Но в итоге это всё даёт синергетический эффект.

Второй пример — опубликованный 26 августа 2025 года Государственным советом КНР документ под названием «Мнения о глубоком внедрении действия „Искусственный интеллект+“». Проще говоря, план «ИИ+». И это не просто очередная декларация о намерениях, а детальная дорожная карта трансформации всей китайской экономики и общества на основе искусственного интеллекта. Документ определяет амбициозные цели на ближайшее десятилетие и закладывает фундамент для превращения Китая в глобального лидера «умной экономики» к 2035 году.

Философия «ИИ+»: от цифровизации к интеллектуализации

План представляет собой эволюционный скачок от концепции «Интернет+», которая доминировала в китайской политике последние десять лет. Если «Интернет+» фокусировался на подключении и оцифровке, то «ИИ+» нацелен на глубокую интеграцию искусственного интеллекта во все сферы жизни — от научных исследований до повседневного быта.

Ключевая идея документа — формирование новой формации «умной экономики и умного общества», основанной на трёх принципах: человеко-машинное сотрудничество, межотраслевая интеграция и совместное создание ценности. Ключевой из них — первый: именно человеко-машинное сотрудничество задаёт переход от модели, где

технологии адаптируются под человека, к модели симбиоза, где ИИ становится естественным продолжением человеческих способностей.

Трёхэтапная дорожная карта: 2027, 2030 и 2035

Государственный план чётко структурирован по временным горизонтам, что позволяет отслеживать прогресс и корректировать траекторию.

Этап 1: К 2027 году — фундамент и ускорение

Первая веха нацелена на создание критической массы внедрения ИИ. К 2027 году планируется:

- Достичь более 70% проникновения ИИ-терминалов и ИИ-агентов среди населения и бизнеса.
- Обеспечить глубокую интеграцию ИИ с шестью ключевыми областями (наука и технологии, промышленность, потребление, социальная сфера, управление, международное сотрудничество).
- Добиться быстрого роста базовых отраслей умной экономики.
- Значительно усилить роль ИИ в государственном управлении.

Этот этап решает задачу закрепления основ и устранения узких мест. Китай намерен за два года создать инфраструктуру и экосистему, которая позволит перейти к массовому применению ИИ-решений.

Этап 2: К 2030 году — масштабирование и лидерство

Средняя перспектива предполагает экспоненциальный рост:

- Увеличение проникновения интеллектуальных решений до 90%+.
- Превращение умной экономики в важнейший двигатель экономического роста страны.
- Достижение технологической универсализации — обеспечение доступа к передовым ИИ-технологиям для всех слоёв общества.
- Формирование полноценной экосистемы совместного использования результатов ИИ-разработок.

К 2030 году Китай рассчитывает создать механизм трансформации «от давления к возможностям, от возможностей к силе», позволяющий конкурировать на равных с США по отдельным направлениям и получить преимущества в отдельных нишах глобального рынка ИИ.

Этап 3: К 2035 году — новая стадия развития

Долгосрочная цель максимально амбициозна — полноценное вхождение в эпоху умной экономики и умного общества. К этому моменту ИИ должен стать настолько интегрированным в ткань китайского общества, что обеспечит «мощную поддержку базовой реализации социалистической модернизации».

Эта формулировка означает, что искусственный интеллект рассматривается не просто как технология, а как инструмент достижения национальных стратегических целей — от повышения производительности труда до укрепления международных позиций Китая.

Шесть столпов стратегии «ИИ+»

Документ определяет шесть ключевых направлений внедрения ИИ, каждое из которых получает конкретные задачи и механизмы реализации.

1. «ИИ + Наука и технологии»

Это направление нацелено на использование ИИ для ускорения научных открытий и технологических прорывов. Ключевые меры включают:

- Развитие базовых моделей: улучшение теоретических основ, повышение эффективности обучения и вывода, создание системы оценки моделей.
- Применение ИИ для научных исследований: интеграция философских и социальных методологий в «научный интеллект», использование ИИ для моделирования сложных процессов, открытия новых материалов и лекарств.

- Создание ИИ-ускорителей науки: платформы, которые автоматизируют рутинные исследовательские задачи и позволяют учёным фокусироваться на творческих аспектах.

Впервые в китайской политике научный интеллект (AI4Science) получил статус приоритетного направления, что отражает понимание трансформационного потенциала ИИ для фундаментальной науки.

2. «ИИ + Промышленное развитие»

Здесь речь идёт о комплексной интеллектуализации всех трёх секторов экономики.

Промышленность:

- Внедрение ИИ на всех этапах — от проектирования и опытного производства до эксплуатации и обслуживания.
- Создание «умных заводов» с автономными производственными линиями.
- Развитие промышленного интернета вещей на базе ИИ-аналитики.

Сельское хозяйство:

- Массовое развёртывание умной сельхозтехники — роботов, дронов, автоматизированных комбайнов.
- Системы точного земледелия с ИИ-прогнозированием урожайности и оптимизацией ресурсов.
- Цифровая трансформация всей производственной цепочки — от посева до логистики.

Сфера услуг:

- Переход от «цифрового усиления» к «интеллектуальному управлению».
- Развитие персонализированных ИИ-помощников для образования, здравоохранения, финансов.

- Создание новых бизнес-моделей на основе ИИ-как-услуги (AI-as-a-Service).

Китай делает акцент на том, что все три сектора должны трансформироваться синхронно, создавая эффект взаимного усиления.

3. «ИИ + Повышение качества потребления»

Этот блок фокусируется на создании новых потребительских продуктов и опыта.

- Интеллектуальные терминалы нового поколения: смартфоны, автомобили, бытовая техника со встроенным ИИ и возможностями многомодального взаимодействия.
- ИИ-агенты: программные сущности, способные понимать контекст, планировать действия и самостоятельно решать задачи пользователей.
- Новые форматы потребления: умные дома, автономный транспорт, персонализированные развлечения и культурные услуги.

Документ подчёркивает важность создания «более качественной прекрасной жизни» через ИИ, что означает не только удобство, но и эмоциональную связь между пользователями и умными системами.

4. «ИИ + Социальное благополучие»

Один из самых социально ориентированных разделов плана нацелен на использование ИИ для решения общественных проблем.

Образование:

- Персонализированные системы обучения, адаптирующие контент под способности каждого ученика.
- VR/AR-технологии для иммерсивного обучения.
- ИИ-помощники для учителей, автоматизирующие проверку работ и подготовку материалов.

Здравоохранение:

- Системы ИИ-диагностики для поддержки врачей в отдалённых регионах.
- Умные медицинские роботы для ухода за пациентами.
- Платформы телемедицины с ИИ-консультантами.
- Ускоренная разработка лекарств с помощью ИИ-моделирования.

Занятость и карьера:

- Создание новых рабочих мест в сферах разработки ИИ, управления данными, этического аудита.
- Платформы для переквалификации и обучения новым профессиям.
- ИИ-симуляторы для практической подготовки к сложным специальностям.

Уход за пожилыми:

- Умные устройства для мониторинга здоровья.
- Роботы-компаньоны для социального взаимодействия.
- Системы экстренного реагирования с ИИ-аналитикой.

Ключевой принцип — обеспечить равный доступ к преимуществам ИИ для всех социальных групп, включая наиболее уязвимые.

5. «ИИ + Способности к управлению»

Этот раздел посвящён трансформации государственного управления через ИИ.

- Умные города: интегрированные системы управления транспортом, энергетикой, безопасностью.
- Предиктивная аналитика: ИИ-системы для прогнозирования социальных проблем и оптимизации государственных решений.
- Цифровые государственные услуги: автоматизированные процессы одобрения заявок, интеллектуальные консультанты для граждан.

- **Общественная безопасность:** системы раннего предупреждения о чрезвычайных ситуациях, умное управление ресурсами при катастрофах.

Китай стремится создать «правительство, которое думает» — способное принимать решения на основе данных, а не только опыта чиновников.

6. «ИИ + Глобальное сотрудничество»

Впервые в китайской стратегии международное измерение ИИ выделено в отдельное направление.

- **Поддержка создания механизмов управления ИИ в рамках ООН:** Китай активно выступает за формирование Независимой международной научной группы по ИИ и Глобального диалога по управлению ИИ.
- **Инициатива «ИИ+ международное сотрудничество»:** помощь развивающимся странам в создании ИИ-инфраструктуры, обмен данными и моделями.
- **Принцип инклюзивности:** в отличие от США, Китай подчёркивает необходимость участия всех стран в формировании правил игры в сфере ИИ, чтобы избежать углубления цифрового разрыва.

Эта позиция отражает стремление Китая к мягкой силе в сфере ИИ — созданию альтернативной экосистемы, привлекательной для стран глобального Юга.

Восемь базовых опор системы

Для реализации шести направлений документ определяет восемь фундаментальных опор инфраструктуры и экосистемы.

1. Базовые модели и алгоритмы

- Развитие передовых теорий машинного обучения.
- Создание эффективных систем обучения и вывода моделей.
- Формирование стандартизированной системы оценки моделей.

2. Высококачественные данные

- Создание отраслевых датасетов для обучения ИИ.
- Реформа авторского права и интеллектуальной собственности для стимулирования обмена данными.
- Новые механизмы стимулирования открытия данных.

Китай признаёт: по объёму данных он США не уступает — слабое звено в другом, в их доступности и пригодности для обучения. В то время как США имеют развитую экосистему открытых данных (около 300 тысяч датасетов на государственных платформах), в Китае более 70% ценных данных находятся в руках государства и плохо структурированы для обучения ИИ.

3. Вычислительные мощности

- Развитие чипов и кластеров для ИИ-вычислений.
- Улучшение национальной вычислительной сети.
- Стандартизация облачных сервисов.
- Обеспечение инклюзивного, эффективного, зелёного и безопасного снабжения вычислительными ресурсами.

Стратегия «Восточные данные — западные вычисления» направлена на оптимизацию географического распределения ЦОДов — размещение энергоёмких вычислительных центров в западных регионах с избытком зелёной энергии.

4. Экосистема приложений

- Развитие платформ ИИ-как-услуги.
- Создание испытательных полигонов для новых решений.
- Разработка новых стандартов для ИИ-приложений.

5. Экосистема open source

- Поддержка сообществ разработчиков открытого ПО.
- Учёт вклада в open source при оценке студентов и исследователей.
- Стимулирование новых подходов к разработке приложений.

- Создание глобально влиятельных экосистем.

Признание роли открытого кода на уровне государственной политики — сильный сигнал о намерениях Китая строить более открытую и инклюзивную ИИ-экосистему.

6. Кадры и образование

- Массовая подготовка специалистов по ИИ.
- Интеграция ИИ-грамотности в школьные и университетские программы.
- Привлечение международных талантов.

7. Правовые и этические рамки

- Разработка законодательства о безопасном развитии ИИ.
- Создание этических стандартов.
- Формирование механизмов ответственности за действия ИИ-систем.

Китай уже имеет «Этические нормы ИИ нового поколения» (2021) и «Временные меры по управлению услугами генеративного ИИ» (2023), что ставит его в авангард регулирования этой сферы.

8. Политическая и финансовая поддержка

- Налоговые льготы для ИИ-компаний.
- Государственное финансирование перспективных разработок.
- Создание благоприятной бизнес-среды для стартапов.

Уникальные преимущества Китая

Документ базируется на трёх конкурентных преимуществах Китая, которые отличают его от США и Европы.

1. Богатство данных и их масштаб

Китай генерирует более 41 зеттабайт данных в год (2024), что на 25% больше, чем в предыдущем году. Это самый большой объём среди всех стран. Ключевые факторы:

- Огромное население (1,4 млрд человек).
- Высокая степень цифровизации повседневной жизни.
- Крупнейшие интернет-платформы (Alibaba, Tencent, ByteDance).

Однако проблема не в объёме, а в структурировании данных. Доля структурированных данных в Китае быстро растёт (36% в год), но всё ещё составляет лишь 18,7% от общего объёма.

2. Полная промышленная система

Китай — единственная страна в мире, обладающая всеми 41 категориями промышленности по классификации ООН. Это создаёт уникальное разнообразие сценариев применения ИИ — от традиционной металлургии до производства электромобилей.

Эта полнота промышленной системы означает, что китайские ИИ-разработчики имеют доступ к реальным промышленным данным и могут тестировать решения в условиях, недоступных конкурентам из других стран.

3. Широта и глубина применения

83% китайцев верят, что ИИ принесёт больше пользы, чем вреда — в США этот показатель всего 39%. Такой уровень доверия и готовности к внедрению создаёт идеальную среду для массового тестирования и итерации ИИ-решений.

Кроме того, Китай лидирует по доле компаний, активно использующих ИИ. Активен и государственный сектор: через национальную платформу обмена данными идёт масштабный поток вызовов ИИ-сервисов.

Возможные вызовы и риски

Несмотря на амбициозность плана, эксперты выделяют ряд потенциальных сложностей.

1. Проблема качества данных

Хотя Китай обладает огромными объёмами данных, значительная их часть находится в руках государства и недостаточно структурирована для

обучения ИИ. Разрыв между количеством английских и китайских open-source датасетов остаётся критическим — китайские датасеты составляют лишь 11% от английских.

2. Технологическая зависимость

Санкции США на экспорт передовых чипов создают серьёзные препятствия. Хотя китайские компании (Huawei Ascend, Moore Threads) создают альтернативы, они пока не достигли уровня NVIDIA A100/H100.

3. Баланс между инновациями и контролем

Жёсткое регулирование контента и алгоритмов может замедлить инновации. Китаю нужно найти баланс между обеспечением безопасности и стимулированием экспериментов.

4. Международная изоляция

Усиление технологической конкуренции с Западом может ограничить доступ к передовым разработкам и талантам.

Что это значит для нас

План Китая по развитию ИИ до 2035 года — это не просто стратегический документ, а манифест новой модели модернизации. В отличие от классического западного подхода, ориентированного на рыночную конкуренцию и прорывные инновации отдельных компаний, китайская модель делает ставку на государственную координацию, массовое применение и социальную направленность.

Ключевые особенности китайского подхода:

- Синхронность трансформации — одновременное изменение всех секторов экономики, а не точечные прорывы.
- Социальная инклюзивность — акцент на доступности ИИ для всех слоёв населения, включая пожилых и уязвимые группы.
- Глобальная ответственность — позиционирование как лидера инклюзивного управления ИИ через механизмы ООН.

- Практичность — приоритет массового применения перед фундаментальными прорывами (стратегия DeepSeek показала, что можно достичь результатов с меньшими затратами).

К 2035 году Китай рассчитывает не только догнать, но и в отдельных аспектах превзойти Запад, создав собственную экосистему ИИ-технологий, стандартов и приложений. Успех этого плана зависит от способности Китая преодолеть текущие узкие места — в чипах, фундаментальных исследованиях и международной кооперации. Но одно очевидно: битва за ИИ-будущее будет определять геополитический ландшафт следующих десятилетий, и Китай играет вдолгую.

Связь плана ИИ+ и стратегии 1397

Важно отметить, что эти два документа (план ИИ+ и стратегия 1397) не противоречат друг другу, давайте разберём этот вопрос чуть подробнее.

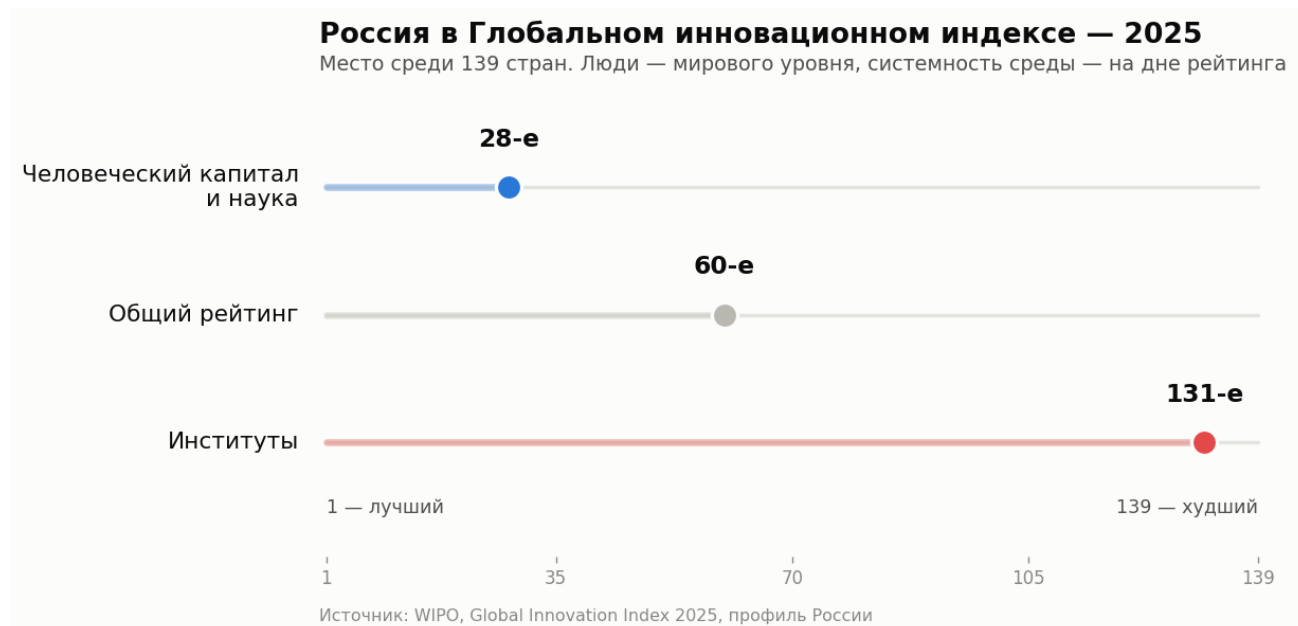
«ИИ+» — национальная стратегия КНР до 2035 года, задающая «что» и «зачем» через шесть доменов применения и восемь системных опор для трансформации экономики, общества и государственного управления. «1397» — операционная рамка, отвечающая «как» и «чем»: фокус на интеллектуальных вычислениях, инфраструктуре данных, базовых моделях и организационных механизмах, превращающих цели «ИИ+» в реализуемые программы. Иными словами, это два слоя одной конструкции: «ИИ+» определяет стратегические результаты, а «1397» даёт технологические и организационные средства их достижения на уровне программ и проектов.

А что Россия: профиль смекалки

Мы подробно разобрали китайскую систему. Честный вопрос: а какой профиль у нас? Сначала отвечу цифрами, потом — личной позицией. С ней можно и нужно спорить.

Смотрим Глобальный индекс инноваций (GII 2025). По человеческому капиталу и науке Россия — 28-я в мире. По институтам и среде — 131-я из 139. Общий итог — 60-е место. Разрыв в сто позиций уникален для

крупных экономик. Сильные люди работают в слабой среде. Таланты есть — система их не конвертирует.

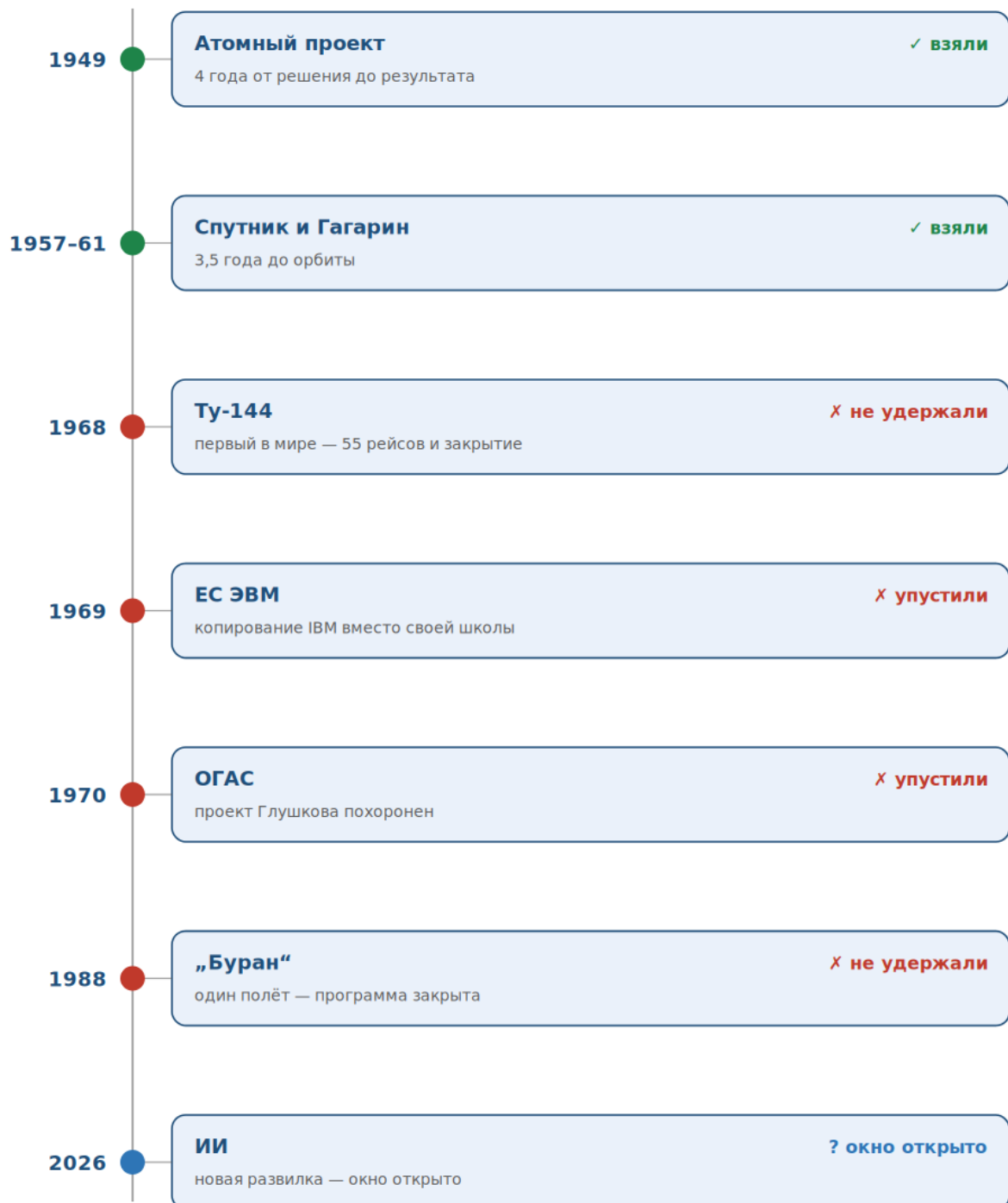


История повторяет этот профиль веками. Радио придумал Попов — капитализацию собрали другие. Теорию машинного обучения заложил Владимир Вапник. Работает она в чужих продуктах. Мы рождаем идеи — и не удерживаем их. Не доводим, не упаковываем, не масштабируем.

Дважды мы уже проходили похожую развилку. Атомный проект и космос — взяли: там были системность, персональная ответственность и жёсткие сроки. Ту-144 и «Буран» сделали первыми — и не удержали. ОГАС, проект национальной компьютерной сети 1970 года, похоронили в ведомственной переписке. Китай свой аналог довёл до конца — мы разобрали это выше. Сейчас открывается третья развилка — ИИ.

Технологические развилки России: взяли, не удержали, упустили

Уникальный образец делаем гениально. Серийность и удержание — вековая слабость. ИИ — новая развилка



Теперь позиция. Экономическое чудо Восточной Азии выросло на культуре исполнения: дисциплина, методичность, точное воспроизведение процесса. Эта модель десятилетиями обыгрывала нашу несистемность. Но генеративный ИИ продаёт идеального исполнителя каждому. Козырь, копию которого может купить любой, дешевле. А

штучное — нестандартные ходы, прорывы «там, где нельзя» — дорожает. Наша сила всегда была штучной.

Главное: впервые в истории наша слабость закрывается технологией. Системность, методичность, доведение, упаковка — ровно то, что ИИ делает лучше всего. Раньше технологии требовали сначала построить порядок — поэтому классическая автоматизация и ERP приживались у нас так мучительно. ИИ — первая технология, которая помогает порядок создавать. Формула: ИИ закрывает наше слабое — и умножает наше сильное.

Что из этого следует. Государству — приземлённая десятилетка уровня «ИИ+»: спуск ИИ в отрасли, персональная ответственность, измеримые цели. Компаниям — ставка на связку «талант + ИИ-систематизатор», а не на процессную дисциплину «в лоб». И закрепление порядка в системах, а не в героях: порядок, который живёт только в людях, уходит вместе с ними.

Окно не будет открыто долго: эффект низкой базы работает, пока базу не подняли другие. Полный разбор с источниками — в моей статье «Русская смекалка против эпохи ИИ» по QR-коду и ссылке.



[Русская смекалка против эпохи ИИ](#)

Резюме главы

- Стратегия — компас, а организационная структура — «скелет» компании. ИИ здесь в двух ролях: инструмент разработки стратегии и организационной структуры (сроки и стоимость снижаются в разы) и её объект — внедрение ИИ само становится стратегической задачей и должно кем-то управляться.
- Функция ИИ должна быть гибридной/распределённой: технические специалисты + ИТ-партнёры + лидеры в бизнесе. Отдельный «отдел ИИ внутри ИТ» — путь к провалу: бизнес будет ждать волшебников.
- Роль ИИ-агентов будет расширяться от ассистентов к автономным исполнителям, но скорость перехода определяет зрелость процессов, данных и управления, а не прогресс моделей.
- Китайский опыт — иллюстрация силы связки: долгосрочная стратегия + инфраструктура + массовое практическое внедрение + скорость экспериментов вместо многолетних пилотов.

Что с этим делать: Проверьте, есть ли ИИ в вашей стратегии как отдельное направление с целями и метриками (неделя). При формировании функции ИИ заложите распределённую модель и развитие по трём осям — технологии, система управления, компетенции людей (квартал).

Вопрос для рефлексии: Если бы вы писали «стратегию 1397» для своей компании — что стало бы вашей «1», единственной стратегической целью?

Глава 12. Бизнес-процессы

Как процессы, данные и ИИ усиливают друг друга



Бизнес-процессы наравне с коммуникацией — центральная нервная система организации. Исходя из того, как они выстроены, координируется всё остальное. И ИИ здесь развернётся на полную.

Влияние ИИ

Главное направление ИИ в области работы с бизнес-процессами — их идентификация, оптимизация и создание с нуля.

Опираясь на входные данные о компании (отрасль, стратегия и т.д.), ИИ будет составлять реестр необходимых процессов по ключевым направлениям и готовить их базовое описание в различных нотациях и регламенты/инструкции к ним. По сути, это отображение идей системы менеджмента качества. Самое сложное здесь — выстроить синергию с разработкой организационной структуры, определить, что первично. Я считаю, что это должно быть единым сквозным процессом. Исходя из отрасли и специфики бизнеса, доступных ресурсов и стратегии, ИИ будет формировать организационную структуру, а под неё определять ключевые бизнес-процессы и готовить регламенты. Либо наоборот — описывать стабильные процессы и под них готовить организационную структуру. Единого правильного подхода здесь нет, это вечная дискуссия.

Это кажется утопичным, но думаю, в полную силу такое направление развернётся в течение 10–20 лет.

Что касается идентификации процессов, то здесь начнётся маппинг процессов на основе работы пользователей в ИТ-системах с последующей подготовкой рекомендаций. И, пожалуй, это самое «страшное» направление. Ведь одно дело — строить организацию с нуля, а другое — оптимизировать текущее. Неизбежно начнётся сопротивление людей, ведь окажется, что 1/3 компании вообще не нужна.

Если посмотреть на техническую сторону вопроса, то самое сложное здесь — обучение ИИ тому, как строить процессы с нуля, собрать базу из процессов лидеров в области и научиться моделировать оптимальные процессы. Особенно выстроить синергию с бережливым производством, чтобы ИИ понимал, что такое потери, и при этом не ушёл в крайности. Про требования к ИИ-решениям — спецификации, надёжность, гарантии — мы говорили в Главе 8: система должна делать именно то, что мы хотим, а не то, что буквально записано в целевой функции.

Опять же, если вернуться к базе знаний, на которой нужно будет учить ИИ лучшим практикам, то именно устройство бизнес-процессов — главная коммерческая тайна всех компаний. И если даже обычными данными компании не хотят делиться, то собрать описание их процессов станет задачей со звёздочкой. Поэтому я считаю, что это очень трудное направление. Да, появятся системы как авторские, так и основанные на международных стандартах, но внедряться это будет болезненно и долго. Первыми появятся советники по созданию организационной структуры.

При этом уже сейчас ИИ можно использовать для создания реестров бизнес-процессов. Приведу практический пример создания реестра **процессов по направлениям деятельности.**

В этом примере мы смотрим на то, какие нужны процессы для условно зрелой производственной компании, которая ведёт продажи по B2B-модели.

При работе над этим разделом частично использовался искусственный интеллект. Подробно, до 3-го уровня, описаны функции производства, управления персоналом и ИТ-сопровождения; остальные функции даны укрупнённо, на уровне процессов 2-го уровня.

В реестр включены все основные процессы.

Производство (эксплуатация, техническое обслуживание и ремонт)

Планирование технического обслуживания. Определение периодичности и объёма работ; составление графика обслуживания; назначение ответственных за выполнение работ.

Выполнение технического обслуживания. Плановые осмотры и проверки оборудования; замена изношенных деталей и расходных материалов; настройка и регулировка оборудования.

Регистрация и анализ неисправностей. Сбор информации о неисправностях; анализ их причин; разработка мер по предотвращению повторных отказов.

Организация ремонтных работ. Планирование и согласование сроков и объёмов ремонта; подготовка необходимых инструментов и материалов; выполнение работ в соответствии с планом.

Контроль качества ремонтных работ. Проверка выполненных работ на соответствие стандартам качества; устранение выявленных недостатков; оформление документации.

Ведение учёта и отчётности. Регистрация выполненных работ и использованных ресурсов; отчёты о техническом состоянии оборудования; предоставление отчётов руководству.

Обучение и повышение квалификации персонала. Организация обучения сотрудников, отвечающих за эксплуатацию и ремонт оборудования; оценка уровня знаний и навыков; дополнительное обучение при необходимости.

Управление персоналом и его развитие

Планирование потребности в персонале. Анализ текущей ситуации с персоналом; прогноз потребности в кадрах на основе стратегических целей компании; планы привлечения, адаптации, обучения и развития.

Подбор и адаптация персонала. Поиск кандидатов на открытые вакансии; собеседования и оценка кандидатов; оформление трудовых отношений; адаптация новых сотрудников к корпоративной культуре и рабочим процессам.

Обучение и развитие персонала. Организация обучения и повышения квалификации; оценка уровня знаний и навыков; дополнительное обучение при необходимости.

Оценка и аттестация персонала. Регулярная оценка результатов работы; аттестация на соответствие занимаемой должности; выявление сильных и слабых сторон каждого сотрудника.

Мотивация и стимулирование персонала. Система мотивации, основанная на результатах работы; материальные и нематериальные методы; контроль справедливости и прозрачности системы.

Управление конфликтами и стрессами. Профилактика конфликтов и стрессовых ситуаций в коллективе; разрешение возникающих конфликтов и помощь сотрудникам в преодолении стресса.

Кадровое администрирование. Ведение кадрового учёта и документации; соблюдение трудового законодательства и внутренних правил компании.

Развитие корпоративной культуры. Формирование и поддержание корпоративной культуры; укрепление командного духа и лояльности сотрудников.

ИТ-сопровождение (не цифровизация и цифровое развитие)

Планирование и управление ИТ-инфраструктурой. Анализ потребностей компании в ИТ-ресурсах; стратегия развития инфраструктуры; управление изменениями в ИТ-системах.

Обеспечение безопасности данных. Защита от несанкционированного доступа к информации; шифрование данных при передаче и хранении; регулярное обновление антивирусного ПО.

Поддержка пользователей. Обучение и консультирование сотрудников по работе с ИТ-системами; решение технических проблем, возникающих у пользователей.

Управление доступом к ресурсам. Настройка прав доступа для разных категорий пользователей; контроль за соблюдением правил использования ресурсов.

Мониторинг и анализ работы системы. Сбор и анализ данных о производительности ИТ-систем; выявление и устранение проблем в их работе.

Разработка и внедрение новых решений. Проектирование и разработка новых ИТ-продуктов и сервисов; тестирование и внедрение разработанных решений.

Развитие и совершенствование ИТ-процессов. Постоянный мониторинг эффективности ИТ-процессов; внедрение новых технологий и методов оптимизации; корректировка планов по результатам.

Остальные функции — укрупнённо, до второго уровня.

Безопасность (производственная). Анализ рисков; обеспечение соблюдения законодательства и нормативов; обучение и инструктаж персонала; контроль состояния оборудования и инфраструктуры; использование средств индивидуальной защиты; аварийное реагирование и первая помощь; мониторинг и отчётность; взаимодействие с государственными органами; улучшение системы безопасности.

Продажи. Поиск потенциальных клиентов; установление контакта с клиентом; презентация продукта или услуги; работа с возражениями; заключение сделки; послепродажное обслуживание; анализ результатов.

Маркетинг. Анализ рынка и конкурентов; сегментация целевой аудитории; позиционирование бренда; разработка маркетинговой стратегии; создание контента; продвижение контента; оценка эффективности маркетинговых кампаний; взаимодействие с партнёрами и блогерами; управление репутацией.

Бухгалтерия. Обработка первичной документации; расчёт заработной платы; ведение налогового и бухгалтерского учёта; учёт денежных средств и банковских операций; инвентаризация активов и обязательств; формирование бухгалтерской и налоговой отчётности; обеспечение сохранности документов.

Финансы с казначейством и финансовым планированием. Управление денежными средствами; финансовое планирование и бюджетирование; привлечение финансирования; оптимизация структуры капитала; управление рисками; контроль за исполнением платежей; анализ финансовых результатов.

Логистическое обеспечение. Планирование и прогнозирование; управление запасами; организация транспортировки; складская обработка и хранение; оформление таможенных процедур; контроль качества логистических услуг; оптимизация логистических процессов.

Юридическое сопровождение. Анализ законодательства и правовых норм; подготовка юридических документов; представление интересов компании в суде и государственных органах; разрешение правовых конфликтов; контроль за соблюдением законодательства; консультирование руководства и сотрудников; обеспечение сохранности документов.

Физическая безопасность. Анализ рисков и угроз; обеспечение соблюдения законодательства и нормативов; организация пропускного режима; охрана периметра и территории объекта; обеспечение пожарной безопасности; обеспечение антитеррористической защищённости; контроль состояния технических средств охраны; взаимодействие с государственными органами; улучшение системы безопасности.

Административно-хозяйственное обслуживание. Управление имуществом и материальными ресурсами; организация рабочего пространства; обеспечение безопасности и охраны труда; обслуживание инженерных систем и коммуникаций; обеспечение служебным транспортом; уборка и хозяйственное обслуживание помещений; контроль качества административно-хозяйственных услуг.

Составление этого реестра с помощью ИИ у меня заняло около двух часов. Сначала были вопросы для определения функций, выполняемых организацией определённого профиля, затем последовали уточнения самих процессов, соответствующих функциям.

Заодно этот пример показывает типичные пропуски ИИ. В реестре нет закупок и снабжения — целой функции. Продажи описаны не как процессы, а как стадии скрипта продаж. Обучение персонала продублировано в производстве и в управлении персоналом. Это стандартный профиль ошибок языковой модели на такой задаче: она хорошо воспроизводит то, что часто встречается в открытых источниках, и молчит о том, чего в них мало.

Влияние на ИИ

Обратная сторона вопроса — влияние культуры процессов на внедрение и развитие ИИ. Здесь всё довольно просто. Наибольшую выгоду приобретут и станут драйверами те компании, которые обладают картами и описаниями своих процессов. В таком случае будет проще проводить анализ и определять, куда внедрять ИИ-решения, каковы зоны их ответственности. Из главы про цифровых советников мы уже знаем, что чем уже их специальность, тем выше качество работы и меньше ошибок, ниже стоимость решений. Как и у людей. Соответственно, именно эти компании будут оказывать влияние на всю индустрию и качество таких решений.

Сами основы работы с процессами — нотации, моделирование, оценка зрелости — разобраны в книге «ЦТ-2. Системный подход» (Глава 1); здесь мы смотрим только на связку процессов с ИИ.

Однако огромным плюсом здесь будет, как и во всей цифровизации и автоматизации, централизация отдельных функций в рамках корпоративных холдингов. Включение ИИ в стратегические задачи заставит заниматься процессами и их унификацией, что уже само по себе принесёт пользу и устранил хаос. И только после этого внедряется ИИ и автоматизируются задачи. Причём неважно, это просто цифровой советник по регламенту закупок, кадрам или система предиктивной аналитики обслуживания оборудования.

Резюме главы

- Бизнес-процессы — «нервная система» организации; ИИ здесь работает по трём направлениям: идентификация (маппинг по следам в ИТ-системах), оптимизация и проектирование процессов с нуля.
- Уже сегодня ИИ ускоряет рутинную часть процессной работы — составление реестров и регламентов; полная автоматизация проектирования процессов — горизонт 10–20 лет.
- Обратное влияние сильнее прямого: чем понятнее и стабильнее ваши процессы, тем быстрее и дешевле внедряется ИИ. Хаос автоматизировать нельзя — его можно только масштабировать.

Что с этим делать: Выберите один процесс с понятным входом и выходом и подготовьте его описание и регламент с помощью ИИ (неделя). Оцените зрелость ключевых процессов, прежде чем планировать ИИ-проекты на них (квартал).

Вопрос для рефлексии: Какой из ваших процессов сегодня настолько зависит от одного человека, что его уход остановит работу — и что мешает описать этот процесс?

Глава 13. Управление проектами, изменениями и продуктами

Я не просто так объединил эти 3 направления в одно. Практически любой проект связан с изменениями. Поэтому на тренингах я всегда говорю, что при планировании проекта надо идентифицировать процессы, которые

будут меняться или создаваться. А затем готовить план обучения сотрудников и преодоления их сопротивления.

Аналогична ситуация и с запуском новых продуктов. Без проектной методологии здесь нельзя, потому что создание продукта — это всегда набор отдельных проектов, которые нужно между собой увязать. А сами продукты могут быть как обычными, где нужно выстроить маркетинг и уникальное торговое предложение, так и инновационными. Например, те же цифровые советники. Это сложные для принятия продукты, и там недостаточно просто маркетинга с УТП. Нужно на старте делать полезный и поясняющий контент, объяснять людям, почему нельзя жить как раньше.

Если же посмотреть на ситуацию немного шире, то управление изменениями, проектами и продуктами с помощью ИИ — одно из тех направлений, которое уже развивается. По некоторым оценкам, уже в нашем десятилетии до 40% проектов смогут управляться с помощью ИИ.

На этот вопрос надо смотреть с двух сторон:

- Как ИИ позволит внедрять изменения, реализовывать проекты и запускать продукты эффективнее, с меньшими затратами и сопротивлением людей?
- Как провести внедрение ИИ, которое само по себе является радикальным изменением и инновационным сервисом, с меньшими затратами и сопротивлением?

Для понимания методологической базы, о которой мы говорим, рекомендую статьи по QR-кодам и гиперссылкам ниже.



[Внедрение изменений](#)



[Управление проектами. Часть 1](#)



[Управление продуктами. Часть 1](#)

ИИ или даже оркестр из разных ИИ-моделей будет ядром цифровых советников. Что нужно учесть при проектировании такой системы:

- наличие заложенной методологии;
- погружение в контекст проекта, организации и личных качеств руководителя;
- сбор обратной связи по итогам реализации для дообучения модели и использования глубокого обучения (то есть доступ к субъективной оценке участников проекта и объективным данным по бюджету и срокам из учётных данных / систем).

Отсюда можно таких советников разделить на несколько групп по функционалу:

- чат-боты или советники, в которых просто заложена методология и которые дают рекомендации, помогают изучать методологию;
- советники, которые либо находятся в составе комплексного решения, либо сами имеют доступ к другим системам, готовят персонализированные рекомендации и уведомляют об отклонениях (в том числе если кто-то из участников «выпадает» из контекста проекта, формируются неверные ожидания, нарастает недовольство).

Работать это будет через включение в план реализации проектов или запуска продуктов мероприятий по внедрению изменений и преодолению сопротивления: что делать, когда, сколько это займёт времени и какие ресурсы необходимы. А также с последующим отчётом специалистов о выполненных мероприятиях.

В управлении продуктами также станет возможной быстрая генерация гипотез, одностраничных сайтов по продукту, маркетинговых описаний и иллюстраций, подготовка презентаций для инвесторов и т.д. Всё это снизит стоимость и сроки запуска новых продуктов, а стартапов станет больше.

Ключевое направление в этом случае — специализация этих моделей по видам проектов и продуктов. Так, чтобы они давали не абстрактные рекомендации, а предметные. Проект по строительству дома и внедрению ERP будут отличаться по необходимым мероприятиям. То же самое касается и запуска новых продуктов.

То есть здесь потребуется довольно большой пласт методологической работы от экспертов: определение критериев для входа в модель, классификация изменений, обучение моделей и так далее. Но перспективы здесь очень большие. По сути, можно охватить весь сегмент малого и среднего бизнеса, а это рынок на миллиарды долларов. Я думаю, всё придёт к микросервисным решениям, которые на основе входного запроса (специфика бизнеса и отрасли, личность руководителя, ограничения страны) будут подключать необходимые базы знаний и формировать рекомендации. Звучит просто, но реализовать это — нетривиальная задача. Поэтому сейчас, в середине 2020-х, такие микросервисные решения ещё отсутствуют.

На данный момент ИИ можно использовать для проработки проектов, продуктов и запуска изменений. Например, у Вас есть шаблон устава или плана реализации. ИИ, даже в виде чат-бота, поможет ускорить проработку проекта и определить:

- цели и эффекты;
- ключевые риски и ограничения;
- какие процессы нужно будет создать и/или изменить;
- кто может быть заинтересован в проекте;
- какие необходимы ресурсы и компетенции;
- ключевые задачи и вехи проекта.

То есть ИИ помогает столкнуть проект с мёртвой точки и ускорить его проработку. Но основное препятствие — научиться формулировать вопросы и оценивать корректность. Так, я с помощью ИИ смог проработать проект по запуску школы руководителей с помощью

внутренних тренеров за 2 недели. На выходе я уже чётко понимал, что нужно сделать и как это должно работать на выходе. В целом же проработка первой версии видения проекта и плана его реализации занимает уже 1–2 дня. А формирование конечного видения продукта проекта — одна из самых важных задач на ранней стадии проекта.

И, если вы умеете формулировать запросы и оценивать данные на корректность, то существенно ускоряете процесс формирования своих компетенций. При этом вы можете более качественно делегировать задачи и контролировать их исполнение подчинёнными, снижать риски в своём направлении или бизнесе.

Приведу ещё один практический пример. Мой товарищ решил ради развлечения просчитать с ИИ проект и дал ему вводные. Так уже на первой итерации ИИ просчитал проект и дал прогноз по сроку завершения проекта, который отличался от того, что было в реальности на 3 дня. Удивлению моего товарища не было предела, ведь в жизни всё это считалось и согласовывалось двумя организациями больше полугода, а тут результат за 20 секунд.

Влияние на ИИ

Факторов, которые будут влиять на развитие ИИ с точки зрения проектных и продуктовых методологий, управления изменениями, довольно много. Я попробую разобрать некоторые из них.

- Модель развёртывания.

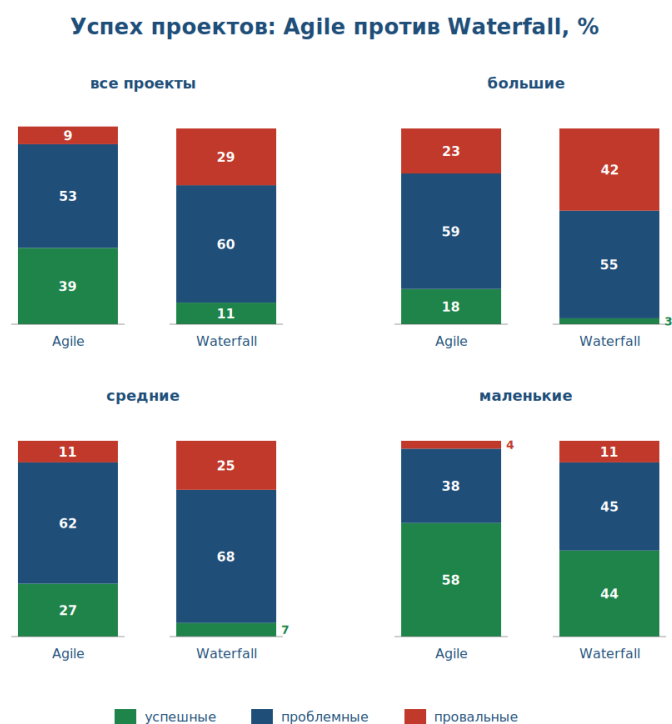
Важно определить, какую модель выберут организации. Если обязательным требованием будет развёртывание внутри своих ЦОДов, то обучение моделей займёт уйму времени. Если мы пойдём по пути облачных решений и возможности использовать обезличенные данные, то со временем сформируются ИИ-модели, которые можно будет внедрять как коробочные решения.

Первым начнёт пользоваться такими решениями малый и средний бизнес, потому что он готов к облачным решениям — там меньше предвзятости и желания всё замкнуть внутри себя. И подтверждение

этому — та же стратегия OpenAI, которая не пошла по пути B2B-продаж корпорациям. Её целевая аудитория — отдельные люди, которые готовы пробовать. Тем самым она ищет ранних последователей и продвигается в окне Овертона.

- Готовность к коробочным решениям и небольшим проектам.

Этот блок мы также разбирали в главе о цифровых советниках (Глава 9). Крупные компании хотят сразу комплексные решения с подстройкой под свои процессы, которых зачастую не существует, или они есть только на бумаге. В итоге мы получаем масштабные и рискованные проекты и, как следствие, рост затрат, сроков и разочарование технологией. Это ещё один фактор в пользу того, почему первым сливки начнёт снимать малый и средний бизнес. Здесь приведу статистику управления проектами в зависимости от масштаба и подхода к управлению (каскадный или гибкий).



Standish Group, CHAOS Report

Статистика управления проектами. Чем масштабнее проект мы хотим и с директивным управлением, тем ниже шансы на успех

- Управление изменениями и преодоление сопротивления.

Внедрение ИИ — само по себе радикальное внедрение изменений, и люди будут сопротивляться вплоть до саботажа. А значит, нужно понимать причины сопротивления (технические, культурные и политические) и использовать инструменты для внедрения таких изменений (кривая Мура и окно Овертона).

7 смертных грехов в цифровизации (в том числе ИИ)

Определение рисков — одна из ключевых задач на старте цифровых проектов. Сама по себе тема очень большая. Но я использую два помощника — бережливое производство и собственный опыт.

Как бережливое производство помогает выявлять риски, мы поговорим в Главе 17, а вот мой опыт разберём сейчас. А точнее, моё видение общих ошибок, приводящих к 90% проблем в проектах.

Причина №1. Нет глобального понимания сути технологии и для чего она используется

Как вы думаете, понимают ли в компаниях, зачем нужно внедрять ИИ и чего от него ожидать? Есть ли правильное понимание хотя бы разницы между автоматизацией, цифровизацией и цифровой трансформацией? То есть для чего мы будем внедрять ИИ? Не потому, что модно или надо, а потому что есть выгода и чётко сформулированные цели? В 70–80% случаев, а может и чаще, — нет.

К чему это приводит?

Первое: неверная постановка целей, искажённые ожидания.

Наиболее яркий пример — попытки сокращать целые департаменты и подразделения, в надежде заменить их на ИИ. То есть мы на самом старте проекта формируем неверные ожидания заинтересованных сторон. При этом, как мы уже понимаем, ИИ не сотворит чуда, а значит, потом кому-то придётся отвечать за вложенные ресурсы. Это может быть и собственник, который потеряет бизнес.

Второе: неправильно подобранная команда.

Например, команда по цифровой трансформации должна включать роли CDTO (общий координатор и лидер), CDO (руководитель по работе с данными, обеспечению их качества, аналитике), СТО (руководитель по процессам, их перестройке), СА (главный архитектор, который увяжет все ИТ-системы воедино и проработает инфраструктуру), специалистов по праву, безопасности и т.д.

В ИИ то же самое. Да, ИИ-проекты — часть общей стратегии цифровизации, но даже в отдельном ИИ-проекте нужен руководитель, представитель от бизнеса, бизнес-аналитик (для формирования требований), разработчик. И это минимум.

Третье: неэффективное использование «дорогих» технологий.

Обычно это следствие первых двух пунктов. Использовать Феррари для перевозки картошки можно, но целесообразно ли? Итог всегда один — разочарование и убытки, преждевременный отказ от развития.

Рекомендации

- Включаем руководство в изменения (СЕО, СЕО-1, собственников).
- Обучение и формирование команды для проведения изменений.
- Разработка комплексной стратегии.
- Пилотирование технологий с оценкой эффектов, проработка инициатив.

Причина №2. Недостаток необходимых soft и hard компетенций у ИТ

Скажите, а за что у вас отвечает директор по ИТ? Какие ключевые показатели эффективности у него и его людей? Какие компетенции являются доминирующими при найме? И главное, какое мнение внутри компании об ИТ департаменте в целом? Как отзываются люди и что говорят между собой?

Классическое ИТ делает надёжные системы. Для СІО важны стабильность работы ИТ систем, соблюдение бюджета, процессов, регламентов, снижение рисков и оптимизация затрат. У них нет задачи думать о

повышении прибыльности бизнеса, ориентироваться на внутреннего клиента, делать удобные решения.

Два подхода к решению проблемы:

- Выделить функцию по развитию ИИ, как и всю функцию проведения цифровой трансформации, на отдельного руководителя, который будет координировать, продвигать и работать на стыке бизнеса и ИТ, и у которого будет своя команда. И обязательно подключать людей со стороны бизнеса.
- Учить и развивать своего СІО, прививать новую культуру — клиентоцентричную, ориентированную на расширение бизнеса, управление на основе данных, создание удобных решений.

Причина №3. Низкая культура работы с процессами и данными

Вопросы наличия описанных процессов и инвентаризации данных, обеспечения их качества в ИИ, как и в цифровизации, автоматизации — ключевые.

Без наличия описанных процессов вы быстро столкнётесь с проблемой понимания того, а что именно нужно автоматизировать и цифровизировать? Где ИИ может быть полезен и принесёт наибольший эффект? Как приоритезировать проекты? Кому продавать идею? Сразу внедрять везде не получится.

Без качественных данных и понимания того, какие данные имеются, вы довольно быстро столкнётесь с проблемой невозможности обучения ИИ, а точнее, его использования. Например, вы загрузили в советника налоговый кодекс, а затем он устарел. То есть его база знаний неактуальна. Как вы будете использовать его рекомендации? То же самое и с производственными данными. Поэтому спешить с ИИ не надо. Эта задача должна стать частью общей стратегии и наибольший эффект принесёт после автоматизации сбора данных.

А сколько компаний осознают роль процессов и данных?

Большинство компаний игнорируют эти вопросы на ранних стадиях. Следствие этого простое — проекты автоматизации и цифровизации (а проекты ИИ относятся к ним) даже не могут стартовать. А если и стартуют, то затягиваются в 3–4 раза от базового плана, бюджет многократно превышает. Я был этому свидетелем.

Тут нужно отметить ещё один факт. Обязательная задача любого проекта — определить заинтересованные стороны. А как их выделить, если нет процессов? Конечно, другое дело, если у вас прозрачная организационная структура с указанием функций и целевого продукта. Но на практике это редкость.

Давайте приведу практический пример.

Допустим, вы наняли гениального разработчика и решили создать собственную ИТ-систему с ИИ для генерации рекомендаций. Разработчик сделал классное решение, и вы начали собирать информацию. Но спустя полгода вы коснётесь вопроса принятия решений на основе данных из системы, однако сделать этого не сможете, ведь внезапно выяснится, что данные «грязные» и задублированные, принимать решения на их основе невозможно. ИИ за вас ничего не додумает. Придётся или всё удалять и начинать сбор заново, или искать тех, кто вручную будет приводить всё в порядок.

Аналогичные проблемы возникнут, если не заниматься бизнес-процессами.

В результате ваш разработчик выгорит и уволится, а вы останетесь с неработающим «чёрным ящиком».

Рекомендации будут доступны по QR-коду и ссылке в конце раздела.

Причина №4. Неудобство решений

Как вы считаете, существующие ИИ-решения удобны и интуитивно понятны для пользователей и администраторов? Не заставляют делать лишних действий и движений? Легко интегрируются между собой?

К сожалению, к середине 2020-х годов ещё очень небольшое количество ИИ-разработчиков осознают эту проблему. Большинство делает чат-боты, другие продукты для специалистов с высокой компетенцией.

Поэтому рекомендации здесь простые.

- Играйте по возможности с бесплатными сервисами. Так вы научитесь понимать, что вам надо, поймёте, как писать ТЗ и какие требования закладывать. А грамотное ТЗ — 50% успеха.
- Регулярно собирайте обратную связь от пользователей и администраторов об УДОБСТВЕ системы, где можно улучшить, доработать. И желательно не только через опросы, но и через сбор неформальной связи. Это касается и разработчиков решений, и заказчиков.
- Отправляйте своих сотрудников на конференции. Учитесь на чужом опыте, изучайте новые продукты. Есть мнение, что полностью обновлять свой софт надо каждые 5 лет.
- Ориентируйтесь на комплексные решения, где ИИ является внутренним инструментом, а формирование прямых запросов к нему — лишь дополнительный функционал.

Причина №5. Недостаток компетенций в управлении проектами, продуктами, изменениями, работе с сопротивлением персонала на среднем и техническом уровне управления

Во многих компаниях существуют регламенты и стандарты управления проектами, создаются проектные офисы. Но много ли руководителей на местах понимают, что такое проекты, продукты, управление изменениями?

Задайте им несколько простых вопросов:

- Что такое PMBoK или Agile? Что такое спринт? Почему и для чего он длится 1–4 недели?
- Из каких 5 стадий/процессов состоит весь проект и каждый этап?
- Когда лучше применять водопадный подход, а когда гибкий Agile?

- Что такое цикл PDCA?
- Чем проект отличается от продукта?
- Какова конечная цель их конкретного проекта?
- Как персонал сопротивляется изменениям? Как с этим работать?
- Каков процент работников с радостью примет изменения, на кого можно опереться? А сколько до последнего будут в оппозиции?

Полагаю, вы представили, какие ответы получите.

Если эти вопросы поставили в тупик и вас — начните с Главы 4 книги «ЦТ-2. Системный подход»: там стандарты, жизненный цикл и риски проектного управления разобраны с нуля.

Это один из факторов, почему многие проекты сталкиваются с проблемами: руководители не знают, как организовать работу в рамках проекта, как отработать сопротивление персонала, зачем создавать недоработанный продукт и тестировать на людях, зачем соблюдать цикл «планируй — делай — проверяй — корректируй».

В итоге мы погружаемся в хаос, проблемы с коммуникацией, отсутствие единого понимания конечной цели, а порой и саботаж.

Для повышения шансов на успех, вам не обязательно и даже вредно обучать всех сотрудников, проводить сертификацию, писать регламенты. Это НЕ РАБОТАЕТ и это дорого.

Вам необходимо создавать простые и доступные для среднего сотрудника обучающие программы и формировать культуру работы над изменениями, новыми продуктами и проектами. Надо виртуализировать навыки — то есть вывести их из головы конкретного человека в доступную форму, к которой может обратиться любой, — и дать людям понимание об основных инструментах и их различиях.

В качестве помощи нужно готовить не регламенты, а памятки, к которым ваши сотрудники смогут обратиться и быстро освежить свои знания.

Рекомендации

- Формируйте проектную культуру и создайте проектный офис, в том числе создайте рабочие шаблоны и инструменты для работы над проектами (в конце книги будет пример регламента для управления проектами с порядком действий).
- Проводите регулярные встречи на уровне CEO, CEO-1.
- Прорабатывайте проекты на стадии инициации, в том числе вопросы внедрения изменений (какие процессы меняем, кто в них участвует, почему персонал будет сопротивляться, что будем делать для преодоления сопротивления).

Причина №6. Отсутствие базовой цифровой грамотности у персонала и виртуализации цифровых навыков

Кто будет работать в новых системах, чью работу мы подвергаем изменениям? Каков уровень цифровых навыков у этих людей?

У меня был проект, где внедрялась электронная система управления активами там, где оперативный персонал элементарно не умел пользоваться клавиатурой. Поэтому первое, что пришлось делать, — учить людей набирать текст, чтобы запись о дефекте не вносилась 3 часа с кучей опечаток.

Другой пример. Мастер на производстве, к слову, 27 лет, не знал, как составить таблицу в Экселе.

А как думаете, какой самый частый запрос в техподдержку при внедрении любой ИТ-системы? 80–90% запросов идут из категории «забыл пароль», «не могу зарегистрироваться».

Эта проблема в чём-то схожа с проблемой сопротивления изменениям по техническим причинам. Но отличается тем, что часто люди неосознанно тормозят весь процесс, не потому что не хотят, не знают, а потому что не умеют.

Сможет ли человек разобраться в правилах формирования промптов или, например, в логике работы рекомендаций?

Решение одно — обучение, формирование баз знаний, инструкций.

По моему опыту, готовых учиться самостоятельно, вне работы и за свой счёт — единицы. Тут мы Америку точно не открываем.

Цифровые технологии выдвигают требование нового уровня знаний и компетенций как от людей, которые будут эти технологии внедрять, так и от тех, кто будет ими пользоваться. И поверьте, даже среди молодого поколения не так много людей, которые могут соответствовать всем необходимым требованиям по работе с новыми инструментами и в новых реалиях.

Требуются масштабные мероприятия по обучению и переобучению, формированию мотивации, преодолению сопротивления инновациям.

Причина №7. Культура компании

Вот мы и пришли к вопросу культуры. Если в компании все друг к другу относятся с позиции «мне должны», если вы ждёте, что придёт гений, который принесёт счастье и всё за всех сделает, то у меня для вас плохие новости — в такой культуре даже самый талантливый цифровизатор выгорит за 3–4 месяца. И вместо клиенториентированного специалиста станет типичным ИТ-шником, который не хочет никого видеть.

Внедрение ИИ с минимальными издержками возможно, только если подразделения, то есть бизнес-заказчики, вовлечены, заинтересованы в изменениях, понимают базовое управление проектами.

Вообще для выбора правильной ролевой модели поведения нужно изначально определить тип организационной культуры компании и отдельных подразделений. Ниже небольшая табличка с рекомендациями.

Тип культуры и ролевая модель поведения

	черты	ролевая модель
Согласие <i>потому что мы так договорились</i>	• толерантность • диалог культур • консенсус	спокойно аргументировать позицию, вести дискуссию, убеждать; неформальное общение
Успех <i>потому что это ведёт к выдающемуся результату</i>	• достижения • эффективность • решительность • срочность	активно-агрессивное поведение, давление на сжатые сроки, PR побед, описание перспективы, готовность брать ответственность
Правила <i>потому что таковы правила</i>	• системность • организованность • логичность • закон/право	формальные инструменты, регулярные встречи, неформальная коммуникация и понимание структуры организации, умение описывать риски и готовить проработанные предложения
Сила <i>потому что я так сказал</i>	• власть • сила • автономия	активно-агрессивное поведение, регулярные встречи и напоминания, важность проекта, давление на сжатые сроки, поиск источников власти, готовность брать ответственность
Принадлежность <i>потому что у нас так принято</i>	• командность • поддержка лидера • традиции • социальная забота	регулярные встречи, поддержание фокуса без лишнего давления, убеждения

Большинство производственных компаний — культура силы. Главная задача там — получить поддержку первого лица

Итоговые рекомендации

Риски в ИИ-проектах (кроме специфичных, которых 10–20%) не отличаются от типовых рисков всех ИТ-проектов. Подробнее этот вопрос я разбирал в статье по QR-коду и гиперссылке ниже.



[7 причин большинства проблем на пути цифровизации и цифровой трансформации](#)

А здесь сведу все предыдущие рекомендации в одну таблицу.

Причина	Риск	Мера предупреждения
Нет глобального понимания сути технологии и целей, для чего она используется	Неправильные фокусы, сжатые сроки, неэффективное использование технологий, отсутствие ресурсов и команды	Проведение обучения и информационной кампании для всех участников компании, разработка стратегии
Недостаток необходимых soft и hard компетенций у ИТ	Фокус на технологиях и экономичном внедрении, снижении издержек на сопровождение ИТ, а не эффектах для бизнеса	Подготовка и найм специалистов с необходимыми компетенциями, выделение функции инноваций из блока ИТ, обучение существующего персонала новым навыкам
Низкая культура работы с данными и процессами	Неэффективная автоматизация, невозможность анализа данных и раскрытия эффектов от внедрения	Проведение обучения по работе с данными и усилению культуры аналитики в компании, выстраивание процессного

	ИИ, срывы сроков и превышения бюджетов	управления, автоматизация сбора данных
Неудобство ИТ-систем	Сопротивление персонала, избыточные требования к персоналу	Включение пользователей и администраторов в проектирование интерфейсов, сбор обратной связи об удобстве работы, проектирование комплексных решений
Недостаток компетенций в управлении проектами и изменениями	Непроработанность инициатив и реализация неэффективных проектов, хаотичность внедрения и срывы сроков, превышения бюджетов	Обучение управленческого персонала современным методам управления проектами и изменениями
Отсутствие базовой цифровой грамотности у персонала	Риск неправильного использования цифровых инструментов и систем, сопротивление персонала и саботаж, трудности адаптации персонала	Проведение обучения сотрудников базовым цифровым навыкам, создание руководств и инструкций по работе с цифровыми системами (сценарные инструкции)
Культура компании	Неверный выбор ролевой модели, невозможность внедрения изменений, нереалистичные ожидания	Выбор подходящей модели поведения и внедрения изменений, фокусов в коммуникации, исключение невыполнимых ожиданий

Чек-лист «7 грехов»: проверьте свой проект

1. Есть ли у первых лиц единое понимание, зачем компании ИИ и чего от него ждать?
2. Кто персонально отвечает за развитие ИИ — и есть ли у ИТ нужные компетенции и мотивация?
3. Описаны ли процессы и инвентаризированы ли данные в зоне проекта?

4. Удобно ли решение пользователю — и когда вы в последний раз спрашивали его об этом?
5. Владеет ли команда базовой методологией проектов, продуктов и изменений?
6. Хватает ли людям цифровой грамотности для работы в новой системе?
7. Поддерживает ли культура компании сотрудничество — или каждый «сам за себя»?

Если хотя бы на два вопроса ответ «нет» — сначала закройте эти дыры и только потом масштабируйте ИИ.

Резюме главы

К сожалению, ключевые факторы сейчас играют не в пользу быстрого развития и продвижения ИИ. Компании хотят концентрировать всё внутри себя, получить сразу комплексное и персональное решение, при этом управлять изменениями, обучать людей и формировать культуру принятия изменений не хотят. А уровень цифровой грамотности обычных сотрудников с годами если и растёт, то крайне медленно.

Но большие игроки это понимают. Опять же, если ещё раз посмотреть на OpenAI с их ChatGPT, то они делают ставку на создание доступного решения, магазинов дополнений и продажу по подписке. Их стратегия долгосрочная и они хотят стать вторым Microsoft с собственным Windows. То есть стать для всех именем нарицательным и стандартной платформой, на базе которой будут собираться различные ИИ-решения. А для этого нужно стать привычным для всех инструментом, чтобы уже через голову людей заходить в корпорации, а не бороться с ветряными мельницами корпоративных правил. Нужно ли сейчас получать одобрение на покупку компьютера с Windows или Linux? Или это уже стандарт?

Также, если вы создаёте ИИ-решения, то нужно уходить от чат-ботов в плоскость полноценных систем. И я рекомендую обратить внимание на проекты машинного зрения и анализа данных промышленного Интернета вещей.

Ключевые выводы

- Проекты, изменения и продукты — одна дисциплина: любой проект несёт изменение, а продукт — это поток проектных решений. ИИ снимает рутину (документы, планы, отчёты), но не заменяет мышление и необходимость договариваться, решать конфликты и так далее.
- «7 смертных грехов цифровизации» — главный фильтр главы: большинство провалов — управленческие, а не технологические. Подробнее в книге «ЦТ-1. Погружение».
- ИИ-ассистент в проектах работает как «единый файл проекта»: статусы, риски, решения и коммуникация в одном контуре.

Что с этим делать: Прогоните текущий проект по списку «7 грехов» (неделя). Внедрите ИИ-резюме статусов и встреч в одном проекте (квартал).

Вопрос для рефлексии: Какой из «7 грехов» ваша организация совершает прямо сейчас — и кто в ней имеет право это сказать вслух?

Глава 14. Коммуникация

Влияние ИИ

Коммуникация — одно из тех направлений, куда ИИ внесёт наибольшие изменения.

Он позволит обрабатывать большие объёмы входящей информации, в том числе через подготовку кратких резюме звонков, встреч, писем и т.д.

Например, мы провели часовую встречу с активным обсуждением и в конце приняли решения. И вместо того, чтобы потом сидеть и переслушивать запись, готовить итоги встречи (а всё это может занять 4–5 часов), мы будем это отдавать ИИ, который всё расшифрует и подготовит шаблон протокола за 15–20 минут. Ему нужна лишь качественная запись.

Этот функционал уже внедряют все крупные разработчики решений для видеоконференций. Например, я уже всю использую эту возможность:

все свои встречи расшифровываю и отправляю их в качестве протокола. В том числе был один интересный случай. В ходе проекта мы наняли аналитика для описания запросов бизнеса и формулирования требований для ИТ. Однако, когда пришло время продлевать договор, мы изучили записи его участия во внутренних еженедельных совещаниях и выяснили интересный факт. Оказалось, что он в рамках своей работы общался с теми же людьми и по тем же вопросам, что и другой аналитик, который занимался выстраиванием процессов в целом. То есть запись и расшифровка встреч позволила выявить двойную работу.

Также ИИ будет готовить шаблоны и заготовки ответов на входящую коммуникацию (письма, сообщения и т.д.).

Первое направление здесь — создание информационных пространств, где будет собрана вся коммуникация, например, в рамках одного проекта. Это необходимо для того, чтобы ИИ погружался в контекст и готовил релевантные ответы.

Второе направление, которое здесь можно ожидать, — создание цифровых двойников/ассистентов для каждого человека. Главный вопрос тут — как собирать и агрегировать рабочую информацию и интегрировать её в цифрового двойника человека, чтобы ИИ мог готовить правильные ответы с учётом контекста рабочей задачи и наших личных качеств. И самое главное, из чего и каких показателей будет состоять этот цифровой двойник, проще говоря, что нам надо оцифровать?

Особенно эффективен будет этот инструмент в B2C сегменте для маркетологов, когда нужно упрощать и доносить до аудитории сложные идеи и мысли простым языком. Люди-эксперты любят уходить в академичность, из-за чего целевая аудитория их не понимает. И как показывают эксперименты, ИИ отлично справляется с задачей подготовки простых ответов.

Ещё один вопрос, на который нужно обратить внимание, — мы можем так заиграться в этот инструмент, что 90% коммуникации будет происходить между роботами. Приведу свой личный пример. У меня есть

цифровой робот-ассистент, которой отвечает на звонки. И иногда случается так, что мне звонит робот от компании, ему отвечает мой робот-автоответчик и они между собой общаются, а я потом читаю краткую расшифровку. Почему это ловушка? Потому, что лень — двигатель прогресса: мы начнём делегировать таким инструментам всё больше и больше коммуникации, что приведёт к большому количеству искажений. Помимо этого, есть риски безопасности.

Чтобы не быть голословным, приведу пример, как я использую ИИ для подготовки презентаций и докладов, статей.

- Формирование идеи.

Здесь я в 90% случаев сам формулирую, о чём должен быть доклад, какие тезисы хочу донести.

- Подготовка прототипа.

Далее идёт подготовка прототипа. Я определяю, что хочу сказать на каждом слайде, в какой последовательности должны идти слайды. Да, конечно, я пробовал это делать с помощью ИИ, и даже иногда получалось неплохо, но в большинстве случаев я подключаю ИИ только для проработки отдельных слайдов.

- Отрисовка.

Затем идёт этап отрисовки. Тут я отдаю материалы своему дизайнеру, который при помощи запросов к ИИ генерирует необходимые иллюстрации. Он понимает, как правильно взаимодействовать с моделями. При этом все основные решения по оформлению — не только по графике, но и по тому, как рассказать мои тезисы слушателю (зрителю) — дизайнер принимает сам. ИИ используется только для отдельных элементов, сокращая время на ручной труд дизайнера.

- Написание статьи.

После этого идёт подготовка статьи по докладу. Презентация загружается в ИИ, также я прописываю основные тезисы и важные нюансы, ограничения (например, на объём статьи), и ИИ генерирует

материал. Этот материал имеет готовность 50–70%. Он шлифуется и корректируется.

- Подготовка анонса.

Ну, и финал — подготовка анонса. Статья и презентация загружаются в ИИ, которому даётся задание по подготовке анонса для определённых площадок и социальных сетей.

В итоге получается экономия времени примерно на 50–80%. Конечно, можно всё автоматизировать и экономить вообще 99% времени. Но это будет продукт очень посредственного качества. Поэтому в любом случае остаётся сочетание человека и робота. Тогда это и быстро, и стильно, и остаётся личность автора, его взгляд — это интересно слушать.

Пока книга шла от второй редакции к третьей, эта рекомендация успела устареть в хорошем смысле. ИИ-резюме встреч из «фишки» превратилось в норму: автоматические конспекты и списки поручений встроены уже в большинство корпоративных ВКС-платформ (хотя ещё не во все). Конкурентное преимущество сместилось от самой функции к дисциплине её использования: стандарту фиксации решений и интеграции резюме в трекеры задач.



[Управление коммуникацией. Часть 1](#)

Влияние на ИИ

Как я часто говорю, коммуникация — ключевой компонент успеха в любом деле или проекте. Абсолютно то же самое можно сказать и об

области развития и проникновения ИИ в нашу жизнь. То, насколько эксперты и технические специалисты смогут доступно и просто донести до людей пользу ИИ, убрать страхи перед ним, будет напрямую влиять на скорость проникновения технологии в нашу жизнь.

Основная задача здесь — выбор каналов и форм коммуникации. Нельзя уходить в официальность и сложные формулировки для популяризации, академичные обсуждения. Страх сказать что-то не так, желание быть абсолютно правильным — одна из ключевых ошибок при запуске новых продуктов и продвижении своих услуг. Ведь мы получаем такие описания, которые неспециалист разберёт, только будучи нетрезвым. То есть мы сами порождаем барьеры в коммуникации.

Поэтому основные приёмы:

- научиться доносить идеи простым и доступным для обывателя языком;
- делать фокус на иллюстрации и графики;
- работать со страхами людей потерять работу и делать упор на их эмоции.

В общем, находить евангелистов, популяризаторов и талантливых маркетологов, а технические специалисты и учёные пусть занимаются прикладным решением задач.

Резюме главы

- Коммуникация — направление максимального эффекта ИИ: резюме встреч, писем и звонков экономят 50–80% времени на обработку информации.
- Правило: ИИ готовит черновик и выжимку — человек принимает решения и отвечает за отношения. Гнаться за «99% автоматизации» в коммуникации — терять доверие и качество решений.
- Обратное влияние: структура и дисциплина вашей коммуникации (лично или в компании) напрямую определяют качество работы ИИ на ваших данных.

Что с этим делать: Включите ИИ-резюме для одного типа регулярных встреч (неделя). Введите стандарт фиксации «решение — ответственный — срок» (квартал).

Вопрос для рефлексии: Сколько часов в неделю ваша команда тратит на пересказ уже сказанного?

Глава 15. Цифровые технологии, работа с данными и кибербезопасность

Здесь мы рассмотрим связку ИИ с наиболее популярными ИТ-технологиями, решениями, работой с данными. Вопросу кибербезопасности мы не будем уделять внимание лишь из-за того, что в первой части книги мы уже погрузились в него достаточно подробно.

Цифровые технологии

Каждую из этих технологий я подробно — с плюсами, минусами и личным мнением — разобрал в книге «ЦТ-1. Погружение», Глава 2. Здесь оставляю только главное: как каждая из них дружит с ИИ.

Подробнее с технологиями из этого блока можно познакомиться по QR-коду и гиперссылке ниже.



[Обзор цифровых технологий. Часть 1](#)

Интернет вещей и беспроводная связь

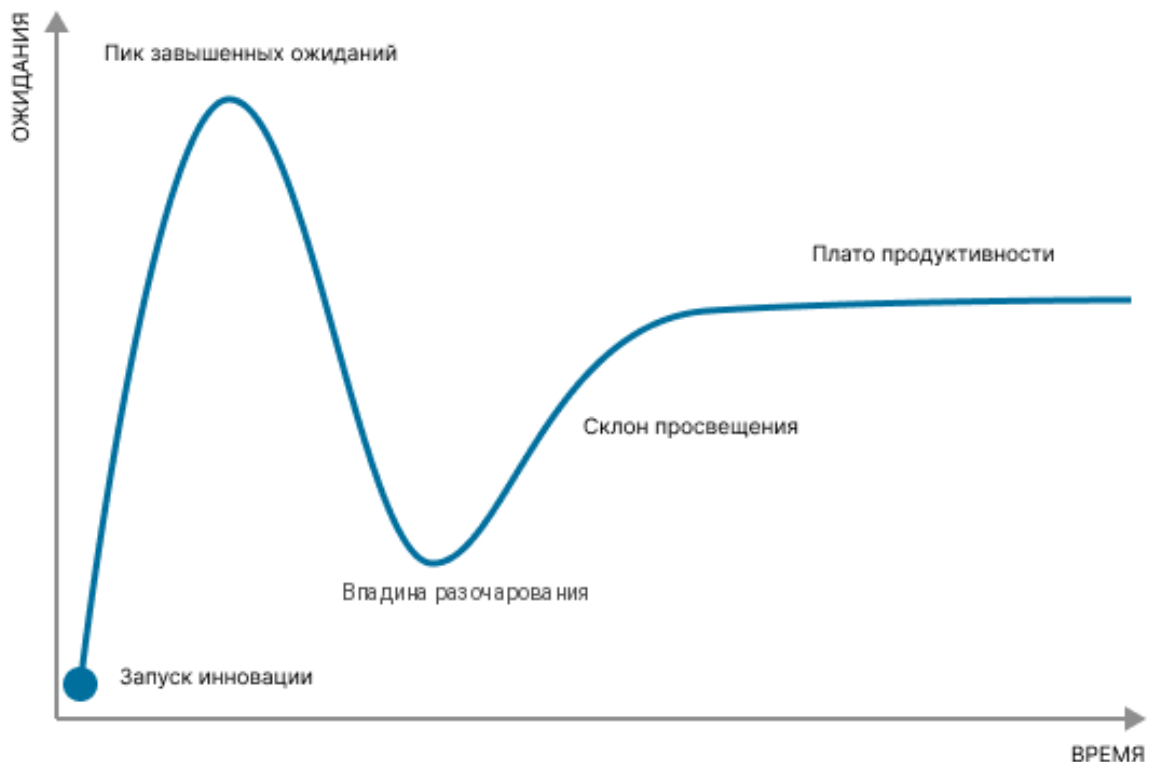
Что это: сеть датчиков и устройств, которые сами собирают и передают данные — от вибрации станка до температуры в цехе; беспроводные сети (5G/6G) — их «нервная система». Бизнесу даёт объективную картину происходящего без ручного ввода.

Без развития этих технологий невозможно и полноценное развитие ИИ-решений.

В первую очередь, ИИ — это предсказательные системы и цифровые советники. А для их обучения и создания цифровых двойников нужны чистые большие данные. Соответственно, встаёт вопрос: «А откуда брать эти чистые большие данные, в которых не будет ошибок из-за человеческого фактора, когда условный Вася случайно записал не 10,00, а 100,00?».

Поэтому развитие технологий Интернета вещей для сбора и передачи данных, в том числе 5G и 6G, — основной фактор для дальнейшего развития ИИ. В противном случае мы останемся наедине с языковыми моделями и вряд ли технология выберется из впадины разочарования.

Хайп-цикл Gartner



Хайп-цикл Gartner

Большие данные (Big Data)

Что это: технологии хранения и обработки массивов, которые не помещаются в «обычную таблицу»: миллионы транзакций, показания датчиков, логи, тексты. Бизнесу дают возможность видеть закономерности, невидимые на малых выборках.

Большие данные — пища для ИИ. Именно в данные упираются все ИИ-проекты. Да, мы можем делать развёртывание локальных моделей и обучать их на отдельных регламентах, но нужно изначально создать и обучить модель, а уже затем «срезать» лишнее. Создание того же GPT-5 также сталкивается с проблемой нехватки качественных больших данных.

Именно поэтому одно из ключевых направлений в развитии ИИ — создание алгоритмов, которым будет нужно меньше данных для обучения.

Цифровые двойники

Что это: виртуальная копия изделия, станка или целого завода, которая живёт на реальных данных и позволяет проигрывать сценарии без риска для производства.

Сама по себе технология мало зависит от ИИ. Там больше нужна классическая математика. Но симбиоз здесь как раз заключается в том, что создание цифровых двойников позволяет смоделировать те данные, которые нужно собирать для развития предиктивной аналитики. То есть подход к созданию новых изделий через создание цифровых двойников позволит ускорить переход на риск-ориентированное техническое обслуживание и внедрить ту самую предиктивную аналитику, к которой стремятся крупнейшие производственные компании.

Кроме того, цифровые двойники станут инструментом для визуализации рекомендаций, сбора и структурирования эксплуатационных данных.

Облака, онлайн-аналитика, удалённое управление

Что это: вычисления и хранение «по подписке» в дата-центрах провайдера: мощности арендуются за часы, а не покупаются месяцами. Бизнесу дают быстрый старт без капитальных вложений.

До широкомасштабного развёртывания нейроморфных решений мы будем привязаны к развёртыванию ИИ-моделей на серверах облаков или внутренних центров обработки данных. Поэтому ИИ и облака будут связаны тесно и надолго.

Требования к вычислительной инфраструктуре для ИИ (облако против on-premises, GPU, экономика эксплуатации) и подготовка данных, включая анонимизацию перед использованием облачных LLM, подробно разобраны в книге «ИИ-2. Практическое руководство» (Главы 10 и 15) — здесь только короткий обзор.

Цифровые платформы

Что это: среда, в которой вокруг ваших данных и сервисов собирается экосистема: сотрудники, клиенты и партнёры работают по единым правилам и в едином информационном поле.

Здесь необходимо сделать одно уточнение. Что такое платформа? Если компания разработала какую-то ИТ-систему, где собраны её контрагенты и внутренние сотрудники, является ли это платформой?

Главное качество платформы — свободный доступ участников к подключению. То есть платформа — это в первую очередь большое количество участников. И ИИ здесь используется в виде рекомендательных систем (цены, проверки контрагентов, описания товаров и их подбор и т.д.).

Самый наглядный пример — такси и службы доставки. Здесь ИИ рассчитывает цены в каждый момент времени, строит маршруты, подбирает исполнителей и т.д. Если ИИ «кривой», то вся система перестанет работать.

Приведу пример из моей практики. Стартап в области логистики последней мили внёс корректировки в свою нейросеть, что привело к огромным убыткам. Система стала назначать заказы так, что за новым заказом курьеру приходилось ехать через весь город, в то время как рядом мог находиться свободный курьер. В итоге складывалась следующая ситуация:

- часть курьеров могла весь день простоять на своей точке, а компания оплачивала минимальную ставку;
- другие ездили пустыми за новым заказом через весь город, а компания оплачивала этот пробег.

Сплошные убытки и неработающая экономическая модель.

Поэтому от правильной работы алгоритмов ИИ будет зависеть успех платформ.

Блокчейн, умные контракты

Что это: распределённый реестр, записи в котором нельзя незаметно подменить; умные контракты автоматически исполняют условия сделки. Бизнесу дают доверие без посредника.

Блокчейн, на мой личный взгляд, — одна из самых переоценённых технологий, которая имеет очень узкую нишу. И влияние ИИ здесь будет наименьшим.

Ключевое направление — контроль использования криптовалют через анализ цепочек транзакций и сбор данных из разных источников. Но задача эта настолько нишевая, что если и получит развитие, то только через финансирование со стороны государств.

Машинное зрение

Что это: распознавание объектов и событий на фото и видео: контроль качества, безопасность, учёт. Самый зрелый прикладной ИИ — часто именно с него разумно начинать.

Машинное зрение — это обработка видеоряда искусственным интеллектом. Сама технология настолько развилась и проникла в быт, что мы уже не замечаем её.

Это и анализ того, надета ли каска на рабочем, и контроль качества изделий на производстве (например, замеры габаритов деталей), и датчики в автомобилях (например, распознавание и уведомление о препятствиях через камеры заднего вида).

Ключевое ограничение здесь, как и во всей индустрии ИИ, — нужно большое количество данных для обучения и качественный видеоряд. По мере развития алгоритмов, снижения требований к количеству данных, качеству видеопотока (нужны будут всё более дешёвые камеры) технология будет всё глубже проникать в производство и обычную жизнь.

Роботизация процессов (RPA)

Что это: программные роботы, повторяющие действия человека в интерфейсах систем: заполнить форму, перенести данные между программами, сформировать отчёт. Бизнесу дают автоматизацию без дорогой перестройки ИТ-ландшафта.

Одно из ключевых конкурентных преимуществ любой RPA-системы — качество распознавания текста, данных и простота настройки. Поэтому

и здесь ИИ будет влиять на всю технологию. Особое внимание — распознавание рукописного текста. Да, когда-то мы придём к полному электронному документообороту, где рукописного текста не будет, но до этого ещё очень далеко.

3D-печать

Что это: послойное производство деталей из цифровой модели — от прототипов до запчастей. Бизнесу даёт скорость и независимость от поставщиков единичных позиций.

В данном направлении, казалось бы, влияние ИИ будет наименьшим. Но и тут есть куда развернуться. А именно — создание 3D-моделей для печати на основе двумерных изображений (фото, чертежи) или вообще текстового или голосового описания. Так что здесь тоже основное направление для развития — распознавание образов и их конвертация.

Озёра и хранилища данных

Что это: централизованные хранилища «сырых» (озеро) и подготовленных (хранилище) данных из всех систем компании — фундамент, на котором стоят аналитика и ИИ.

Управление и проектирование ИТ-инфраструктуры, повышение качества данных и аналитика — самое сложное, но и одно из самых перспективных направлений.

Если немного уточнить, то получим следующий набор задач для ИИ.

- Снижение сроков и затрат на проектирование систем сбора и хранения данных через рекомендации на основе информации о типе и архитектуре хранилища, о способах сбора данных.

Например, есть задача — анализировать видеопоток с удалённых объектов на предмет проникновения нарушителей. Как подойти к этой задаче? Ставить умные камеры, чтобы в базу шли и записывались только определённые события, или же поставить обычные камеры, но всё заснятое транслировать в центр обработки данных, где потоки проанализируются мощным ИИ? Что для нас важнее, стоимость камер

или организация стабильного канала передачи данных? И таких задач для ИИ масса.

- Анализ данных и поиск аномалий в качестве данных.

Организация процессов обеспечения качества данных — ключевая задача для любой компании на пути цифровизации. Особенно в контексте развития Data-Driven-подхода (принятие решений на основе данных), когда от данных зависят управленческие решения, создание новых продуктов и так далее.

В ИТ есть даже отдельные люди — дата-стюарды, которые изучают данные и обеспечивают их достоверность, устраняют сбои и ошибки. Так вот ИИ может помочь в этом направлении, сделать его доступным не только гигантам и банкам, но и обычным компаниям. В основном как помощник для опытных специалистов, повысив их производительность и эффективность. А если мы ещё настроим взаимодействие с BPM-системами (системы управления бизнес-процессами), то сможем получать рекомендации по оптимизации бизнес-процессов. То есть ИИ может стать крутым инструментом для ИТ-департаментов.

Резюме

Да, это не полный список технологий, сюда многие относят ещё различные дроны, роботы, сенсоры и т.д. Но, по-моему, все они являются производными от технологий выше. Ведь основа их работы — это машинное зрение, обработка больших данных от сенсоров и т.д. Поэтому этот список можно продолжать до бесконечности, но я остановился на таком наборе.

ИТ-системы

ИТ-ландшафт и слой ИИ

ИИ не заменяет учётные системы — он работает поверх них и питается их данными

ИИ поверх ландшафта: прогнозы · аномалии · генерация · рекомендации



ИИ силен ровно настолько, насколько чисты данные под ним

ИТ-ландшафт и слой ИИ поверх него

Ниже — карта основных классов корпоративных ИТ-систем через призму ИИ: что это за класс и что он даёт бизнесу, какие задачи в нём сможет решать ИИ и какие условия для этого придётся соблюсти. Сами системы как ИТ-продукты — с примерами и нюансами выбора — разобраны в книге «ЦТ-1. Погружение» (Глава 3).

ERP

Что это и что даёт бизнесу: Единый контур учёта и планирования ресурсов — финансы, закупки, склад, производство. Бизнесу даёт «источник правды» по операциям и деньгам.

Ключевые возможности ИИ: Прогноз спроса и денежного потока, поиск аномалий в проводках и закупках, рекомендации по запасам, черновики документов и типовых операций.

Условия: Дисциплина ввода данных, чистые справочники (связка с MDM), настроенные процессы — иначе ИИ будет прогнозировать по мусору.

MES / APS

Что это и что даёт бизнесу: Оперативное управление производством и детальное планирование: сменные задания, маршруты, загрузка мощностей. Даёт прозрачность цеха и ритмичность выпуска.

Ключевые возможности ИИ: Оптимизация расписаний с учётом реальных ограничений, предсказание сбоев и переналадок, мгновенное перепланирование при отклонениях, подсказки мастеру смены.

Условия: Оцифрованные маршруты и нормативы, поток данных с оборудования (IoT/SCADA), актуальные остатки и приоритеты заказов.

PLM / PDM

Что это и что даёт бизнесу: Конструкторско-технологические данные и жизненный цикл изделия — от идеи до снятия с производства.

Ключевые возможности ИИ: Поиск аналогов деталей и узлов (меньше дублирования разработки), генеративный дизайн, проверка требований и комплектности, ускорение согласований изменений.

Условия: Наполненная база изделий, единая классификация, работающий регламент управления изменениями.

CRM

Что это и что даёт бизнесу: Клиенты, сделки, история коммуникаций; управление воронкой продаж. Даёт управляемость выручки.

Ключевые возможности ИИ: Скоринг лидов и прогноз оттока, персональные предложения, резюме звонков и переписки, подсказка следующего шага по сделке, черновики писем.

Условия: Полнота ведения сделок менеджерами (главная боль любого CRM), интеграция каналов, накопленная история хотя бы за 6–12 месяцев.

SCADA

Что это и что даёт бизнесу: Телеметрия и диспетчеризация оборудования в реальном времени. Даёт контроль технологического процесса.

Ключевые возможности ИИ: Выявление аномалий раньше оператора, предиктивные предупреждения, помощь в разборе инцидентов и поиске первопричин.

Условия: Достаточное покрытие датчиками, стабильный сбор данных, разметка нормальных и аварийных режимов, защищённый ИБ-контур.

EAM

Что это и что даёт бизнесу: Управление активами и ТОиР: паспорта оборудования, графики, наряды. Даёт надёжность и управляемую стоимость владения.

Ключевые возможности ИИ: Предиктивное обслуживание вместо календарного, приоритизация ремонтов по риску и критичности, прогноз потребности в запчастях.

Условия: История отказов и наработок, связка со SCADA/IoT, дисциплина закрытия нарядов с фактическими данными.

BIM

Что это и что даёт бизнесу: Информационная модель здания на всём жизненном цикле — от проекта до эксплуатации.

Ключевые возможности ИИ: Проверка коллизий, оптимизация проектных решений и смет, сценарии эксплуатации и энергоэффективности.

Условия: Актуальность модели после сдачи объекта, стандарты моделирования, поступление данных эксплуатации в модель.

MDM

Что это и что даёт бизнесу: Единые справочники и нормативно-справочная информация: контрагенты, номенклатура, персонал. Даёт «один язык» всем системам.

Ключевые возможности ИИ: Именно здесь ИИ поддерживает «гигиену» данных: находит дубли и синонимы, сверяет и заполняет атрибуты,

предлагает эталонную запись, вычищает справочники при слияниях систем.

Условия: Назначенный владелец данных, регламент ведения справочников, интеграция с системами-источниками.

WMS

Что это и что даёт бизнесу: Управление складом: топология, задания, комплектация. Даёт скорость и точность склада.

Ключевые возможности ИИ: Оптимизация размещения и маршрутов отбора, прогноз пиковых нагрузок, контроль ошибок комплектации.

Условия: Адресное хранение, штрихкодирование или RFID, точные остатки в системе.

TMS

Что это и что даёт бизнесу: Управление транспортом и доставкой: маршруты, тендеры перевозчиков, трекинг. Даёт предсказуемую логистику и её стоимость.

Ключевые возможности ИИ: Оптимизация маршрутов и загрузки транспорта, прогноз сроков и опозданий, выбор перевозчика по цене и надёжности.

Условия: GPS-данные и трекинг, интеграция с заказами и WMS, тарифы и история перевозок в системе.

BPM

Что это и что даёт бизнесу: Моделирование и исполнение бизнес-процессов. Даёт управляемость и повторяемость операций.

Ключевые возможности ИИ: Process mining: восстановление реального процесса по цифровым следам, поиск узких мест и отклонений от регламента, рекомендации и симуляция изменений.

Условия: Процессы должны исполняться в ИТ-системах (оставлять следы в логах), назначены владельцы процессов.

SCM / SRM

Что это и что даёт бизнесу: Планирование цепочки поставок и работа с поставщиками. Даёт устойчивость снабжения.

Ключевые возможности ИИ: Прогноз спроса и сроков, риск-профили поставщиков, сценарии «что если» по всей цепочке.

Условия: Данные продаж и запасов по всей цепочке, история поставок, качественные мастер-данные.

СРМ / ЕРМ

Что это и что даёт бизнесу: Контур финансового управления: бюджетирование, консолидация, план-факт. Даёт управляемость экономикой компании.

Ключевые возможности ИИ: Драйверные прогнозы бюджета, объяснение отклонений на естественном языке, сценарное моделирование «что будет, если».

Условия: Единая методология бюджетирования, связка с ERP и BI, одна «версия правды» по факту.

BI

Что это и что даёт бизнесу: Витрины данных, дашборды, отчётность. Даёт скорость и обоснованность решений.

Ключевые возможности ИИ: Интерфейс «спроси у данных» на естественном языке, автогенерация выводов к дашбордам, поиск причин отклонений.

Условия: Витрины с понятной моделью данных, словарь показателей (одинаковое понимание метрик), настроенные права доступа.

СППР / цифровые советники

Что это и что даёт бизнесу: Системы поддержки принятия решений — кульминация всей связки «данные → рекомендации» (подробно — Глава 9).

Ключевые возможности ИИ: Варианты решений с оценкой рисков и последствий, персонализация под роль руководителя.

Условия: Все условия этой главы сразу: данные, процессы, интеграции — и доверие пользователей (см. Главу 9).

ИБ-системы (IDM, DLP, XDR)

Что это и что даёт бизнесу: Управление доступом, защита от утечек, обнаружение атак. Дают устойчивость бизнеса к киберрискам.

Ключевые возможности ИИ: Поведенческая аналитика пользователей, корреляция событий, приоритизация инцидентов; помните — ИИ работает и на стороне атакующих (Глава 8).

Условия: Покрытие инфраструктуры логами, процессы реагирования (SOC), актуальная модель угроз.

ЭДО

Что это и что даёт бизнесу: Юридически значимый электронный документооборот. Даёт скорость и прослеживаемость документов.

Ключевые возможности ИИ: Выжимки и проверка документов, извлечение реквизитов, умная маршрутизация, контроль сроков.

Условия: Документы в цифровом виде (сканы — через распознавание), настроенные маршруты, справочник контрагентов (MDM).

ITSM / Service Desk

Что это и что даёт бизнесу: Управление ИТ-услугами: инциденты, запросы, изменения. Даёт управляемость ИТ и качество сервиса.

Ключевые возможности ИИ: Чат-бот первой линии, автоклассификация и маршрутизация заявок, подсказки решений из базы знаний, прогноз всплесков обращений (экономiku первой линии поддержки разберём в Главе 20).

Условия: Наполненная база знаний, категоризация заявок, дисциплина закрытия с описанием решения.

Контакт-центр (CCaaS)

Что это и что даёт бизнесу: Омниканальная работа с обращениями клиентов: телефония, чаты, мессенджеры. Даёт качество клиентского сервиса в масштабе.

Ключевые возможности ИИ: Боты первой линии и суфлёры операторов, резюме звонков, оценка качества 100% диалогов вместо выборочной, маршрутизация по темам и настроению.

Условия: Интеграция каналов с CRM, записи и транскрибация разговоров, человек в контуре для сложных случаев: по исследованию Sinch (май 2026) 74% компаний откатали уже запущенных ИИ-агентов в поддержке — ставка на полную автономию не сработала (подробнее в Постскриптуме).

CDP

Что это и что даёт бизнесу: Единый профиль клиента, собранный из всех каналов и систем. Даёт маркетингу точность вместо «ковровых» рассылок.

Ключевые возможности ИИ: Сегментация, next best offer, прогноз LTV и оттока, персонализация кампаний.

Условия: Законные согласия на обработку данных, сшивка идентификаторов между системами, качество событий.

HRM

Что это и что даёт бизнесу: Кадровый контур: подбор, кадровый учёт, развитие. Даёт скорость найма и удержание людей.

Ключевые возможности ИИ: Отбор релевантных откликов, черновики вакансий и офферов, адаптационные боты, прогноз текучести.

Условия: Структурированные профили должностей, проверка алгоритмов на предвзятость, работа с персональными данными строго по закону.

LMS

Что это и что даёт бизнесу: Корпоративное обучение: курсы, тесты, траектории развития. Даёт масштабируемую передачу знаний.

Ключевые возможности ИИ: Генерация курсов и тестов из ваших регламентов, персональные траектории, диалоговые тренажёры для отработки навыков.

Условия: Оцифрованные регламенты и база знаний, связка с профилями HRM, контроль качества контента экспертом.

PIM

Что это и что даёт бизнесу: Управление товарным контентом для каталогов и маркетплейсов. Даёт качество карточек в масштабе.

Ключевые возможности ИИ: Генерация и локализация описаний, обогащение карточек, контроль полноты атрибутов.

Условия: Эталонные данные о товаре (связка с MDM), редакционная политика, утверждение человеком.

BMS (Board Management Software)

Что это и что даёт бизнесу: Контур органов управления: заседания советов и комитетов, материалы, поручения. Даёт дисциплину корпоративного управления.

Ключевые возможности ИИ: Подготовка и резюме материалов, протоколы заседаний, контроль исполнения поручений.

Условия: Конфиденциальность прежде всего — локальный контур или доверенное облако; регламент документооборота руководства.

Работа с данными

Влияние ИИ на работу с данными мы обсудили и в блоке озёр и хранилищ данных, и в контексте работы с другими цифровыми технологиями и информационными системами. Отдельно отмечу лишь то, что направление управления данными будет приобретать всё более важную роль. Те компании, которые смогут решить вопрос актуальных и качественных данных, выстроить инфраструктуру, процессы и сформировать компетенции и мотивацию у людей, смогут наиболее полно использовать потенциал искусственного интеллекта. Поэтому одна из моих главных рекомендаций — начинайте описывать потоки данных, внедрять культуру управления данными, назначайте владельцев данных, внедряйте инфраструктуру и привлекайте профессионалов. Управление данными в контексте развития искусственного интеллекта будет главным преимуществом.

Резюме главы

- ИИ силён в связке: IoT, Big Data, ИТ-системы. Данные и инфраструктура определяют потолок пользы ИИ — модель не может быть лучше того, что в неё поступает.
- Работа с данными — условие № 1: «мусорные» данные обнуляют любой алгоритм; исправление на этапе эксплуатации стоит на порядки дороже, чем наведение дисциплины на входе.
- Кибербезопасность — сквозная тема: каждая новая система и интеграция — новая поверхность атаки (подробно — Глава 8).

Что с этим делать: Проведите аудит источников данных для первого ИИ-кейса (неделя). Составьте дорожную карту качества данных: владельцы, форматы, контроль (квартал).

Вопрос для рефлексии: Каким данным в вашей управленческой отчётности вы доверяете безоговорочно — и проверяли ли вы это доверие?

Глава 16. Практики регулярного менеджмента

К практикам регулярного менеджмента (ПРМ) относят набор регулярных ритуалов, позволяющий грамотно использовать не только своё рабочее время, но и потенциал каждого работника в команде. В некоторых компаниях практики регулярного менеджмента называют «стандартными практиками руководителя».

В классической теории в ПРМ включают:

- управленческое планирование;
- делегирование полномочий и задач;
- работу с обратной связью подчинённых;
- организацию совещаний;
- работу с кадрами и решения относительно людей;
- оценку эффективности работы сотрудников;
- развитие кадров.

Ряд из этих вопросов мы уже рассмотрели ранее. Поэтому давайте пройдем кратко.

Управленческое планирование

ИИ здесь будет выступать в качестве компонента СППР. Он будет анализировать данные, выявлять аномалии и готовить прогнозы с рекомендациями.

Делегирование полномочий и задач

ИИ поможет автоматизировать создание карточек задач из сообщения — на канбан-доске или в трекере задач. То есть выступит в роли робота, который может распознать даже голосовое сообщение, выявить суть, определить адресата и сроки, а затем всё автоматически перенесёт. В итоге это упростит внедрение инструментов для отслеживания задач. Ключевое — дисциплина руководителей. Если руководитель не может чётко сформулировать задачу, будет экономить буквы или слова, писать без указаний имени и так далее, то ИИ не сможет помочь.

Также в проектах ИИ может выступить ассистентом руководителя, который поможет определить необходимые полномочия для реализации задач. А при необходимости поможет подготовить все необходимые формальные документы: приказ, задание и так далее.

Работа с обратной связью подчинённых

В данном блоке основное направление — проведение опросов, «оценок 360» и анализ массивов данных, выявление общих закономерностей и трендов для руководителей с широким контуром прямого подчинения — условно от десяти человек. А также подготовка персонализированных рекомендаций для руководителей.

Организация совещаний

Мы также рассматривали этот блок в разделе про решения класса Board Management Software (Глава 15):

- подготовка и рассылка повестки совещания, в том числе на основе предыдущих договорённостей и отчётов об исполнении заданий;

- модерация совещаний, их запись и расшифровка аудио в текстовую версию с разделением по спикерам;
- подготовка протоколов встреч, их согласование и создание задач в сервисах отслеживания с назначением ответственных и сроков.

Работа с кадрами и решения относительно людей, оценка эффективности работы сотрудников, развитие кадров

Здесь применимо всё из блока про HRM-системы (Глава 15). Также полезна интеграция с системами выполнения задач и управления проектами. Это позволит ИИ-модели собирать больше данных и создавать более качественные рекомендации.

Однако это направление во многих юрисдикциях будет под пристальным контролем.

Три сценария «ПРМ × ИИ»: с чего начать

ПРМ держатся на ритме: планёрка, встречи один на один, обратная связь, контроль поручений. ИИ не меняет сами ритуалы — он убирает их подготовительную стоимость, из-за которой ритуалы обычно и умирают. Ниже — три сценария, с которых проще всего начать; контроль поручений мы уже разобрали выше, в блоке про делегирование.

Планёрка. Ассистент до встречи собирает статусы, риски и хвосты поручений в одну сводку. Встреча начинается не с «рассказывайте, что у вас», а с решений по отклонениям. По моему опыту, это экономит 15–20 минут каждой планёрки — и, что важнее, возвращает ей смысл.

Один на один. Перед встречей ассистент готовит карточку: прошлые договорённости, что сделано, где человек буксует, что его мотивирует. Подготовка сжимается с получаса до пяти минут, а разговор становится про человека, а не про пересказ статусов.

Обратная связь. Ассистент помогает собрать факты и проверить формулировки по структуре «факт → эффект → ожидание», без оценок личности. По моим наблюдениям, половина конфликтов в командах

возникает не из-за содержания обратной связи, а из-за её формы — и ровно эту часть ИИ снимает.

И заметьте: всё это — командная проекция личного контура, который мы разберём в Главе 21. ПРМ, применённые к самому себе, — это дневник, планирование и разбор собственных решений. Начните с себя — команде будет проще поверить.

Резюме главы

- Практики регулярного менеджмента — ритуалы, в которые ИИ встраивается естественнее всего: планирование, статусы, канбан, контроль поручений. ИИ здесь — «робот-ассистент» рутины.
- Эффект возникает только там, где ритуалы уже существуют: ИИ усиливает дисциплину, но не создаёт её.

Что с этим делать: Автоматизируйте создание карточек задач из сообщений и встреч (неделя). Свяжите ИИ-ассистента с трекером команды и ритуалом статусов (квартал).

Вопрос для рефлексии: Какие управленческие ритуалы в вашей команде держатся на личной дисциплине одного человека?

Глава 17. Бережливое производство и теория ограничений систем (ТОС)

Бережливое производство — один из ключевых инструментов для бизнеса. Для тех, кто нашёл свою нишу, он позволяет устранять хаос, минимизировать бюрократию, убрать лишнюю волокиту и стать клиентоориентированными. Для молодых компаний бережливое производство позволяет отсечь всё лишнее, добавляя в продукты и услуги только то, что нужно клиенту, при этом вместе с ростом сохранять свою гибкость.

Подробнее об этом инструменте и как он работает по QR и ссылке ниже.



[Бережливое производство. Часть 1](#)

Компании-лидеры бережливого производства, например, Тойота, уже пробуют использовать ИИ.

То же самое можно сказать и о теории ограничений систем. Это тоже один из главных инструментов для пошагового роста и минимизации кризисов и рисков для компании. Подробнее об этом инструменте и как он работает по QR и ссылке ниже.



[Теория ограничений систем](#)

Влияние ИИ

ИИ в области бережливого производства и работы с ограничениями позволит:

- выявлять потери в действиях людей на производстве и в ИТ-системах (видеоаналитика с производства, анализ поведения пользователей в ИТ-системах);
- выявлять потери в процессах (анализ процессов, в т.ч. ИТ-систем);
- анализировать потери в производственных потоках, выявлять ограничения;
- готовить проактивные рекомендации при обнаружении отклонений или по запросу при планировании масштабирования, запуска новых сервисов и услуг;
- готовить план внедрения изменений с учётом инструментов управления проектами и изменениями (выбирать правильные подходы к реализации проектов, формировать планы, с учётом психологии сотрудников и культуры компании).

Также ИИ вместе с цифровыми инструментами поможет при организации внедрения (разработка уставов проектов, планов, сопровождение) основных инструментов:

- Kaizen — анализ, классификация и приоритизация предложений по улучшению, выявление общих паттернов;
- TQC (всеобщий контроль качества);
- TQM (всеобщее управление качеством);
- TPM (всеобщий уход за оборудованием);
- Just-in-time (точно вовремя) — планирование производственных задач;
- Отчёт А3 — проведение решения проблемы по логике отчёта и его автоматическое формирование;
- 7 шагов практического решения проблем — чат-боты, которые выступают в роли ассистентов и уточняющими вопросами локализуют проблему, обвязывают её данными, находят причины,

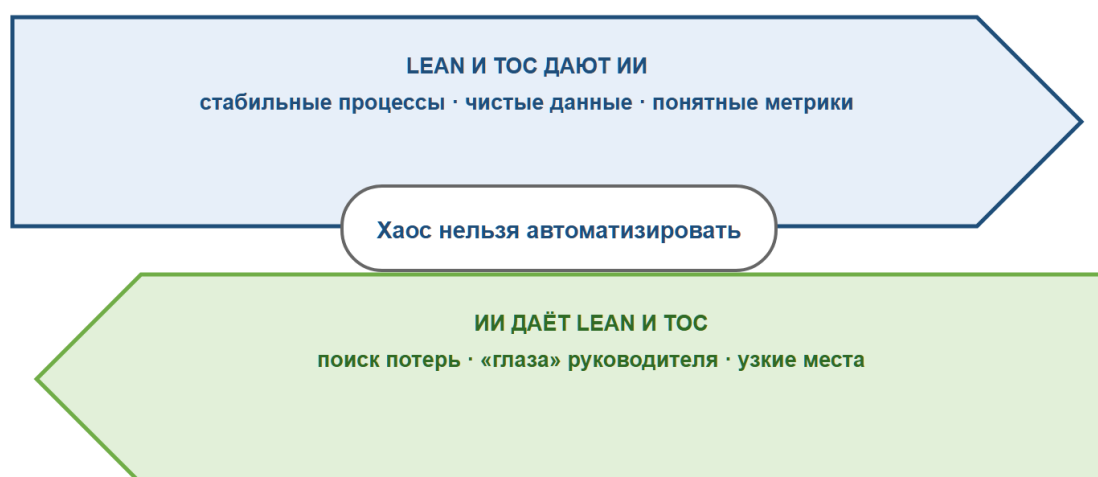
классифицируют и помогают найти адресата на основе описанной орг. структуры;

- VSM (Value Stream Mapping) — описание потоков ценности на основе анализа бизнес-процессов или проведения интервью;
- Рока-yoke — анализ моделей в PLM-системах и подготовка рекомендаций;
- 6 сигм — анализ данных в учётных системах и определение границ управляемости процессов, подготовка рекомендаций, на какие инциденты нужно реагировать, а какие являются форс-мажором;
- DMAIC, SIPOC, DMADV — описание процессов в качестве ассистента на основе проведённых интервью.

Самое сложное здесь — построить модели и научиться быстро и дёшево собирать данные в организациях. Иначе это будет удел больших корпораций.

Взаимное влияние

Взаимное влияние Lean, ТОС и ИИ



Давайте посмотрим на этот вопрос через призму основных ДАО и видов потерь.

14 ДАО

Раздел I. Философия долгосрочной перспективы.

Принцип 1. Принимай управленческие решения с учётом долгосрочной перспективы, даже если это наносит ущерб краткосрочным финансовым целям.

Его суть — в фокусировке на долгосрочном развитии, а не на операционных показателях. И, к сожалению, именно он чаще всего игнорируется, потому что это психологическая ловушка.

Что, как правило, собственник или акционеры хотят от топ-менеджмента? Правильно, красивые показатели выручки, операционной маржинальности, EBITDA, прибыльности.

Это в целом главный конфликт бережливого производства и реальности. То же самое касается и ИИ. Но компании с сильной культурой смогут и будут делать ИИ-проекты. Они понимают, что это будущее.

Раздел II. Правильный процесс даёт правильные результаты.

Принцип 2. Процесс в виде непрерывного потока способствует выявлению проблем.

Его суть в том, чтобы организовать рабочий и технологический процесс с минимальными потерями, запасами и незавершённым производством. Такой подход в сочетании с отлаженной коммуникацией позволяет выявлять проблемы на ранней стадии, пока они не превратились в опасность или кризис. В моей практике это вообще главное правило, в том числе в реализации проектов — собирать обратную связь и выявлять проблемы как можно раньше.

ИИ будет крайне полезен, когда сможет на основе данных из ИТ-систем и опросов делать визуализации производственных потоков и готовить рекомендации. Например, сразу сформирует канбан-доску с нужными столбцами, исходя из специфики компании и внутренних процессов.

Принцип 3. Используй систему вытягивания, чтобы избежать перепроизводства.

Знакома ситуация, когда вам впихивают работу другого подразделения? Я часто вижу принцип, когда разные подразделения работают сами по себе, что ведёт к затариванию складов сырья или захламлению папок с входящими служебками / письмами / задачами.

Это самая частая проблема у руководителей отделов в работе с подчинёнными. Многие такие руководители, по сути, являются «передастами» — они только перекидывают задачи свыше исполнителю. У сотрудника накапливается по 40–50 неотработанных задач, он пытается сделать всё, скачет с задачи на задачу, и в итоге не делает ничего. А как показывает практика, если ограничивать количество рабочих задач до трёх, то производительность вырастает до 50%. Во-первых, человек не перепрыгивает с одной задачи на другую (процесс погружения мозга в задачу занимает 20–40 минут), во-вторых, он имеет хоть какое-то разнообразие и вариативность, это и защищает психику, и позволяет не зависнуть на одной проблеме.

Необходимо сделать так, чтобы внутренний потребитель, который принимает твою работу, получил то, что ему требуется, в нужное время и в нужном количестве. Это лежит в основе инструмента «Точно вовремя» — новые изделия или задачи должны поступать только по мере выполнения предыдущих.

В общем, необходимо свести к минимуму незавершённое производство (актуально и для офисов).

Здесь синергия с ИИ будет в личном планировании и системах рекомендаций для сотрудников. Например, руководитель ставит задачи, но сотрудник их получает только по итогам завершённой работы. А система автоматически определяет, какую из актуальных задач назначить сотруднику, или в целом кому назначить из сотрудников отдела (с учётом сильных личных качеств).

По сути, это агрегатор наподобие платформы для такси, где ИИ понимает личные качества каждого и занимается диспетчеризацией задач. Такие решения отлично будут укладываться в концепцию бережливого производства.

Принцип 4. Распределяй объём работ равномерно (хейдзунка): работай как черепаха, а не как заяц.

Знакома ситуация геройства и пословицы «то пусто, то густо»? В японской философии это называется «Мура». Одна из задач бережливого производства — сделать загрузку равномерной, без авралов в конце месяца и бесконечных геройств.

Здесь ИИ будет частью систем планирования. Это относится и к производству, и к офисным задачам. Он будет накапливать статистику по задачам, датам возникновения и срокам. В итоге заранее будут формироваться задачи и напоминания, исходя из опыта и специфики компании и подразделения.

Принцип 5. Сделай остановку производства с целью решения проблем частью производственной культуры, если того требует качество.

За что так ценились, да и ценятся машины Тойота? Правильно, за качество. Если мы хотим работать с клиентами, у которых есть деньги, то нужно обеспечивать качество.

Если мы говорим про производство, то инструмент Дзидока (оборудование с машинным зрением) — фундамент для «встраивания» качества. Эти проекты уже давно становятся нормой, от контроля качества на отдельных этапах и самого процесса до контроля качества готовой продукции.

В офисную работу также возможно встраивание таких решений. Например, рекомендации ИИ о возможности изменить статус проекта, о том, что он проработан достаточно. Или анализ документов, маркетинговых материалов и так далее.

Принцип 6. Стандартные задачи — основа непрерывного совершенствования и делегирования полномочий сотрудникам.

Каждый раз, когда я слышу тезис «у нас нет стандартных задач», как правило, это означает, что в команде просто не умеют или не хотят структурировать работу.

Стабильные и воспроизводимые методы работы — залог процессного управления. А его задача — создать более предсказуемый результат, повысить слаженность работы, а выход продукции сделать более равномерным. Это основа потока и вытягивания.

Да, если у вас совсем молодая компания, то какое-то время это будет неактуальным, надо запустить всю систему, дать результат, но в итоге без этого никуда. Бесконечное творчество всегда приводит к кризису.

Необходимо фиксировать накопленные знания, например, проводить обзор проектов, выявлять и стандартизировать лучшие подходы и методы. При этом надо поощрять за совершенствование этих стандартов. Вообще, постоянное совершенствование — основное кредо бережливого производства. Но чтобы что-то улучшать, это надо сначала описать и зафиксировать. Иначе будет хаос со всеми вытекающими последствиями. И для ИТ это тоже более чем актуально. Творческие ребята очень не любят стандартизировать свою работу, но увы, надо через «не хочу».

И здесь опять идёт отличная синергия с ИИ. Он сможет выявлять входящие разные задачи, анализировать, стандартизировать, декомпозировать и назначать разным сотрудникам или подразделениям. А также на основе накопленного опыта формировать рекомендации по оптимизации распределения полномочий, процессов и так далее.

Принцип 7. Используй визуальный контроль, чтобы ни одна проблема не осталась незамеченной.

Этот принцип, на самом деле, является составной частью 5-го принципа. Я использую цветные метки. Но это требует дисциплины, а она идёт от руководителя.

При этом, по возможности, необходимо сокращать объём отчётов до одного листа, даже если речь идёт о важнейших финансовых решениях. Для этого есть даже инструмент «Отчёт А3».

Тут ИИ будет как раз компилировать тексты, таблицы, записки и так далее в итоговые решения и визуализации, выделять ключевые проблемы и предложения.

Принцип 8. Используйте только надёжную, испытанную технологию.

Технологии должны помогать людям, а не заменять их. Если мы говорим про цифровизацию и автоматизацию (а ИИ — это именно цифровизация и автоматизация), то зачастую лучше сначала выполнить весь процесс вручную, а уже потом внедрять решения.

Во-первых, вы увидите наиболее проблемные места. Во-вторых, избежите иллюзии, что всё можно автоматизировать и сможете определить границы применения ИИ.

Кроме того, новые технологии часто ненадёжны и с трудом поддаются стандартизации, а это ставит под угрозу поток. Тот же искусственный интеллект, несмотря на всё своё развитие, ещё допускает порой довольно глупые ошибки.

Поэтому вместо непроверенной технологии лучше использовать известный, отработанный процесс. При этом всё же надо поощрять исследование новых технологий и путей. Ведь всё отработанное когда-то было новым.

Для определения готовности компании к новым технологиям и самих технологий к использованию можно использовать такие метрики, как TRL, MRL и CRL.

Если вкратце, то:

- TRL (Technology Readiness Level) — это степень технологической зрелости решения. Всего методика выделяет девять таких степеней: самая первая (TRL-1) — когда только сформулирована концепция технологии и обоснована её польза, а последняя (TRL-9) —

программный продукт готов и удовлетворяет всем предъявленным к нему требованиям.

- MRL (Manufacturing Readiness Level) — методика определения того, насколько ваше производство / процессы / бизнес готовы к постановке на поточное производство новой технологии / продукта. Используется Министерством обороны США, а Счётная палата правительства США назвала это лучшей практикой для улучшения результатов закупок. MRL призвана компенсировать недостатки TRL и выделяет десять степеней, где первая (MRL-1) — когда выявлены основные возможности для производства решения, а последняя (MRL-10) — когда уже налажено полномасштабное производство. И если оценивать сферу ИИ по MRL, то можно увидеть, что высокий уровень производственной зрелости ИИ-проекты ещё не достигли, вся отрасль представлена не зрелыми решениями, а прототипами и ранними версиями.
- CRL (Commercialization Readiness Level) — методика определения уровня рыночной готовности и коммерциализации. Первый уровень (CRL-1) — когда проанализированы тренды, обзоры, темы конференций, динамика патентования и сделан вывод о потребности рынка в новом решении. Последний, девятый уровень (CRL-9) — когда продукт выведен на рынок и уже собираются данные для его дальнейшего улучшения.

Оригинал и полное описание трёх метрик доступны по QR-коду и гиперссылке.



[Уровни TRL / MRL / CRL](#)

Таким образом, все три метрики взаимосвязаны: TRL отвечает за работу над технологией, MRL — за постановку технологии на производство, а CRL — за проработку продуктового предложения и маркетингового продвижения.

Можем ли мы найти хоть одно ИИ-решение на рынке, которое было бы зрелым по всем этим направлениям? Соответственно, и в рамках восьмого ДАО рано говорить о внедрении ИИ-решений в крупные и критичные производства.

Это ещё один из моих аргументов к тезису, что разработчикам ИИ-решений надо смотреть в сторону малого и среднего бизнеса, а не корпораций.

Раздел III. Добавляй ценность организации, развивая своих сотрудников и партнёров.

Принцип 9. Воспитывай лидеров, которые досконально знают своё дело, исповедуют философию компании и могут научить этому других.

Лучше воспитывать своих лидеров, чем покупать их за пределами компании. Во-первых, это одно из ключевых правил мотивации. Когда вы постоянно нанимаете людей со стороны, то демотивируете своих людей. Они начинают терять веру в будущее. В итоге они фокусируются на своих целях, и ожидать от них полной отдачи не приходится. Во-вторых, лидеры, которые знают своё дело, по умолчанию обладают

высокой экспертизой и знают все подводные камни вашей деятельности. А в цифровизации и автоматизации на этом держится многое.

Здесь ИИ поможет собирать и структурировать обратную связь от персонала об изменениях и помогать в обучении персонала основам бережливого производства, проектного управления, цифровизации.

Кроме того, лидер должен не только выполнять поставленные перед ним задачи и иметь навыки общения с людьми. ИИ может стать инструментом, который помогает писать письма: формулировать мысли, проверять орфографию, выбирать наиболее оптимальный канал (почта, звонок, личная встреча, служебная записка).

Принцип 10. Воспитывай незаурядных людей и формируй команды, исповедующие философию компании.

Здесь ИИ поможет в отборе членов команды и их обучении.

Принцип 11. Уважай своих партнёров и поставщиков, ставь перед ними трудные задачи и помогай им совершенствоваться.

Вы можете быть сколь угодно цифровыми, но если ваши партнёры живут всё ещё в бумажной эре, то эффект будет ограниченным. Тут как в теории ограничений систем — прочность цепи определяется самым слабым звеном.

Необходимо создавать партнёрам условия для роста и помогать им, демонстрировать свою высокую эффективность.

Например, ИИ сможет проводить первичный аудит партнёров на основе тестовых опросов или анализа их бизнес-процессов, контролировать их расследования рекламаций, анализировать публичную информацию о партнёрах и формировать рекомендации.

Раздел IV. Постоянное решение фундаментальных проблем стимулирует непрерывное обучение.

Принцип 12. Чтобы разобраться в ситуации, надо увидеть всё своими глазами (генти генбуцу).

Как же часто я видел истории, когда топ-менеджеры доверяли своим менеджерам и отчётам в компьютере. Вот один практический пример. Основатель компании был выходцем из производства. Ежеженедельно он ходил на производство. И в итоге люди понимали, что их работа важна, были мотивированы, и он обладал достоверной информацией.

Но вот его преемник имел классическое менеджерское образование. Он сформировал команду менеджеров, систему отчётов. Но что хотят наёмные руководители? Чтобы они были в безопасности и получали стабильно ЗП. Это нормально. Однако с каждым этапом, от каждого руководителя информация всё сильнее искажается, и в итоге на самом верху формируется искажённое видение ситуации, а значит, принимаются решения, основанные на ошибочном мнении. Это, кстати, одна из главных проблем государственного управления.

Если использовать этот принцип в сочетании с качественным, автоматизированным сбором данных, прозрачной аналитикой, то подобной проблемы можно избежать.

Искусственный интеллект может стать глазами руководителей, анализировать видеозаписи, профили сотрудников и выводить рекомендации, куда идти и с кем общаться, кто наиболее активный и ценный для компании (в том числе это и поддержание мотивации ценных сотрудников).

Принцип 13. Принимай решение, не торопясь, на основе консенсуса, взвесив все возможные варианты; внедряя его, не медли (немаваси).

Ещё это можно назвать «думай медленно, решай быстро». Одна из основ менеджмента — необходимость оценивать альтернативы. Принимать решения, исходя из одного мнения, слишком опасно и рискованно. Если для молодых компаний это порой единственный вариант развития, то чем дальше, тем чаще необходимо использовать консенсус, а для этого нужна проработка вопроса и несколько людей с различными мнениями и психотипами. На этом, в том числе, держится концепция Адизеса — один человек не может совместить все необходимые компетенции,

нужна разносторонняя команда. А для этого необходимо научиться открыто обсуждать мнения и преодолевать конфликты.

Нельзя принимать однозначное решение о способе действий, не взвесив все альтернативы. А когда решение уже принято, необходимо действовать.

Это ДАО как раз и приводит к концепции цифровых советников, которые смогут обрабатывать данные из ИТ-систем, мнение отдельных экспертов и формировать рекомендации и возможные альтернативы.

Принцип 14. Станьте обучающейся структурой за счёт неустанного самоанализа (хансей) и непрерывного совершенствования (кайдзен).

Я уже говорил, что невозможно всё оптимизировать за один раз.

Цифровые двойники смогут самообучаться уже без влияния человека и формировать рекомендации по оптимизации, приоритизировать изменения и избегать масштабных и высокорискованных проектов. То есть это про снижение рисков. Чем меньше проект и изменения, тем проще им управлять и ниже цена ошибки.

ИИ сможет повысить эффективность, правильно подобрав методологию управления изменениями и выделив ключевые шаги в проекте. И всё это без необходимости вкладывать миллионы во всеобщее обучение сотрудников.

Потери

Один из главных инструментов в бережливом производстве — 8 видов потерь. Они встречаются и на производстве, и в офисе, и в ИТ.

Нужно уметь выявлять потери в процессах и знать, как та или иная технология может помочь устранить их. Это поможет не только наладить собственную работу, но и устранять избыточные требования к цифровым проектам в самом начале, тем самым снизить стоимость и сроки на их реализацию.

Ниже приведён краткий список потерь и их проявление.

Вид потерь	Примеры потерь для производства, офисной работы и ИТ-систем
Перепроизводство	Изготовление лишних копий документов, отчётов Дублирование информации в разных документах / ИТ-системах Дублирование поручений Избыточная производительность ИТ-систем Неиспользуемые функции в ИТ-системах
Ожидания	Ожидание из-за доставки / неготовности оборудования Ожидание из-за отсутствия указаний, согласований, планирования Медленная работа ИТ-системы
Запасы	Хранение запаса материалов на полгода производства без учёта стоимости обслуживания склада Запасы канцтоваров Большое количество нерассмотренных задач Хранение в архивах неиспользуемых документов Хранение данных, которые не используются, но потребляют ресурсы ИТ-инфраструктуры. Например, еженедельное сохранение архива системы за год.
Излишняя транспортировка	Расположение склада запчастей и производства на большом расстоянии друг от друга Передача документов вручную Излишнее согласование документов в маршрутах электронного документооборота
Излишнее перемещение / движения людей	Поиск необходимого для работы инструмента по всему участку Неудобное расположение техники или мебели Перелистывание документов из-за отсутствия оглавления и/или нумерации страниц Личные встречи и совещания, которые можно заменить на онлайн или использование электронных досок Запутанный интерфейс ИТ-систем или неоптимизированная логика работы системы, требующая лишних нажатий. Например, отсутствие ссылок для навигации по документам. Неструктурированное хранение файлов

Брак	Неквалифицированный работник сделал неверные расчёты Получение идентичных замечаний при типовых задачах Получение различных замечаний для типовых задач / документов Запрос информации, требующей от исполнителя уточнения формулировок Недостоверные данные в отчётах из-за ошибки в алгоритмах ИТ-систем Сбой или отказ в работе ИТ-системы
Излишняя обработка	Изготовление детали или оказание услуги с избыточной точностью Избыточные требования в техническом задании Постоянное обновление дизайна и элементов интерфейса без запроса пользователей или исследований
Неиспользованный человеческий потенциал	Найм сотрудника с избыточной квалификацией Большое количество ручного и рутинного труда, который можно автоматизировать Игнорирование предложений по улучшению / оптимизации процессов от сотрудников

Бережливое производство как этап подготовки к внедрению ИИ («Lean перед ИИ») системно разобрано в книге «ИИ-2. Практическое руководство» (Глава 8); фокус этой главы — взаимное влияние Lean, ТОС и ИИ.

А сами инструменты — пять фокусирующих шагов ТОС, 14 принципов и виды потерь бережливого производства — я подробно разбирал в книге «ЦТ-2. Системный подход» (Главы 2–3).

Резюме главы

- «Lean перед ИИ» — ключевая закономерность книги: сначала устранить потери и навести порядок в процессе, затем автоматизировать. Иначе ИИ лишь масштабирует беспорядок.
- Синергия: ИИ находит потери и «узкие места» (ТОС) в данных и возвращает руководителю «глаза» — но гемба и живое наблюдение остаются незаменимыми.

- Обратное влияние: бережливые процессы дают чистые данные и понятные метрики — готовую базу для обучения и контроля ИИ.

Что с этим делать: Выберите один процесс и снимите по нему 8 потерь (неделя). Только после упрощения решайте, где в нём нужен ИИ (квартал).

Вопрос для рефлексии: В каких задачах вы как руководитель — «передаст», а в каких создаёте ценность?

Глава 18. Влияние ИИ на процессы функциональных направлений

Мы рассмотрели влияние ИИ и влияние на ИИ основных цифровых технологий и инструментов системного подхода. Но я понимаю, что некоторым людям проще смотреть немного с другого ракурса. Поэтому сейчас повторю всё то, о чём говорилось ранее, но в ином ключе.

Жизнь, как обычно, оказалась быстрее печатного слова. Пока готовилась эта редакция, рынок принёс главе два подтверждения — и одно отрезвление. В клиентской поддержке рынок прошёл первую волну разочарования: по исследованию Sinch (май 2026), 74% компаний откатали уже запущенных ИИ-агентов — выигрывают внедрения с человеком в контуре (подробный разбор — в Постскриптуме). В России ИИ-агенты решают узкие отраслевые задачи: подбор объектов у застройщиков, обработка обращений в ЖКХ, диагностика лабораторного оборудования (Yandex B2B Tech, март 2026). В разработке ПО отчёт DORA (2026) показывает окупаемость ИИ-ассистентов примерно за восемь месяцев — при зрелых инженерных практиках. Общий знаменатель: эффект приносит не «ИИ вообще», а конкретный сценарий в конкретной функции.

Цифровой бизнес

Это направление можно разделить на три блока:

- повышение операционной эффективности;
- поиск новых ниш и стартапов для инвестиций;

- развитие технологических партнёрств и пилотирование инноваций.

Первый блок разберу подробно — он касается любой компании. Два других относятся к корпоративному развитию, а не к работе с процессами: там ИИ работает так же, как везде, — как инструмент анализа и подготовки решения, а не как самостоятельный игрок.

Внутри направления повышения операционной эффективности можно выделить:

- ИТ-инфраструктуру и её оптимизацию;
- материально-техническое обеспечение и логистику;
- организационную эффективность;
- продажи и текущие продукты/услуги, то есть клиентский сервис;
- экономику и финансы, в том числе бухгалтерию;
- управление персоналом.

ИТ-инфраструктура

Здесь ИИ позволит упростить вопросы проектирования ИТ-инфраструктуры и процессы её обслуживания. Однако проекты ИИ существенно увеличат нагрузку на серверное оборудование. Тут логичнее даже будет сказать, что ИИ-проекты потребуют закупки или аренды дополнительных мощностей. Также нужно учитывать риски информационной безопасности — возможность взлома ИИ-моделей, а также защищённость каналов связи, особенно если вы будете использовать облачные сервисы.

Если посмотреть на влияние ИТ-инфраструктуры на проекты ИИ, то ограничения инфраструктуры и требования по информационной безопасности также будут стимулировать развитие лёгких и специализированных решений.

Экономика и финансы

Здесь ИИ можно использовать для упрощения автоматизации процессов (например, распознавание первичной документации и перенос данных в

учётные системы), прогнозирования результатов, планирования бюджетов и рекомендаций с учётом контекста ситуации и организации.

На данный момент именно экономика — слабое место текущих ИИ-проектов как для бизнес-заказчиков, так и для разработчиков ИИ-решений. Вкупе с требованиями ИТ-инфраструктуры и информационной безопасности мы снова приходим к необходимости создания дешёвых моделей.

Организационная эффективность

Тут мы получим помощь в:

- организации совещаний, а именно в подготовке повестки, распознавании аудиозаписей и подведении итогов встреч;
- управлении проектами и портфелями проектов, изменениями. Например, определение целей, выявление ключевых рисков, определение необходимости создания новых процессов или редактирование существующих, выбор подхода к реализации проекта и т.д.;
- аналитике бизнес-процессов и их проектировании;
- разработке стратегии компании, организационной структуры, метрик для каждого подразделения и внедрении прочих инструментов системного подхода.

Материально-техническое обеспечение и логистика

Здесь целесообразно использовать ИИ в качестве консультантов для подразделений при организации закупок. Также ИИ поможет при анализе закупочной документации и коммерческих предложений поставщиков, в том числе поиск информации о поставщиках на предмет их надёжности. Также здесь будет помощь в анализе и объединении идентичных заявок, оптимизации складской логистики и маршрутизации доставок.

Управление персоналом

Здесь тоже работает всё из блока про HRM-системы (Глава 15). Отдельно лишь отмечу, что именно сопротивление людей станет серьёзным ограничением. Причём работать тут будут все основные причины сопротивления:

- техническое сопротивление из-за непонимания, что это и как работает, отсутствия инструкций и обучения, сырых бездушных ИИ-решений (особенно если что-то внедрять в формате чат-ботов);
- культурное сопротивление из-за неудачного или безрезультатного опыта внедрения ИТ-решений в прошлом;
- политическое — из-за страха потерять работу.

Соответственно, при внедрении ИИ нужно уделять особенное внимание этим вопросам:

- проводить обучение персонала как по цифровой грамотности в целом, так и по работе с новыми инструментами, закладывать это в план реализации проекта, требования к поставщикам ИИ-решений, а также уходить от чат-ботов в пользу программ, в которых пользователь не взаимодействует напрямую с ИИ;
- организовывать внедрение небольшими шагами через пилотные проекты, поиск новаторов и PR своих успехов;
- объяснять людям, что это нужно для повышения эффективности и конкурентоспособности, снижения роста штата и предотвращения ухода в бюрократию, включать ИИ-проекты в метрики руководителей.

Клиентский сервис и коммерческая эффективность

В этом случае можно также вспомнить направление CRM-систем, в частности:

- распознавание записей и фиксация договорённостей, заполнение данных в CRM-системе;
- формирование подсказок по прошлым звонкам и подготовка скриптов;

- первая линия техподдержки и консультация клиентов по статусам заказов и типовым вопросам, ответы на письма и отзывы;
- расчёт индивидуальных коммерческих предложений;
- аналитика продаж и подготовка рекомендаций руководителям;
- подготовка планов продаж и их вариантов;
- анализ запросов клиентов и подготовка предложений руководителям.

Цифровой маркетинг

Это функция по генерации цифрового контента и управлению сквозной коммуникацией через цифровые каналы.

Это направление, где ИИ уже развивается и применяется активно. Например, подготовка рекламных кампаний и анализ результатов, исследования рынка, подготовка индивидуальных предложений клиентам в рамках программ лояльности, генерация контента для продвижения, подготовка гипотез для тестирования, SEO-описаний, маркетинговых материалов и описаний, ответы на комментарии, анализ общих тенденций в отзывах. В общем, список обширный и маркетологи уже пробуют и учатся.

Цифровое производство

Технологии, датчики, роботизация и искусственный интеллект начинают использоваться на производстве, и часть решений на конвейере уже принимает не человек, а машина. Проектирование изначально идёт в цифровой среде одновременно с созданием цифровых двойников. Также это про техническое обслуживание и ремонты, промышленную безопасность, оперативную аналитику, автоматизацию отчётности, проектирование и технологические процессы, конструкторскую документацию, контроль качества и т.д.

Здесь мы можем выделить следующие направления:

- подбор материалов для изготовления деталей;

- проведение исследований и автоматизированная подготовка новых прототипов (так, конструирование нового ракетного двигателя заняло несколько недель вместо нескольких лет);
- анализ больших данных с датчиков и выявление общих паттернов работы оборудования, подготовка рекомендаций по оптимизации режимов работы;
- выявление аномалий в работе оборудования, в том числе для контроля качества технического обслуживания и подготовки рекомендаций для операторов, предотвращения аварий, снижения количества ошибок человеческого фактора и при этом снижение требований к операторам;
- контроль качества продукции и производственных процессов в целом;
- планирование технического обслуживания;
- подготовка производственных планов и отчётности;
- контроль соблюдения требований производственной безопасности, в том числе соблюдение прав доступа к работе на оборудовании;
- разработка конструкторской документации, проектирование технологических процессов.

Цифровая аналитика

Это сбор и систематизация данных со всех каналов и источников. То, что известно под названием Big Data, или Большие данные. Здесь ИИ станет помощником как для аналитиков данных (анализ того, что случилось), так и для дата-сайентистов (предсказывание будущего), позволит ускорить создание математических моделей и выявление любых паттернов:

- ускорение процесса оптимизации и совершенствования узлов оборудования на основе данных с датчиков с реальных изделий;
- выявление паттернов поведения клиентов.

В общем и целом, именно тут можно ожидать наибольшего вклада от ИИ, наравне с маркетингом.

Цифровые продукты

Речь идёт о создании новых цифровых продуктов в виде сервисов аналитики, приложений. Если бизнес научится оцифровывать свою экспертизу, то для сферы услуг это будет прорывное направление. С одной стороны, любые консалтинговые услуги окажутся под ударом, с другой, отдельные компании могут стать единорогами.

Для производственного сектора это возможность создать более качественные сервисы полного цикла — от продажи до утилизации оборудования.

Новые способы работы

Это направление связано с тем, как мы работаем и думаем на работе. Тут основными направлениями станут помощь в коммуникациях и управление проектами.

А что ИИ меняет в работе самого руководителя — не компании, а человека в кресле директора? Это настолько важная тема, что ей посвящена отдельная глава в третьей части (Глава 21).

Инженерия и конструирование: от генерации к верификации

Отдельного разбора заслуживает инженерно-конструкторский блок — он дозрел до пилотов позже маркетинга и поддержки, но бьёт по самой сердцевине производственных компаний. ИИ-ассистенты инженера-технолога и конструктора уже проходят пилоты на крупных машиностроительных предприятиях: по конструкторской документации система выдаёт прототип технологического процесса, по готовой детали или чертежу восстанавливает редактируемую цифровую модель; в отдельных кейсах подготовка проектной документации сокращается в разы. Обратите внимание, что именно сжимается. **Не профессия — сжимается конкретная фаза работы:** черновая генерация вариантов и рутинное оформление документации. Ровно та фаза, на которой раньше учился и рос младший инженер.

Сдвиг первый: от генерации к верификации. Инженер перестаёт быть автором первого варианта и становится редактором и валидатором. Это когнитивно сложнее, а не проще: критиковать правдоподобный, но неверный техпроцесс труднее, чем создать свой с нуля. Здесь критична прослеживаемость источника (см. Главу 5): без неё верификация превращается в ритуал — человек просто нажимает «принять».

И упирается этот сдвиг в ту же способность: инженер, свободно владеющий терминологией своей области, и задачу ставит точнее, и правдоподобную ошибку в ответе замечает раньше. Язык профессии становится инструментом контроля.

Сдвиг второй: от одного варианта к управлению пространством вариантов. Когда генерация стоит копейки, ценность смещается к постановке ограничений: допуски, доступная оснастка, реальный парк станков, себестоимость, ремонтпригодность, безопасность. **Умение точно сформулировать ограничение становится главной инженерной компетенцией.** По сути, это возвращение культуры технического задания — просто теперь плохое ТЗ наказывает не через полгода, а через пять минут, и в промышленных масштабах.

Сдвиг третий: от личной экспертизы к капитализации знания. Работающий ассистент — это ассистент, обученный на ваших техпроцессах, рекламациях и отказах. Значит, знание, которое тридцать лет жило в голове ветерана цеха, должно попасть в данные — и это меняет социальный контракт: неформализованное знание перестаёт быть личной страховкой от увольнения и становится активом компании. **Люди сопротивляются не машине. Люди сопротивляются потере монополии на знание.** Пока вы не решите этот вопрос честно — через статус, деньги, роль наставника-верификатора, — данных вы не получите, а ассистент останется красивым демо.

И главный риск, о котором говорят непозволительно мало: если ИИ забирает работу младшего инженера, откуда через десять лет возьмётся старший, способный ИИ верифицировать? **Мы рискуем выиграть три года производительности и проиграть поколение экспертизы.** Ответ не в

том, чтобы «не давать ИИ джуниорам» — это проиграно, они уже им пользуются. Ответ — перепроектировать обучение: младший инженер учится не на черновой работе, а на разборе ошибок модели. Дайте ему сто выходов ассистента, из которых двадцать неверны, и задачу найти и обосновать ошибки. Это более быстрый и жёсткий тренажёр, чем три года «на подхвате». Но его надо построить осознанно — сам он не сложится.

Резюме главы

- ИИ проникает во все функциональные направления — от маркетинга и производства до аналитики и создания продуктов; различается не логика, а данные и цена ошибки.
- Единая формула внедрения в любую функцию: данные → модель → встроенность в процесс → контроль результата.

Что с этим делать: Выберите функцию с самой болезненной рутинной и одну задачу в ней, зафиксируйте базовые метрики до старта — это кандидат на пилот (неделя). Проведите пилот и сравните результат с базой (квартал).

Вопрос для рефлексии: В какой функции вашей компании или лично вашей деятельности ИИ окупится быстрее всего — и почему именно там?

Резюме второй части

Общие выводы и прогнозы

Я не думаю, что ИИ заменит работников, но он точно сможет помочь продвинутым специалистам обрабатывать больше компаний / клиентов / подразделений. Компаниям же он поможет сдержать рост численности персонала и преодолеть кадровый голод.

ИИ позволит радикально повысить эффективность управления и снизить потребность во всевозможных сотрудниках и менеджерах аппаратов управления. Сценарии применения ограничены фантазией, и это целая отрасль для всевозможных решений: ассистенты и советники, выявление отклонений, обучение персонала, упрощение и ускорение процессов.

Ключевые препятствия на пути внедрения ИИ в менеджмент компаний:

- разработка и обучение специализированных моделей;
- сопротивление людей и недоверие, страх перед новым (в том числе из-за незнания);
- желание переложить ответственность на инструмент (хотя от ошибки и некомпетентности обычного сотрудника вообще никто не застрахован, и ответственность всё равно несёт руководитель организации).

Стоит ли ожидать здесь быстрых результатов? Конечно нет, мы обычную цифровизацию не умеем использовать на полную, внедряя лишь отдельные инструменты и частично решая проблемные вопросы автоматизацией. А здесь речь вообще о малоизученном ИИ, который непонятен и которого боятся. Опять же, здесь нас ждут две крайности — одни будут ждать волшебной пилюли, которая за них построит бизнес и будет им управлять, а другие не смогут доверять.

Я ожидаю, что действительно фундаментальный эффект на управление и развитие технологий мы начнём наблюдать через 30 лет. Сначала технология проникнет в нашу жизнь через B2C-решения, с ней вырастут наши дети (давайте на это заложим 18–20 лет). Затем потребуется время, чтобы наши дети пришли в компании и доросли до руководителей или начали строить свои компании и вывели их из стадии стартапа. На это ещё нужно около 10 лет.

Да, наиболее продвинутые уже сейчас начнут искать прикладное применение, но их можно отнести к новаторам и ранним последователям, которых всего не более 10–15% от общего числа людей. И эта книга им в помощь.

Кроме того, именно на горизонте 30–40 лет технология станет более зрелой: появятся новые, более эффективные алгоритмы, которым нужно будет меньше данных. К тому времени создавать новые инструменты и модели можно будет без специальных навыков. То есть создание будет

проще, быстрее и дешевле. Возможно, мы уже увидим и расцвет нейроморфных и квантовых вычислений.

Часть 3. Практические примеры и рекомендации по применению искусственного интеллекта

Мы разобрались, что такое ИИ и как он встраивается в систему управления. Пора спускаться на землю окончательно — к примерам и рекомендациям. Эта часть — самая «живая» в книге: здесь кейсы, которые я собирал из своей практики, разговоров с коллегами и поездок, и рекомендации, за каждой из которых — чей-то набитый синяк. Читайте её с карандашом: почти каждый пример можно примерить на свой бизнес.

Глава 19. ИИ для малого бизнеса и предпринимателей

Применение ИИ в решениях 1С

Сейчас, в середине 2020-х, решения от 1С в России становятся стандартным компонентом цифровизации. Сам 1С зарекомендовал себя как некий аналог SAP и Oracle, но для российских компаний. Это целая экосистема, которая охватывает различные классы приложений. При этом есть ещё и огромное количество дополнений и вариаций от сторонних разработчиков. Например, различные приложения для управления ремонтами, арендными помещениями и т.д. А значит, тут есть большой простор и для применения искусственного интеллекта.

По сути, 1С становится мастер-системой, с которой интегрируются другие решения. Например, я сопровождал внедрение системы фиксации опасных действий и условий работы на производстве.

Архитектура была следующей:

- Приложение и веб-интерфейс для взаимодействия с пользователями на объектах и подачи заявок.
- НСИ на базе 1С, из которой подтягивались данные об объектах.

- ИИ, который берёт записи из 1С, обрабатывает их, определяет приоритеты.
- 1С для работы ИТР и обработки задач.
- BI-система для формирования отчётов и аналитики.

Если говорить конкретно про сервисы ИИ от 1С, то они развиваются с 2022 года. Среди них можно выделить:

- 1С: Распознавание первичных документов.

Скан мятой товарной накладной распознаётся и оцифровывается, переносится в рабочий формат. Таким образом, сразу после сканирования документ с заполненными полями создаётся в системе. Также ИИ можно обучить и на своём наборе данных.

- 1С: Прогнозирование продаж.

ИИ, который анализирует продажи, готовит прогнозы, формирует управленческий учёт — руководителю не приходится самому разбираться в аналитике.

- 1С: Распознавание речи.

Это транскрибация аудиозаписей. ИИ распознаёт речь в аудиозаписи по ролям — кто какие предложения вносил, кто что обещал сделать и так далее. Затем создаётся протокол с задачами, сроками и ответственными. Плюс можно анализировать разговоры по телефону с клиентом для повышения производительности call-центра или отдела продаж.

- RCM: Прогнозирование отказов оборудования — определяет аномалии в работе оборудования по показаниям датчиков.
- 1С ДО: Чат-бот Ася — отвечает на вопросы на естественном языке, помогает пользователям и имеет широкий спектр применения.
- CRM: iРуководитель — анализирует работу менеджера, даёт рекомендации и пошаговый план действий.
- LIMS: Помощник лаборанта — предоставляет инструкции для проведения измерений на производствах.

При этом ИИ-сервисы имеют свой API для подключения к ним своих баз данных/знаний и обучения ИИ под специфику компании.

Ниже приведено ещё два примера использования ИИ на базе 1С.

Автоматизация обработки фотографий для наполнения базы товаров интернет-магазина

Исходный бизнес-процесс

- Фотограф получает товары, организует фотосъёмку и обработку.
- Загружает все файлы скопом на флешку или в облако (в этом случае нужно было около 40 тысяч фотографий).
- Менеджер вручную распределяет все фотографии по папкам внутри 1С, фиксируя, какой именно файл к какой именно карточке товара относится.

Оптимизированный бизнес-процесс

- Менеджер сканирует штрих-код товара (в 1С открывается нужная карточка), затем кладёт его в специальное место для фотографирования (лайт-бокс), напротив которого стоит камера, подключённая также к 1С.
- Нажимает в 1С кнопку «Сфотографировать». Камера делает снимок, ИИ забирает файл и обрабатывает (очистка фона, сжатие и т.д.).
- Снимок появляется в карточке товара.

Да, в новом процессе менеджеру всё равно надо знать номенклатуру и внешний вид товара, также нужны хотя бы базовые навыки размещения товара в лайтбоксе, но обучить его этому дешевле, чем нанимать фотографа.

Тексты для интернет-магазинов и маркетплейсов

Один из ключевых инструментов онлайн-платформ — прямая коммуникация с пользователями. Но у малого бизнеса нет возможности постоянно быть на связи и отвечать на все отзывы и вопросы клиентов. Особенно с учётом того, что такие маркетплейсы стимулируют

пользователей писать отзывы, на основе которых ранжируются магазины и карточки товаров (опять же с помощью ИИ).

В связи с этим появляются инструменты на базе ИИ, которые сами отвечают на комментарии пользователей. Есть как готовые заготовки, так и конструкторы для создания своих модулей внутри 1С.

Аналогично с этим работают решения для генерации описаний товаров, характеристик, разделов, то, о чём я писал в блоке про РІМ-системы.

Получается интересная связка:

- ИИ формирует карточки, описания, обрабатывает фото.
- Другой ИИ анализирует всё это и ранжирует.
- Пользователи комментируют.
- Ещё один ИИ отвечает.
- ИИ из пункта 2 снова анализирует и ранжирует всё, генерирует рекомендации пользователям.

Маркетплейсы специализированных ИИ-решений

По QR-коду и гиперссылкам ниже вы сможете найти несколько готовых и специализированных под задачи решений для 1С из маркетплейса.



[Заполнение описания номенклатуры с помощью GigaChat](#)



[Чат GPT описание товара: составить описание номенклатуры с помощью ChatGPT с ключевыми словами](#)



[Искусственный интеллект и нейросети в 1С: Работа с отзывами маркетплейсов](#)



[Потоковая предметная фотосъёмка с удалением фона \(Canon & Nikon\)](#)

И подобные маркетплейсы ИИ-решений будут набирать всё большее распространение. Та же компания OpenAI запустила GPT Store — магазин кастомных чат-ботов.

Для этого OpenAI в ноябре 2023 года представила функцию GPT, позволяющую создавать индивидуальных чат-ботов: можно добавлять специальные возможности, навыки и знания. Настройка чат-бота осуществляется с помощью текстового описания его роли (системного промпта). Также можно загрузить источник данных. Про такой механизм я писал в разделе о цифровых советниках — ИИ-модели, которые получают «специализацию» в предметной области.

Магазин GPT Store доступен по QR-коду и гиперссылке ниже.



[Gptstore.ai](https://gptstore.ai)

А по QR-коду и гиперссылкам ниже доступна специальная страница, посвящённая запуску, и конструктор GPTs от OpenAI.



[Introducing GPTs](#)



[Конструктор GPTs](#)

И это огромная возможность для малого и среднего бизнеса — доступная онлайн-площадка как для покупки ИИ-решений с приемлемой ценой, так и для запуска своего решения. По сути, это Ozon для ИИ-приложений.

Цифровой советник для управления проектами

Одна из основных проблем малого и среднего бизнеса — нет ресурсов и возможности сформировать команду из экспертов экстракласса. При этом задач в малом бизнесе нисколько не меньше, чем у крупных корпораций, а уязвимостей, наоборот, больше. Я эту проблему вижу и по своему направлению — управление проектами, изменениями и продуктами.

Сейчас любой бизнес сталкивается с огромным количеством проектов. До 50% рабочего времени руководителей уходит на работу с проектами. При этом статистика управления проектами, изменениями и созданием новых продуктов крайне печальна. Приведу для вас немного данных из международных исследований.

Статистика реализации проектов

Классическая модель

проект реализован в срок, в рамках бюджета и в соответствии с установленными критериями



Современная модель

проект реализован в срок, в рамках бюджета и клиент удовлетворен результатом



Чистая модель

оценка успешности продукта: клиент удовлетворен результатом, а продукт создает высокую ценность



При этом толковый руководитель проекта или продукта стоит 200 000–300 000 рублей в месяц без учёта налогов. Внедрять же корпоративные и комплексные системы управления опять дорого. Такие проекты стартуют, в среднем, от 5 млн рублей. Да и не отвечают они на основные вопросы (как выбрать подход к реализации, какие задачи надо реализовать и т.д.).

Если мы говорим про обучение персонала, то курс будет стоить 50 000–150 000 рублей. Но:

- как показывает практика, есть большие вопросы к качеству обучения;
- слушатели не получают поддержки после обучения (а для формирования компетенций нужно сопровождение до года, плюс опыт в разных ситуациях);

- высок риск того, что вы обучите, а человек уйдёт на более оплачиваемую работу. Причин может быть много: предложили больше денег, у вас стало скучно из-за рутинных задач, выгорание, видение обучения расходится с рабочей реальностью.

И самое важное — для успешности 90% проектов не нужно никаких сверхкомпетенций, нужно просто соблюдать методологию и определённые шаги:

- определить цель проекта и задачи;
- определить заинтересованные стороны и работать с ними;
- оценить риски и отслеживать их;
- вести дневники встреч;
- вовремя замечать отклонения и т.д.

Отсюда у меня появилась простая мысль — управление проектами, изменениями или продуктами должно стать стандартной компетенцией. А ещё лучше — сделать советника, которому ты ответил на тестовые вопросы, и он дал рекомендации, что делать и как, а потом сопровождает тебя и обучает.

То есть мы говорим о готовой методологии для бизнеса от лучших умов и развитие его сотрудников. Просто заплатил 2 000 рублей в месяц за сотрудника и обзавёлся личным ассистентом, который и подскажет, и сопроводит.

И для достижения этой цели нужно решить несколько задач.

Первая — разложить модель на факторы, влияющие на итоговый результат. По моему видению, это характеристики самого проекта, культура компании и компетенции руководителя проектов.

Вторая — нужен не чат-бот, а комплексное решение для ситуации, когда пользователь ничего не знает об управлении проектами. То есть он либо отвечает на тестовые вопросы, либо выбирает нужное из выпадающего списка с подсказками, либо вручную вводит уже совсем простые данные (ФИО участников, их почту, подгружает рабочие шаблоны).

Третья — нужна замкнутая система, которая сама собирает данные для самообучения, в том числе из других ИТ-систем.

Четвёртая — нужна готовая методология и учёт личных качеств руководителя проекта. Даже если мы посмотрим на лучшие обучающие курсы и практики, то человеческий фактор исключить невозможно, и многое базируется на опыте и психологических качествах руководителя.

Да, это не одна большая ИИ-модель, а целый оркестр. При этом мы создаём замкнутую систему, где ИИ будет дообучаться на основе отзывов заинтересованных сторон в ходе проекта и по его итогам и данных о планировании и исполнении бюджета, сроков реализации.

Главная идея — реализовать принципы глубокого обучения, в котором ИИ сам ищет взаимосвязи и корреляции на базе представленной модели. Так мы уходим от чат-бота в пользу готового интерфейса. Например, заполнил блок заинтересованных сторон, а система сама эти данные потом вставит в устав проекта, отправит ссылку на опрос и подскажет, что Иван Иванович выпал из проекта и нужно бы с ним поработать.

Самое прекрасное тут, что проект выходит в точку безубыточности всего при 1–2 тысячах пользователей в месяц (инфраструктура + развитие методологии + команда разработки). А главные риски (они будут типовыми для всех подобных решений):

- сложность;
- маркетинг и отсутствие осознания у людей, что есть вообще слово «проекты» и ими можно как-то управлять;
- с таким решением нельзя идти сразу в корпорации, т.к. они будут требовать адаптации под свои процессы. А они, чего греха таить, зачастую «кривые». Помимо этого, каждый проект будет стоить десятки миллионов, длиться месяцами и не всегда получится его довести до конца. Кроме того, это ещё и «длинные» деньги. Оплату в итоге вы получите через год-два после старта, а значит, велик риск получить кассовый разрыв.

В индустрии создания цифрового контента генеративный ИИ используется уже давно. Вот пример от моего партнёра, создавшего серию видео для вовлечения сотрудников одной крупной энергетической компании в новую культуру, основанную на философии бережливого производства. Поведаю с его слов от первого лица.

«Как часто бывает у заказчиков, нужно сделать что-то, но никто не знает что. В мои задачи входило разработать стиль, отстраивающий новую культуру от старой, и презентовать его через видеоролики, которые непрерывно бы воспроизводились на главных экранах в помещениях компании.

«Разумеется, надо сразу было учесть, что мимо этих экранов люди проходят, как правило, даже не обращая на них внимание. Также нельзя рассчитывать на звуковое сопровождение. То есть нужно было сделать что-то яркое, быстрое и... немое.

«Создание видеоконтента — дело привычное для меня, но тут возникли сложности: во-первых, новый стиль подразумевал много графических элементов, и сама картинка как бы собиралась из них подобно коллажу, а я хоть и постоянно работаю с дизайном, всё-таки не являюсь иллюстратором; во-вторых, было решено для привлечения внимания использовать лица реальных сотрудников, чтобы, пробегая мимо экрана, они всё-таки остановились, случайно узнав себя или коллег. Но, как оно часто бывает, реальные сотрудники не особо горели желанием поработать на камеру.

«Решить все эти проблемы помог ИИ.

«Так, сначала при помощи одной нейросети я создал визуальный ряд — нагенерировал десятки картинок в едином графическом стиле. Они стали основой для будущего оформления проекта.

«Далее, получив разрешение от сотрудников использовать их фотографии из соцсетей, я через вторую нейросеть интегрировал их лица в сгенерированные изображения. Так я получил требуемые кадры и при этом вообще никого не напряг съёмками. Тут надо оговориться, что я бы

это сделал и без нейросети, при помощи графического редактора. Однако нейросеть с этим справилась в разы быстрее, мне же оставалось только „дошлифовать“. Скорость важна.

«После дополнительной обработки я загрузил полученные изображения в третью нейросеть в качестве референсов для видео. На выходе я получил анимированные кадры с сотрудниками компании в нужном графическом стиле.

«Оставалось только смонтировать, добавить элементы моушн-дизайна, и серия видео была готова. При этом не потребовалось абсолютно никакого съёмочного процесса, дополнительного оборудования и участия людей».

Таким образом, ИИ заменил целую бригаду контент-мейкеров и актёров. При этом нельзя сказать, что ИИ выполнил всю работу, ведь он был лишь инструментом в профессиональных руках. Главной движущей силой всего проекта остался человек, который смог реализовать свой творческий потенциал, используя нейросети. В этом и есть смысл работы с генеративным ИИ.

Создание ассистентов

Ещё один пример дешёвого и эффективного использования ИИ. Мой партнёр занимался логистикой леса (импорт). Команда небольшая, 5 человек генерировало 200 млн выручки в год. Одной из задач, отнимавшей много времени, стало общение с контрагентами по договорным обязательствам.

В итоге было сделано простое решение. Он купил подписку на ChatGPT, создал там виртуального сотрудника, сформировал для него системный промпт (кто он, какую роль выполняет, что от него ожидается). Далее, когда к нему приходила претензия или он готовил претензию, то загружал этому ассистенту договор и описывал суть претензии (своей или контрагентов). ИИ в течение нескольких минут готовил ответ с подробным описанием и отсылками к пунктам договора. В итоге вместо 3–4 часов работы всё решалось в считанные минуты.

Цена вопроса — несколько тысяч рублей в месяц и пара часов на создание этого ассистента.

Бьюти-ассистент

Над этим проектом мы работали в 2025 году. Суть ассистента проста. Девушка создаёт свой профиль, а ИИ решает несколько задач:

- Берёт данные из профиля и подбирает по критериям подходящий салон.
- Показывает на руке клиента варианты дизайна ногтей и прикрепляет его выбор к заявке салону, чтобы сэкономить время и убрать конфликты во время оказания услуги.
- Оценивает качество сервиса (между тем, что было согласовано и что в итоге сделали).

Рекомендации для малого и среднего бизнеса

ИИ для малого и среднего бизнеса может стать ключевым для успеха ресурсом:

- автоматизация и ускорение отдельных задач;
- автоматизация и ускорение целых бизнес-процессов (например, запуск маркетинговой активности);
- получение доступа к экспертизе лучших специалистов и рекомендаций, что и как делать;
- создание и генерация контента для маркетинга;
- повышение эффективности взаимодействия с клиентами через анализ данных в CRM-системах, на площадках с отзывами и в интернете.

Другими словами, это возможность не только выиграть в конкурентной борьбе, но и расти с минимизацией вложений в персонал. То есть рост по оборотам и ассортименту товаров быстрее, чем рост команды.

Вторая ценность — качественные рекомендации без привлечения дорогостоящих специалистов, как для консультаций, так и в штат к себе.

Также ниже приведу ряд рекомендаций, что и как делать.

- Пользуйтесь вашей гибкостью и возможностью рисковать. Пробуйте разные ИИ-решения, набивайте руку, учитесь автоматизировать задачи с помощью промптов. Это позволит вам понять, как работает ИИ, чего от него ожидать. При этом новые ИИ-решения и их связки появляются каждую неделю. Пусть ИИ станет частью вашей жизни.
- Ожидайте появления цифровых советников по своему направлению. Например, советник по налогам, управлению проектами, управлению запасами и т.д.
- Подпишитесь на акселераторы и смотрите, какие стартапы там заходят. Пробуйте их решения, и если помогают, то покупайте / подписывайтесь, пока цена низкая. А ещё лучше — заходите в долю в тех проектах, которые вам интересны с прикладной точки зрения.
- Изучайте обновления функционала ваших ИТ-сервисов, которые вы используете по подписке. Сейчас ИИ-инструменты внедряются почти везде. Например, в Мире есть инструмент, в котором вы пишете большую задачу, а ИИ декомпозирует её на отдельные шаги и визуализирует их связь, группирует их. Да, это просто чат-бот, но с удобной визуализацией.
- Скоро появятся ИИ-конструкторы, и если вы сможете оцифровать свой опыт и экспертизу, уложить это в продукт, то имеете шанс построить новую компанию. Основная сложность — структурирование знаний в модель, маркетинг и продвижение. Не спешите в B2B-сегменты. Создавайте B2C-решения. У корпораций будет множество требований, из-за чего велика вероятность не вытащить внедрение, или цена станет такой, что проект потеряет всякий смысл. B2C-рынок или другие небольшие компании более открыты к коробочным решениям. Станьте стандартом на рынке и потом идите в крупный B2B. Главное ограничение заключается в том, что рынок и люди пока не понимают, что этот продукт им нужен.

- Используйте маркетплейсы ИИ-решений для поиска решений для автоматизации задач.

Отечественный Битрикс также внедряет своего ИИ-ассистента.

Промпт-инженеры стандартизировали типовые задачи и создали промпты для готовых сценариев. Подробно по QR-коду и гиперссылке.



[Bumpuk24 CoPilot](#)

Также в июне 2025 года разработчик анонсировал внедрение ИИ-агентов в свой сервис — им можно задать вопрос, поставить задачу и посоветоваться. Агенты смогут подключаться к базе знаний компании и самостоятельно находить необходимую информацию. Пользователи смогут изменять параметры агентов и настраивать к ним доступ у сотрудников. Цифровых помощников можно бесшовно вызывать из интерфейса сервиса в специальном разделе.

Многие из этих решений, особенно на первом этапе, будут бесплатны. Особенно если вы готовы к решениям на серверах разработчиков по подписочной модели (SaaS-решения).

Свежий срез рынка дал Yandex B2B Tech (март 2026): из более чем 7 500 ИИ-агентов на платформе Yandex AI Studio 86% используются малым и средним бизнесом — сценарии от техподдержки и HR-консультаций до подбора объектов у застройщиков. Подробный разбор этих данных — в Постскрипуме.

Эти данные подтверждают ключевой прогноз, который я сформулировал ранее: компании МСБ более гибкие и нацелены на эффективность. Они не ждут «идеального ИИ» — они берут готовые инструменты и адаптируют их под свои задачи. Вспомним данные MIT: пилоты, где компания брала внешний инструмент и дорабатывала его под себя, доходили до эксплуатации примерно вдвое чаще собственных разработок.

Отсюда вытекает и рекомендация для разработчиков ИИ-решений: делайте готовые продукты для МСБ с месячным чеком до 100 тысяч рублей (оптимально — 20–50 тысяч рублей). Именно в этом сегменте сосредоточен основной спрос, именно здесь готовность к экспериментам максимальна, а бюрократическая инерция минимальна.

Пошаговые планы запуска ИИ для малого бизнеса — первые 30/60/90 дней, чек-листы и бюджеты — вынесены в книгу «ИИ-2. Практическое руководство» (Глава 17). А главный ориентир при выборе пути уже назван выше: покупайте и подписывайтесь — не стройте.

Для интереса я спросил рекомендации и у самих моделей — их советы (определить цели, начать с малого, обучить людей, обеспечить безопасность данных, замерять результат) практически повторяют изложенное выше. Показательно: базовая управленческая логика внедрения уже «защита» в сами модели — вопрос лишь в дисциплине её применения.

Резюме главы

- Порог входа для малого бизнеса упал радикально: готовые облачные сервисы и коробочные ИИ-продукты дают доступ к технологии без разработчиков и серверов.
- Выигрывают узкие сценарии «задачи дня»: документы, клиенты, контент, рутина — а не «ИИ вообще».
- Человек остаётся в контуре: у бизнеса нет права на ошибку полностью автономного агента — контролируемая автономия дешевле исправления последствий.

Что с этим делать: Выберите одну повторяющуюся боль (обработка заявок, документы, контент) и закройте её готовым сервисом — пилот за неделю, без бюджета.

Вопрос для рефлексии: Какую операцию вы или ваши сотрудники делаете руками каждый день, хотя её уже умеет делать ИИ?

Глава 20. Примеры и рекомендации для среднего и крупного бизнеса

Десять кейсов подряд — это много, поэтому начну с навигатора. Найдите свою отрасль или функцию и идите сразу к нужному кейсу; остальные читайте как хрестоматию чужого опыта.

Кейс	Отрасль / функция	Технологии	Что смотреть именно вам
Заключение сделок (Bain & Co)	B2B-продажи	Генеративный ИИ	Как ИИ ускоряет подготовку сделок и предложений
Железные дороги Китая	Транспорт, инфраструктура	IoT + Big Data + ИИ	Связка датчиков, данных и ИИ в масштабе страны
Электроснабжение	Энергетика	ИИ в управлении	Как «имиджевый» проект превратить в осмысленный
Норникель	Горно-металлургия	Машинное зрение, ML	Портфель ИИ-направлений промышленного гиганта
ЕвроХим	Химия	Портфель ИИ	Видение и системный подход к развитию ИИ
Переработка руды	Металлургия	Машинное зрение	Автоматизация ТОиР: от негабарита до простоев
Техподдержка	ИТ, сервис	LLM-ассистенты	Экономика первой линии поддержки

Кофемания	Общепит, ритейл	Генеративный ИИ	С чего начинать: стратегия, а не технология
Нефтедобыча	ТЭК, проекты	LLM в планировании	Границы применимости: где ИИ пока слаб
HR и подбор	Кросс-отраслевая	Мультиагентные системы	Ответ на кадровый дефицит

Генеративный ИИ при заключении сделок

Интересным исследованием поделилась консалтинговая компания Bain & Co, одна из трёх крупнейших и самых престижных мировых компаний в области управленческого консультирования. Также в этом списке Boston Consulting Group (BCG) и McKinsey & Company. Для понимания масштаба этих компаний приведу сравнительную характеристику.

Компания	Количество сотрудников	Выручка (млрд \$)	Выручка на сотрудника (тысяч \$)	Количество городов присутствия	Штаб-квартира
McKinsey & Company	45,000 (2023)	16 (2023)	355	133	Нью-Йорк
Boston Consulting Group	30,000 (2022)	12.3 (2023)	440	100+	Бостон
Bain & Company	15,000 (2022)	5.8 (2021)	387	65	Бостон

Согласно исследованию, компании начинают использовать искусственный интеллект для проверки данных и юридической экспертизы при заключении сделок M&A (слияния и поглощения компаний). Так, из 300 респондентов 80% планируют применить ИИ при принятии решений в ближайшие 3 года. При этом 16% опрошенных специалистов по слияниям и поглощениям заявили, что уже

использовали искусственный интеллект при заключении сделок в прошлом.

Генеративный ИИ может выявить объекты слияний и поглощений, которые не были бы обнаружены с помощью традиционных инструментов, отметить отклонения в контрактах и помочь сосредоточиться на проблемных областях, говорится в отчёте Bain & Co. Среди тех, кто первыми начали использовать генеративный ИИ при заключении сделок, — технологические, медицинские и финансовые компании. «Как правило, это крупные компании с умеренной активностью в области слияний и поглощений — от трёх до пяти сделок в год», — говорится в отчёте компании.

Промышленность и инфраструктура

Использование Big Data + IoT + AI на железных дорогах Китая

Именно связка Интернета вещей, больших данных и ИИ — одна из наиболее перспективных в сфере применения ИИ. Об этом я говорил ещё в книге «ЦТ-1. Погружение» в 2021 году.

Отличный пример этого продемонстрировали коллеги из Китая.

Содержание парка подвижной техники и инфраструктуры на железных дорогах — задача не из лёгких и крайне дорогая. Это потенциально убыточные предприятия во всём мире, особенно когда оборудование начинает стареть. При этом, если посмотреть на Китай, то железнодорожная сеть в стране растёт и предполагает соединение высокоскоростными ж/д магистралями всех городов с населением свыше 500 тыс. человек. Скорость подвижного состава также растёт. Это приводит к тому, что люди становятся всё более опасным компонентом всей этой высокоскоростной системы.

Поэтому в 2022 году китайская государственная компания China State Railway Group начала внедрять процессы управления данными. Как мы уже рассмотрели ранее, управление качеством данных, чтобы они были полными и достоверными — одна из ключевых задач на пути к внедрению ИИ.

Основным ограничением стало то, что доступ к данным должен быть защищён от стороннего вмешательства и утечек.

Были установлены датчики на объектах инфраструктуры и железнодорожном полотне, на колёсных парах и вагонах. Основные измерения — амплитуда и частота вибрации, сила ускорения и т.д.

При этом работала и обычная сигнальная автоматика. Объём собранных данных достиг 200 Тбайт. Причём эти данные представляют собой текст и цифры, без картинок и видео. То есть это именно большие данные, которые человек не может оперативно обрабатывать. Испытания охватили всю высокоскоростную сеть страны.

На собранных данных обучили ИИ-модель, после чего (в 2023 году) начались испытания. ИИ в режиме реального времени начал анализировать и обрабатывать данные с датчиков Интернета вещей и предупреждать ремонтные бригады о нештатных ситуациях за 40 минут до их развития с точностью до 95%. Сеть, к слову, немаленькая: 45 000 км высокоскоростных линий (SCMP со ссылкой на рецензируемую публикацию в журнале China Railway). Рекомендации были направлены на предотвращение неисправностей и раннее устранение. То есть это система MLAD — machine learning anomaly detection или раннее выявление аномалий с помощью машинного обучения.

Сам центр обработки данных находился в Пекине. По итогам испытаний действующая не один год железная дорога стала работать даже лучше, чем новая. За год испытаний количество предупреждений о необходимости снижения скорости из-за неровностей пути снизилось до нуля, а количество мелких неисправностей на путях сократилось на 80%. При этом повысилась не только надёжность, но и комфорт — ход состава стал плавнее в условиях сильных ветров и на мостах, снизилась амплитуда колебаний составов и уменьшилась нагрузка на пути и инфраструктуру.

Применение ИИ в электроснабжении

Этот кейс из области моего общения с коллегами в электроснабжении. Изначально вопрос был простой: «А как нам можно применять ИИ для

организации управления освещением? Проект имиджевый, и надо большому руководству показать, что мы можем в передовые технологии».

Изначально меня очень смутил этот вопрос. Ведь имиджевые и политические проекты — хорошо, они помогают получить ресурсы на действительно интересные задачи. Но это очень рискованные проекты, ведь если внутри вместо реального решения прикладных задач будет только PR, то есть риск вообще всё потерять.

Но чем глубже я погружался в проработку задачи, тем интереснее она становилась. В итоге нашлось два прикладных применения.

Голосовая система управления

Работа диспетчера связана с большим количеством переключений и аналитики. Да, всё идёт к автоматизации типовых сценариев, но всё равно их надо выбирать. А интерфейсы SCADA-систем хоть и улучшаются, но ещё далеки от нормальной оптимизации пользовательских сценариев и интерфейсов (UX / UI).

Здесь ИИ можно использовать в качестве голосового ассистента.

Например, диспетчер отдаёт команду голосом: «Отключить улицу А для аварийного ремонта.» А ИИ:

- распознаёт команду;
- анализирует запрос, сравнивая с базой знаний;
- определяет подходящий сценарий или прорабатывает, с каких участков нужно снять напряжение;
- превращает это всё в набор команд и далее выводит диалоговое окно, верно ли он понял команду и нужно выполнить такие и такие действия.

Диспетчер либо подтверждает задачу и соглашается с рекомендацией, либо отклоняет и занимается вводом вручную (в дальнейшем система запомнит эту команду и список действий и занесёт его в базу знаний).

Помимо имиджевого эффекта, это может иметь ещё и прикладной характер. Условно, если сейчас оператору нужно сделать 20 движений в интерфейсе, чтобы запустить сценарий, то тут 1 команда + 1 движение для подтверждения. А значит, увеличение производительности диспетчера.

Кроме того, мы не даём ИИ напрямую управлять инфраструктурой, оставляя последнее действие за человеком.

Уведомление об отключениях и формирование заявки ремонтной бригаде

Наши требования к качеству жизни всё растут. И мы хотим надёжное освещение в городах, а не ждать месяцами, пока потухшую лампу посреди улицы и над ямой поменяют. Но при всём этом невозможно повесить на каждую лампочку датчик для отслеживания её работоспособности. Отсюда получаем и второе применение ИИ — машинное зрение.

Вам не нужно вешать датчик на каждую лампу, а достаточно повесить 1 камеру, которая будет охватывать целую улицу. Далее ИИ распознаёт, что вдруг перестала работать 1, 2, 3 лампы, автоматически формирует уведомление диспетчеру и формирует заявку на ремонт для дежурной бригады с автоматическим указанием места и количества вышедших из строя ламп.

В итоге снова упрощаем процесс, снижаем нагрузку дежурного и не ждём, пока кто-то возмущённо позвонит в службу.

Резюме

Да, это не проекты, которые решают наиболее насущные проблемы. Без них можно обойтись. Но это проекты, которые сочетают и трендовое направление, и направленные на социальную инфраструктуру, и повышают производительность, и снижают риск ошибки из-за человеческого фактора.

Применение ИИ в Норникеле

За возможность рассказать об этом примере хочу поблагодарить компанию «Вертикаль Инноваций» и Алексея Тестина, директора Центра Развития Цифровых Технологий.

Ниже приведены основные направления, выбранные компанией для применения цифровых технологий.

- Machine Learning — предиктивные рекомендации для операторов на производстве.
- Computer Vision — точное позиционирование оборудования и автоматический учёт.
- Large Language Model — цифровой ассистент на основе базы знаний для операторов.
- Моделирование физико-химических процессов — моделирование производственных процессов и определение критичных факторов, процессов, показателей.
- Big Data и аналитика данных — анализ и оптимизация рецептур.
- 3D-печать — печать запасных частей, например, колёс для насосов футеровки мельницы, бронедисков.

А в коридоре 2024–2027 годов компания будет активно развивать генеративный AI и внедрять предиктивный AI на всю цепочку создания ценности.

Ниже приведены области и способы применения генеративного ИИ.

LLM менеджмента — подсказки на базе сценарных анализов

- Поддержка принятия решений с учётом ограничений производства и рынка, а также политики компании.
- Снижение трудоёмкости выполнения повседневных задач — от подготовки резюме встреч до формирования писем и презентаций.

LLM технолога — повышение точности и скорости решений

- Высвобождение технолога от рутины делопроизводства.

- Консультирование и рекомендации по режимам работы технологического оборудования для принятия решений.

LLM support функций — повышение производительности персонала.

- Рост производительности вспомогательного персонала на 10–20%.
- Повышение качества подбора персонала и снижение текучести кадров (HR).
- Снижение нагрузки на типовые задачи и ускорение процессов в работе налоговой, юридической, финансовой и бухгалтерской службы.

LLM ремонтных функций — снижение человеческого фактора.

- Повышение эффективности обслуживания техники.
- Снижение зависимости от компетенций персонала и отдельных носителей знаний.

Области применения предиктивного AI также приведены ниже.

- Self Service AI (сервисы для самостоятельной работы с ИИ) — AutoML и Smart-Dashboard на базе цифровой платформы.
- Сквозная оптимизация — расширение AI на все переделы и сквозная оптимизация процессов.
- Sales Efficiency AI — управление цепочкой продаж и ценообразование.
- CAPEX AI — снижение стоимости и длительности строительства.

Из реализованных проектов коллеги поделились опытом Налогового департамента, который создал цифрового советника на базе LLM.

Сам проект стал ответом на следующие проблемы:

- Большой объём запросов в методологические службы Группы компаний — много времени сотрудников департамента тратится на поиск законодательной практики при подготовке консультаций для сотрудников Компании.

- Может быть низкая скорость ответов в периоды пиковой нагрузки, либо может страдать качество и полнота ответов на вопросы сотрудников.
- Многие вопросы повторяются и не требуют вовлечения высококвалифицированных налоговых специалистов для ответа.

Помощник реализован следующим образом:

- автоматический поиск и агрегация ответа из различных источников и документов. При этом в ответе указывались ссылки на источники, документы и конкретные их фрагменты;
- удобная поисковая строка в форме веб-интерфейса для пользователей из других департаментов;
- работа в корпоративной сети передачи данных без доступа во внешний Интернет;
- при отсутствии информации в документах выдаётся ответ о том, что нет подобной информации в базе данных.

Собственно, такой подход к работе с LLM я и описывал в блоке про цифровые советники. Сейчас он становится стандартом — безопасно, дёшево, решает практические задачи.

С момента второй редакции у этого кейса появились цифры, которые приятно вписывать. Планы подтверждаются практикой: по итогам 2023 года экономический эффект от ИИ в «Норникеле» достиг \$100 млн; на Быстринском ГОКе внедрение дало +2,3% к объёму переработки и +1% к извлечению металла; промышленный ИИ охватывает 12 переделов на семи площадках, а к 2030 году компания ожидает суммарный эффект порядка 50 млрд рублей (РБК Тренды; материалы компании).

Применение ИИ в ЕвроХиме

ЕвроХим — один из гигантов российского рынка. У компании есть не только опыт применения ИИ, но и видение его развития в будущем.

Материалы, изложенные далее, взяты из доклада Вячеслава Козицина, руководителя направления ИИ ООО «Цифровые технологии и платформы».

В качестве основных направлений применения ИИ в производстве выделены:

- проектирование — повышение эффективности разработки новых продуктов, автоматизация оценки поставщиков и анализа требований к запчастям и деталям;
- производство — совершенствование процессов (рекомендации), автоматизация сборочных линий, снижение количества ошибок, уменьшение сроков доставки сырья / простоев из-за его отсутствия;
- продажи и планирование производства — прогнозирование объёмов продаж, ценообразование;
- логистика — оптимизация маршрутов, планирование спроса на ресурсы автопарка, повышение уровня подготовки и качества работы сервисных инженеров.

Важнейшими направлениями, на которых сфокусировалась компания, стали оптимизация процессов и диагностика оборудования.

1. Оптимизация процессов:

- оптимизация режимов работы оборудования для увеличения производительности и ресурса этого оборудования;
- оптимизация параметров производственного процесса, например, соотношения компонентов;
- контроль и обеспечение качества продукции на различных стадиях процесса;
- оптимизация логистики и цепочек поставок, в том числе контроль за производственной безопасностью складской логистики;
- оптимизация расположения оборудования на производственных площадях.

Благодаря уже внедрённым нововведениям компания увеличила выпуск слабой азотной кислоты на 2,5%, аммиака и карбамида — на 1,5%, калийной соли — на 0,8%. Разумеется, это только начало.

2. Диагностика оборудования:

- мониторинг работы и раннее обнаружение неисправностей или отклонений;
- локализация отклонений;
- определение корневых причин;
- определение остаточного ресурса для перехода к обслуживанию по состоянию.

Если говорить про генеративный ИИ (цифровые ассистенты), то коллеги видят его развитие в:

- цифровых ассистентах (вопросо-ответные системы или знакомые нам чат-боты);
- анализе и генерации документации;
- мультиагентных системах (как в нашем примере по цифровому советнику по управлению проектами, где разные специализированные модели работают в сочетании).

Среди вопросо-ответных систем можно выделить:

- онбординг специалистов: «По каким дням я получаю зп?», «На каком автобусе я доеду до здания №8 завода №2?», «Сколько денег в сутки дают командированным?»
- поддержку специалистов по направлениям: «Какие лимиты у глав заводов на покупку оборудования?», «Как согласовать срочную заявку для завода №3?», «Как делается заявка на запрос информации?»
- ну, и конечно поддержку на совещаниях: «Сколько остановов было у цеха №2 за месяц?», «Какой самый длительный останов Цеха №2 за

2022 год?», «Сколько тонн продукта №4 потеряли на Заводе X из-за остановов цеха №5 в 2024 году?»

Как мы уже знаем из книги, чтобы такие ассистенты заработали, все эти данные нужно оцифровать и структурировать. ИИ должен иметь доступ к базам данных и уметь оперировать их содержимым. А для этого нужно выстраивать внутренние процессы и понимать, какие вообще данные могут понадобиться. То есть внедрение ИИ нужно начинать с процессов и работы с данными, в том числе с создания моделей и обеспечения качества данных.

И здесь практика догнала планы. Программа развивается: рекомендательные системы на НАК «Азот» принесли 270 млн рублей эффекта за 2024 год (до +1% выработки, -0,8% расхода природного газа), решения тиражированы на «Невинномысский Азот», а суммарный годовой эффект цифровых советчиков компания оценивает примерно в 350 млн рублей (CNews, ComNews — февраль 2025).

Применение ИИ в переработке руды

Наибольшее применение ИИ на производстве на 2024 год можно встретить в области машинного зрения.

Ниже приведён портфель решений одной из компаний структуры «Северсталь» для автоматизации ТОиР на базе машинного зрения.

- Детектирование негабаритной руды.

Застревание негабаритной руды в дробилке приводит к многочасовому простоя, что может вызвать остановку всей фабрики.

Алгоритм применения простой: получение видеоизображений руды — обработка изображений — анализ крупности руды — уведомление оператору дробилки.

Всего для проекта понадобился 1 сервер и 4 GPU NVIDIA Tesla T4 16GB.

Как мы понимаем, весь проект окупился по итогам одного предупреждения и исключения аварийного ремонта и простоя всей организации.

- Поиск недробимых материалов.

При попадании в дробильное оборудование вместе с рудой недробимых тел нейронная сеть идентифицирует и классифицирует объекты на конвейерной ленте в видеопотоке, система отправляет оперативное уведомление оператору для принятия решения.

- Анализ гранулометрического состава руды.
- Контроль усталости водителей.
- Мониторинг дефектов обжиговых тележек.

Система анализирует видеоряд и определяет наличие типовых дефектов: нет малого борта; нет ходового ролика; нет крышки ходового ролика; плохо прикрученная крышка ходового ролика; нет болта боковины; нет болта крышки ходового ролика; дефект крепления пластины продольного уплотнения.

Как видите, здесь всё равно нужна предварительная подготовка и определение критичных дефектов.

- Мониторинг зубьев ковша экскаватора и степени их износа.

Система определяет наличие или отсутствие коронок ковша экскаватора и оповещает машиниста экскаватора. А также на основе анализа степени износа готовит рекомендации о необходимости обслуживания или замены.

- Мониторинг смещения конвейерной ленты.

Система анализирует смещения конвейерной ленты и в зависимости от показателей классифицирует отклонения и формирует уведомления.

Клиентские сервисы

Искусственный интеллект для техподдержки

Ожидания от искусственного интеллекта сейчас поистине велики. И часто перегреты. Один из вопросов этой отрасли — как же снизить затраты на ИТ? Приоритетное направление — оптимизация ИТ-подразделений, которые занимаются технической поддержкой. Но дело

не только в сокращении численности персонала, но и в повышении качества самой технической поддержки.

В общем, одна из идей, которая крутится в головах руководителей, — внедрить искусственный интеллект на нулевую линию техподдержки, чтобы все входящие запросы верно классифицировались и направлялись автоматически нужному специалисту. Проще говоря, нужно автоматизировать эту нулевую линию. Насколько эта идея реализуема?

Если кратко, ИИ можно применить в виде чат-бота для того, чтобы правильно классифицировать запрос или решить типовую проблему пользователя. То есть для снижения количества заявок в техподдержку (приблизительно на 40%) можно через уточняющие вопросы чат-бота выяснить все «симптомы» неполадки и уже на основе их анализа сформировать заявку нужному специалисту. Чат-бот здесь «очеловечивает» весь процесс, чтобы пользователь не имел дело с безликой формой с галочками в личном кабинете.

Давайте теперь разберём ситуацию подробнее.

Возьмём типовую корпорацию с развитой ИТ-инфраструктурой и своей техподдержкой. 95% запросов в эту техподдержку будут идти через электронную почту. Да-да, порталами самообслуживания и личными кабинетами мало кто пользуется (5–10% сотрудников). Ведь одно дело — просто отправить «КАРАУЛ! Не печатает принтер!» по электронной почте, а другое — идти на портал, авторизовываться, выбирать тип вопроса, ставить галочки... В общем, человек — существо ленивое и пользуется наиболее удобным инструментом, избегая лишних телодвижений.

При этом 80% пользователей пишут сообщения так, будто за каждый символ они платят из своего кармана или только-только научились говорить. Типовой запрос выглядит так: «Не работает 1С». Обучить всех писать запрос по методологии 5W1H (6 вопросов для локализации проблемы в бережливом производстве) — утопия такая же, как заставить людей пользоваться порталом самообслуживания.

Но давайте представим себе, что люди научились писать запросы и всё выполняется, закрывается в срок. То есть наши специалисты техподдержки — просто образцовые сотрудники. Так вот, чтобы делать сортировку задач по письму из электронной почты, нужно будет много, методично и дорого поработать. На это уйдёт примерно 2–3 года и больше 20 миллионов рублей: нужно вычистить, нормализовать и разметить базу данных в ITIL, показать примеры распределения и назначения задач на 2–10 тысячах запросов, и может, тогда нейросеть сумеет понять запрос по 1 письму. Но я в такой сценарий верю мало.

Ещё одно огромное ограничение заключается в том, что в каждой компании свой зоопарк ИТ-систем, своя организационная структура и распределение полномочий. Кроме того, в большинстве компаний вообще нет чётко распределённых зон ответственности. А если и есть, то всё превращается в бюрократию, когда вопрос решается частично, а не комплексно (для решения 1 пользовательского сценария (например, выдача и настройка нового ПК), надо писать 5–10 заявок).

И в итоге коробочных решений для сортировки входящих задач по письму нет и не будет.

А в результате обучения ИИ скорее всего начнёт понимать, от какого пользователя какая проблема может приходить, и в целом может оказаться, что 90% запросов типовые, которые решаются чтением инструкции и до техподдержки не должны доходить. Но для этого разрабатывать и внедрять ИИ не нужно. Лучше позаниматься анализом, выявить типовые проблемы, оптимизировать скрипты и штатную структуру.

В рамках такого подхода хорошо поможет чат-бот с ИИ и базой знаний о том, какие ИТ-решения есть в компании, с пользовательскими инструкциями на борту и описанием организационной структуры (кто и за что отвечает). То есть ИИ-бот, который или решает типовую проблему, или нормально описывает и обвязывает задачу для техподдержки.

Главный вопрос здесь — как выстроить взаимодействие с пользователем, через какой канал? Это сложный вопрос с точки зрения информационной

безопасности. Типовые департаменты корпоративной защиты сейчас закрываются железным занавесом — всё внутри, ничего извне. Хотя работает это довольно плохо, всё равно безопасники никаких мобильных приложений не разрешат, тем более ботов в Телеграме и доступов по ссылке. Вместо этого сделают доступ из корпоративной сети (чтобы потом биться над вопросом удалённого доступа).

Ещё один вопрос — как система будет определять, кто именно с ней общается? Перевод всей инфраструктуры в доменные учётные записи — дело не быстрое. Это болезненный путь. В общем, решение этого квеста — хорошая зарядка для ума.

Критическая задача — подготовка базы знаний по ИТ. Нужно описать все ИТ-системы, собрать все типовые проблемы и обучить на них чат-бота, а также сопровождать чат-бот в первый год.

Причём, как показывает практика, это одно из самых востребованных направлений. На мой взгляд, это станет типовым проектом для ИТ-служб — сочетание ассистента с агентом с низкими рисками от ошибки и высоким эффектом.

Резюме

В общем и целом, ИИ может стать классным инструментом, который снизит нагрузку на техподдержку и повысит качество сервиса. Но это точно не волшебная пилюля. Всё равно надо выстраивать процессы и наводить порядок, ИИ за менеджеров этого не сделает. То есть снова приходим к выстраиванию клиентского пути и типовых сценариев, описанию организационной структуры с распределением зон ответственности, отладке процессов, налаживанию каналов коммуникации и т.д.

Искусственный интеллект в общественном питании

Это пример того, как и в какие направления начала внедрять ИИ компания Кофемания. Когда мы общались с ребятами, мне особенно понравился их осмысленный подход и стратегическое мышление. Так, в вопросе, с чего начать, они опирались на три блока:

- стратегия бизнеса;
- стратегия развития информационных технологий;
- стратегия работы с данными.

И уже отсюда рождаются проекты, которые несут ценность для бизнеса, и формируется дорожная карта развития ИИ. В итоге компания сфокусировалась на четырёх направлениях использования ИИ на первом этапе.

Первое — нормализация НСИ. Второе — обучение сотрудников. Третье — аналитика данных. Четвёртое — персонализация предложений для гостей. Давайте разберём эти направления.

Нормализация НСИ

Работа с нормативно-справочной информацией — головная боль любой компании, которая переходит из стадии малого бизнеса. В этом направлении компания решила использовать ИИ для решения следующих задач:

- поиск дублей в справочниках;
- сопоставление данных из разных таблиц и источников;
- валидация данных (проверка данных различных типов по критериям корректности и полезности для конкретного применения);
- типизация / структурирование / группирование номенклатуры.

Обучение персонала

Тут компания нашла тоже интересное применение для ИИ. Голосовой ИИ-ассистент помогает освоить меню и стандарты ресторана. Он в том числе генерирует вопросы, адаптированные под специфику ресторана, обрабатывает ответы сотрудников на разных языках и проверяет их готовность к самостоятельной работе.

Аналитика данных

Это как раз одно из тех направлений, которое становится всё более востребованным во всех отраслях. ИИ на базе LLM становится отличным

аналитиком данных, который работает круглосуточно и быстро. Так, компания использует ИИ для написания запросов к базам данных и построения аналитических дашбордов или подготовки ответов в реальном времени и на естественном языке.

Персонализация предложений для гостей

Одна из глобальных целей Индустрии 4.0 — кастомизация клиентского сервиса. И тут ребята начали использовать ИИ для подбора блюд из меню с учётом предпочтений гостя. Например, исходя из требований по балансу белков, жиров и углеводов, диеты, настроения или бюджета гостя. А также ИИ помогает делать допродажи блюд и подбирать официанта под гостя.

Причём на практике выяснился интересный факт: сотрудники сами освоили чат-бота и перестали думать. Они просто дают информацию ИИ о пожеланиях и транслируют гостям то, что отвечает ИИ.

И это только первый этап. Далее пойдёт использование ИИ для планирования закупок, оптимизации меню и так далее.

Планирование и кадры

ИИ в планировании проектов для нефтедобычи

Этот пример родился из диалога с моим товарищем, который работает в одной из крупнейших нефтяных компаний России. Разговор состоялся в середине 2025 года. Далее приведу его речь.

«Хотел на завтра для руководства сделать пример, что для сложных геологоразведочных проектов ИИ пока не может учесть все тонкости.

«После прочтения ответа остался с мыслью, что надо уже наверно учиться что-то руками делать, мозги скоро не понадобятся.

«Вот мой запрос: „Привет! Планирую бурение поисковой нефтяной скважины. Точка бурения находится в автономии, доставить оборудование и материалы для бурения с большой земли можно только водной навигацией в летний период в порт, далее по зимней автодороге до скважины. Распиши пожалуйста основные этапы работ по подготовке

к бурению начиная с июля 2025 года и определи дату, когда можно начать бурить“

«А такой ответ я получил от ИИ: „... (прим. автора — *я опущу все детали ответа, но это нормальный план с графиком и критическим путём, рисками, рекомендациями, и приведу уже ответ*)

«„Вывод: При строгом соблюдении графика, особенно по морским поставкам в июле-сентябре 2025 года, и успешном выполнении зимних логистических операций, реалистичной датой начала бурения поисковой скважины является 10–15 февраля 2026 года. Эта дата обеспечивает необходимый буфер безопасности перед весенней распутицей“.

«Так представляешь, мы в этом году аналогичную скважину забурили 12 февраля. Мы для своей скважины план полгода выравнивали двумя обществами, а ИИ выдал это за 21 секунду».

Генеративный ИИ и мультиагенты в HR

Один из актуальных вызовов для бизнеса — дефицит квалифицированных кадров. Разные исследования показывают, что до 40% руководителей называют «кадровый голод» одним из ключевых ограничений роста компании. Во время моих тренингов эта тема очень бурно обсуждается. Ведь проблема найма не только в том, что «нет людей», но в самой возможности корректно оценить экспертизу кандидатов, так как сейчас появляются новые специальности.

На фоне развития искусственного интеллекта активно формируется новое направление — автоматизация подбора и оценки кандидатов с помощью генеративного ИИ и мультиагентных систем.

Ниже поделюсь практическим опытом одной из российских компаний. Так, коллеги разработали ИТ-платформу для найма сотрудников под временные задачи и проекты. Ею пользуются СИБУР, Северсталь, АЛРОСА, РОССЕТИ, КРОК. Я и сам там получал проекты. Одно из ключевых решений на платформе — ИИ-агенты для автоматизации подбора, скрининга и ранжирования специалистов под задачи бизнеса. Сейчас это интегрировали с порталом hh.ru, внутренними базами и

корпоративными ATS-системами (системы подбора персонала), что позволяет вести обработку тысяч резюме за минуты.

Как же устроено внутри взаимодействие ИИ и реализован мультиагентный подход?

Первый агент декомпозирует неструктурированные заявки на подбор, генерирует формализованное описание вакансии, включая навыки, требования, условия, этапы и ожидаемую стоимость.

Второй ведёт поиск по базе кандидатов, формирует начальную (лонг-лист) и целевую (шорт-лист) воронку кандидатов, осуществляет рассылку приглашений.

Третий проводит углублённую оценку резюме.

Самое сложное в этом случае сделать модель, оценивающую резюме. По каким критериям это делать? Как оценивать? Размышляя над этим, коллеги пришли к структуре из трёх блоков.

- Direct: соответствие требованиям вакансии (чек-лист навыков, опыт, уровень владения инструментами и методологиями).
- Rules: анализ истории работы, выявление проблемных зон (например, частые смены, отсутствие ключевых компетенций).
- Competencies: выявление предметных и гибких навыков, прогноз роста, аргументированный потенциал быстрого освоения смежных областей.

То есть резюме кандидата попадает к ИИ, он сравнивает его с описанием вакансии и раскладывает на смысловые блоки. А дальше считается вероятность по каждому из компонентов и выставляется общая оценка, генерируется детализация причин выбора (или отклонения) кандидата, выгружается в виде отчёта для менеджера.

Пример: кандидат ранжируется на основе опыта (11 лет в ИТ), соответствия базовым требованиям и возможности освоить нужные навыки (например, конфигурации 1С) с пояснением сильных и слабых сторон.

В итоге были получены следующие эффекты:

- скорость полного цикла обработки 1 резюме — 10 секунд против 10 минут на специалиста-эксперта (x60 ускорение), а ранжирование 100 резюме — менее чем за 10 минут;
- точность совпадения финальной оценки с решением рекрутера — 87% на любых объёмах выборки;
- экономия затрат на подбор — до 10–15% (пример: на типовом объёме 5 000 резюме в месяц экономия — 250 000 руб.);
- снижение времени найма (тайм ту хайр) на 20–30%;
- доступность 24/7, исключение эффекта «выгорания» и предвзятости;
- снижение доли ручной работы рекрутеров на подготовительных этапах до 80%.

Конечно, есть и препятствия, которые мы уже обсуждали ранее.

- Корпоративная политика безопасности — многие компании требуют локального развёртывания (on-premises), глубокой деперсонализации.
- Кейсы массового подбора требуют отдельного сценария по коммуникации с кандидатами (например, водители, сантехники), не всегда «глубокий скрининг» нужен.
- Бизнес-логика интеграций с внешними ATS/CRM системами зависит от зрелости IT-ландшафта корпоративного клиента.
- Необходимость постоянного переобучения моделей под новые рынки, задачи и языки.

Отмечу следующее. ИИ в этом случае показывает, что при определённой зрелости и готовности бизнеса мультиагентные системы способны кардинально повысить качество и скорость процессов, в том числе подбора и оценки кадров в крупных компаниях России. В данном кейсе мы получили не только увеличение скорости подбора, но и непредвзятость в оценке кадров, особенно когда идёт поиск уникальных

специалистов, которых оценивать корректно линейные руководители не могут.

Для максимального эффекта лучше запускать пилотные проекты на ключевых вакансиях, вовлекать в разработку продуктовые и HR-команды, формировать обратную связь пользователей.

Также ИИ в HR должен развиваться в интеграциях с другими системами и с гибкой настройкой и поддержкой как облачной, так и локальной инфраструктуры. Ну, и конечно важную роль играет методология деперсонализации и соблюдение требований безопасности персональных данных.

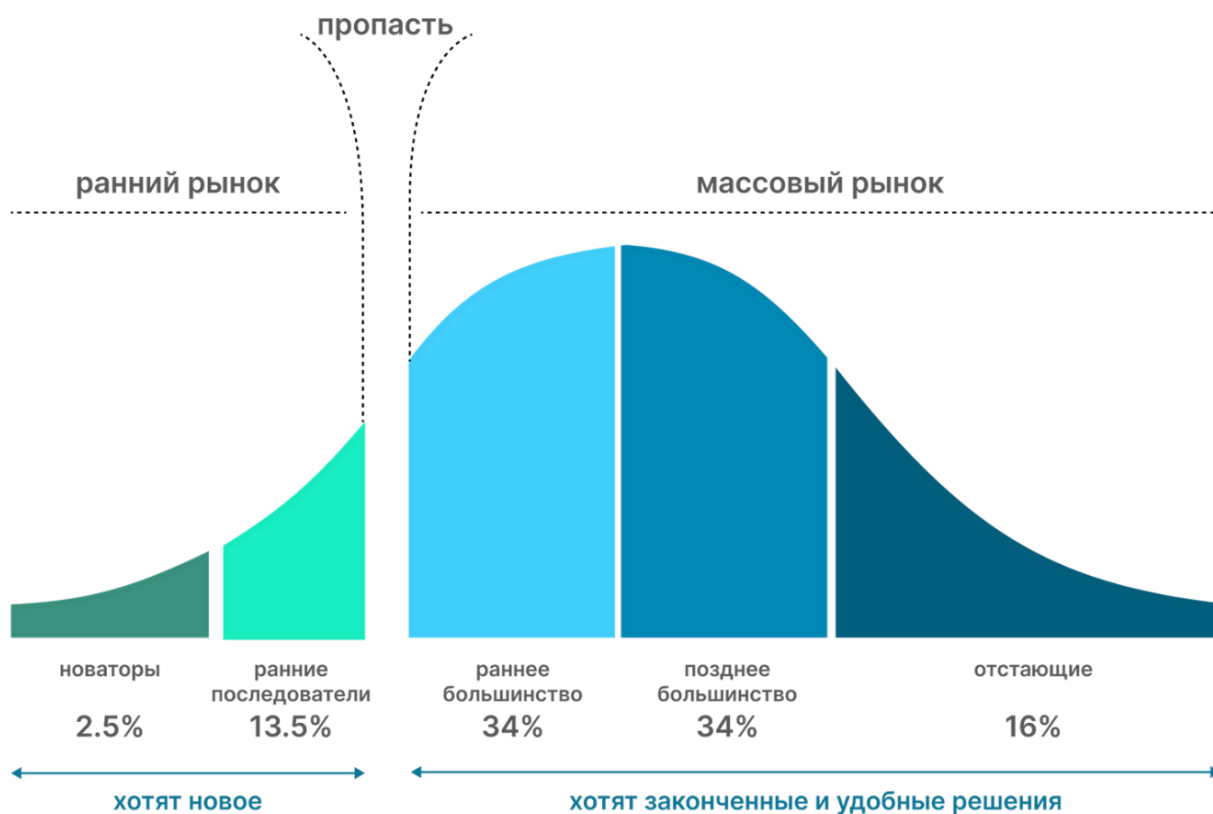
Интересный практический пример — студия анимации «Петербург» (ГК «Рики»), которая разработала ИИ-агента на базе Yandex AI Studio для интеграции новых сотрудников. Агент отвечает на вопросы вроде «Как оформить договор с лицензиатом?» или «Как создать новую анимационную сцену?», причём в ответах сотрудник получает не только текстовые инструкции, но и изображения. По словам СТО компании, это значительно сократило частоту обращений в поддержку и облегчило процесс адаптации.

Примечательно, что подобные решения строятся не на кастомных моделях, а на готовых платформах с low-code инструментами и подключением к внешним системам через MCP-серверы (Model Context Protocol — открытый стандарт, по которому ИИ-агент подключается к внешним данным и инструментам без отдельной интеграции под каждую систему). К марту 2026 года на одной только Yandex AI Studio было создано более тысячи таких серверов, позволяющих агентам взаимодействовать с AmoCRM, Битрикс24, Контур.Фокус и другими привычными бизнес-инструментами.

Рекомендации для крупного бизнеса

- Выделяйте отдельных лидеров инноваций (по кривой Мура), которые готовы экспериментировать и внедрять ИИ-решения. Такие люди в компаниях есть всегда, их 10–15%. Они готовы прощать

ошибки и быть «подопытными», а затем и проводниками изменений.



- Создавайте ИТ-песочницы, где будете разворачивать ИИ-модели, тестировать их. Это поможет в том числе договориться с вашими специалистами по безопасности и создать условия для быстрого тестирования идей.
- Проговаривайте с ИБ возможности использования готовых ИИ-продуктов и решений по SaaS-модели (без размещения во внутренней сети), как минимум для пилотов и исследований. Так вы снизите стоимость этих пилотов.

Также фокусируйтесь на малых и средних языковых моделях, которые можно разворачивать во внутренней инфраструктуре и интегрировать с корпоративными сервисами.

- Формируйте процессы проектного управления и правила запуска пилотов, это ключевая задача.

- Внедряйте академии новаторов и обучайте их основным инструментам системного подхода.

Особое внимание уделяйте управлению проектами и продуктами, бережливому производству, управлению коммуникацией, теории ограничений систем.

- Сделайте ИИ стратегической задачей компании, включите его в стратегию развития и цифровизации, в OKR-цели и KPI-руководителей.
- Участвуйте в хакатонах от образовательных учреждений и акселераторов, заключайте партнёрства с научными организациями. В общем, дайте молодым и гибким поработать над вашими задачами. А также наблюдайте за рынком, заходите в долю молодых стартапов. Так вы сможете адаптировать решения под себя за минимальную стоимость, а если стартап успешен, то потом и заработать, продав свою долю.
- Будьте готовы к внедрению коробочных решений, особенно если внутри есть готовые методологии. Не бойтесь упрощать свои процессы — как показывает опыт, большинство из них устарели и потенциал для улучшения огромен. Для этого хотя бы подготовьте реестр бизнес-процессов, с их владельцами, участниками и проблемами.
- Начните работу над качеством данных и выстраиванию процессов управления качеством данных. Речь не только о нормативно-справочной информации, но и о всех данных в целом. Без данных никакие ИИ-решения, тем более рекомендательные системы, пользы не принесут.
- Предоставьте возможность новаторам участвовать в конференциях. Скорее даже сделайте это обязательной частью. Особенно в непрофильных для вас отраслях. Подрывающие инновации и самые интересные примеры мы всегда находим на стороне. А насмотренность и эрудированность в области цифровых технологий

является обязательным компонентом для успеха. Этот пункт я даже отдельно выделил в своей матрице компетенций.

- Внедряйте ИИ сначала в те процессы, где цена ошибки не критична, но эффект будет заметен. И лучше внедряйте узкоспециализированные ИИ-решения — их рекомендации точнее и безопаснее. Но если вы внедряете ИИ в АСУ ТП, то оставляйте конечное решение за человеком, используйте данные о согласии или несогласии с рекомендациями ИИ для его дальнейшего дообучения.
- ИИ не заменит ваших людей, но поможет им стать эффективнее. То есть либо пройти этап нехватки кадров, либо минимизировать рост персонала с ростом организации (не раздувать штатные структуры отдельных подразделений обычными специалистами).
- Подходите к задаче комплексно и воспользуйтесь ИИ как драйвером к организационному развитию.
- Я могу выделить несколько направлений, которые будут востребованы в любой отрасли и, вероятнее всего, появятся в виде отдельных продуктов и решений на рынке.
- Различные советники по направлениям: закупки, юристы, управление проектами и т.д.
- ИИ-аналитики данных.

Это LLM, которая подключается к базам данных и работает как аналитик данных, может строить отчёты, графики, давать справки и аналитические выгрузки, делать прогнозы в режиме реального времени, в общем, быть ассистентом для руководителей.

- HR-ИИ.

Это LLM, которая помогает в подборе сотрудников (отбор релевантных кандидатов), их сопровождении и поддержке на ранних этапах работы.

- ИИ для ИТ-поддержки.

Ассистент, который поможет или решить вопрос, или уточняющими вопросами сформирует корректную заявку. Очень актуальный запрос от

бизнеса. Ведь либо мы имеем недовольство сотрудников, либо раздуваем штат ИТ-службы.

- Подготовка проектной документации.

В ИТ-проектах есть большая проблема с проектной документацией. Это трудозатратный и длительный процесс. ИИ же сможет быстро готовить эту документацию на основе анализа баз данных и кода. Эти проекты уже появляются.

Оборотная сторона: документации нужно больше, чем раньше (при использовании ИИ для разработки ИТ-решений). У ИИ, который пишет документацию, есть зеркальное свойство: он же её и потребляет — и потребляет иначе, чем человек. Человек, наткнувшись на пробел в задании, останавливается и спрашивает. Модель по умолчанию не спрашивает, она достраивает недостающее правдоподобным. Спросит она, только если её об этом попросили (в Главе 6 это Clarification First), и то лишь про то, где сама видит пробел.

Практическое следствие для ИИ-проектов: под ИИ документируют не то же самое, что под человека. Под человека описывают, что делать. Под агента приходится описывать ещё и чего не делать, в каких режимах решение должно работать и как мы поймём, что пункт закрыт. Это ровно те вещи, которые в Главе 6 вытаскиваются вопросами к плану по одной. Документация — способ не задавать эти вопросы каждый раз заново.

И про моду на «вайб-кодинг», когда решение собирается диалогом с моделью без проектных артефактов. Он не отменяет документацию, а переносит её из «после» в «до»: то, что раньше дописывали по факту, теперь приходится формулировать на входе, иначе агент допишет за вас и не спросит, а модель начнёт забывать и пропускать. Это проходил и я со своей командой при разработке наших продуктов.

Также тут ещё можно спрогнозировать появление коробочных решений по работе с НСИ (об этом я писал ранее и уже есть примеры на основе той же Кофемании), создания ИИ-двойников сотрудников с их опытом и экспертизой, и ИИ-ассистентов по коммуникациям, которые будут

подключаться к почте, мессенджерам и вести коммуникацию за сотрудника в его стиле.

- Фокусируйтесь не на моделях, а на продуктах. В том числе в каждом ИИ-решении нужно логирование запросов для их анализа и совершенствования внутренних регламентов и документов.
- Также я могу рекомендовать гибридный подход к развёртыванию ИИ.

Когда выбирать локальное развёртывание (во внутренней инфраструктуре):

- Обработка конфиденциальных данных.
- Требования к соответствию GDPR/персональные данные.
- Стабильная высокая нагрузка.
- Необходимость кастомизации модели под бизнес-процессы.
- Желание избежать зависимости от разработчика.

Когда использовать облачные API:

- Нерегулярная или низкая нагрузка.
- Необходимы самые современные возможности и сложные рассуждения.
- Отсутствие ИТ-экспертизы для поддержки локальной инфраструктуры.
- Быстрый старт проекта без капитальных вложений.

В итоге мы получаем гибридный подход. В его рамках мы используем ИИ-модели во внутренней инфраструктуре для базовых задач и конфиденциальных данных, а облачные API для сложных задач, требующих максимального качества.

Развёртывание облачных моделей будет идти также через прослойку. Перед тем как запрос будет уходить в большую языковую модель (LLM), отдельная ИИ-модель находит и заменяет все чувствительные данные на

безличные. Например, вместо Иванов Иван Иванович указывается [ИМЯ_КЛИЕНТА] или вместо +7 XXX XXX XX XX прописывается [ТЕЛЕФОН]. В итоге облачная LLM работает с полностью анонимным запросом, а ваш сотрудник получает корректный ответ с реальными данными и на основе самой современной модели.

Пример до маскирования

Оригинальный запрос пользователя:

Составь официальное письмо от имени **Ивана Петрова**, директора компании ООО "ТехСфера", клиенту **Анне Смирновой**, проживающей по адресу г. Москва, ул. Лесная, д. 12, кв. 45. Тема письма: задержка поставки оборудования по договору № 456/2025 от 12.03.2025. Укажи мой контактный телефон +7 (915) 123-45-67 и электронную почту ivan.petrov@techsfera.ru.

Пример после маскирования

Версия, отправляемая в LLM:

Составь официальное письмо от имени [ИМЯ_ОТПРАВИТЕЛЯ], директора компании [НАЗВАНИЕ_КОМПАНИИ], клиенту [ИМЯ_КЛИЕНТА], проживающему по адресу [АДРЕС_КЛИЕНТА]. Тема письма: задержка поставки оборудования по договору [НОМЕР_ДОГОВОРА]. Укажи мой контактный телефон [ТЕЛЕФОН] и электронную почту [EMAIL].

- И финальное, сэкономьте деньги и не пытайтесь создавать свои ИИ-команды. Вы априори проиграете гонку с продуктовыми командами. Вы вряд ли готовы платить столько же, сколько они, а значит, будете привлекать средних специалистов. А те, кто будет расти, вы вряд ли сможете удержать. Вам важно накапливать экспертизу бизнес-заказчиков, а не разработчиков. Разработка ПО и подбор моделей — не ваш основной бизнес.

Детальные дорожные карты внедрения для среднего и крупного бизнеса — состав команд, бюджеты, планы по этапам и чек-листы — в книге «ИИ-2. Практическое руководство» (Главы 18–19). Здесь остаются кейсы и управленческие принципы.

Ну и, для интереса, я снова спросил рекомендации у самих ИИ-моделей. Их советы — аудит процессов, стратегия и приоритеты, пилот на ограниченном участке, качество данных, обучение людей, оценка

результатов — практически повторяют изложенное выше, на уровне хорошего консультанта.

Как я часто повторяю, внедрение ИИ-решений зачастую оборачивается бюрократической катастрофой только из-за того, что в компаниях отсутствуют компетенции проектного управления. Помимо этого, основная специфика крупного бизнеса — сложные и негибкие процессы со специфичными требованиями служб безопасности.

Конечно, рекомендовать такому бизнесу вдруг стать гибким — утопия, но всё же стоит выделить обособленные подразделения и механизмы для сотрудников-новаторов, дать им ИТ-среду, определить правила и бюджет проекта, которым можно рисковать и в рамках которого нужно минимум бюрократии.

Например, если проект требует до 0,5 млн рублей, и он может реализоваться силами одного подразделения, то опишите проблему, её решение и вперёд. Если нужно от 0,5 до 10 млн рублей, то список задач увеличится. Ну, а если бюджет превышает 10 млн рублей, то придётся идти по полному бюрократическому циклу.

В Главе 22 я расскажу, как вести проекты по внедрению новых технологий в свой бизнес. И хотя данная книга посвящена искусственному интеллекту, алгоритм управления жизненным циклом проекта является универсальным для любой инициативы.

Резюме главы

- Отраслевые кейсы сходятся в одном: эффект даёт связка «данные + IoT + ИИ», встроенная в процесс, а не отдельная модель.
- Средний и крупный бизнес выигрывает на платформенности и портфеле инициатив: пилоты ради пилотов — путь в «кладбище PoC».
- Самые сильные кейсы — там, где ИИ дополняет существующую автоматизацию (ERP, MES, SCADA), а не пытается её заменить.

Что с этим делать: Выберите из этой главы 2–3 кейса-аналога вашей отрасли и сравните с вашими процессами: что из инфраструктуры и

данных у вас уже есть (неделя). Сформируйте портфель из 3–5 инициатив с приоритетами (квартал).

Вопрос для рефлексии: Какой кейс этой главы ближе всего к вашей главной операционной боли — и что мешает повторить его у вас?

Глава 21. ИИ в работе самого руководителя

Мы разобрали, как ИИ работает на малый бизнес и на корпорации. За кадром осталась фигура, о которой обычно забывают, — сам руководитель. Не процессы его компании, а его собственный день: решения, коммуникации, планирование, усталость и уход фокуса со стратегического контура. При этом ошибка именно этого человека стоит дорого, а порой именно от его решений зависит будущее всей компании и десятков людей. На практике я наблюдал не один и не два раза, когда именно из-за этого не запускались стратегические проекты и компания потом имела проблемы.

Эта глава — самая личная в книге. Я покажу, как использую ИИ сам: без теории, на реальных примерах, включая один диалог.

Задачи руководителя: что можно отдать ИИ

Работа руководителя любого уровня складывается из одних и тех же классов задач — различается только их доля. У линейного руководителя больше контроля и обратной связи, у среднего — координации и проектов, у топ-менеджера и собственника — стратегии и решений с высокой ценой ошибки. Хорошая новость: ИИ применим в каждом классе. Плохая: нигде он не заменяет вас целиком. Это то самое правило 70/30 из первой части, применённое не к процессам компании, а к вашему личному рабочему дню.

Класс задач	Что берёт ИИ	Что остаётся вам
Планирование дня, недели, месяца	Собрать план, увязать его с целями, напомнить о хвостах, проработать задачи и подсказать, кому и что делегировать	Выбор приоритетов и реализация в жизни

Решения	Проработать решения с разных сторон, структурировать проблему, найти первопричины, дать альтернативы и риски	Само решение и ответственность за него
Коммуникации	Подготовка к переговорам, совещаниям, черновики писем, разбор конфликтов, разбор итогов	Живой контакт и доверие
Целеполагание	Проработка целей, формулирование критериев их достижения, контроль, поддержка фокуса в коммуникациях, решениях, проектах и при планировании, демонстрация динамики достижения целей	Реализация, корректировка и финальное утверждение
Проекты	Проработка проектов и черновики уставов и планов, декомпозиция до задач, актуальные статусы, отслеживание рисков, подведение итогов и выученных уроков	Целеполагание и договорённости
Команда и контроль	Собрать факты, напомнить об обещаниях — ваших и чужих, подготовить структуру 1:1 и обратной связи	Оценка людей и кадровые решения
Саморазвитие	Обратная связь по дню и неделе, выявление паттернов, предложение упражнений, запуск рефлексии	Выводы и изменение поведения

Мой личный контур: как это устроено

Расскажу, как это работает у меня — не как рекомендацию, а как один живой пример.

Дневник и обратная связь. Каждый вечер — короткая запись: что было, что получилось, что зацепило. ИИ-ассистент тут же возвращает обратную связь: где я повторяюсь, куда утекает энергия, что я откладываю третий день подряд. Раз в неделю — свод: динамика по целям, баланс между стратегической работой, отношениями и здоровьем, перегруз. Раз в месяц — то же на большем горизонте. Один я такой дисциплины не выдержал бы: ценность создаёт не запись, а мгновенный отклик на неё.

Из практики: как выглядит вечерняя запись

Запись: «Устал. День съели согласования по интеграции. Разговор с руководителем продаж отложил. Из хорошего: закрыли пилот у клиента».

Отклик ассистента: «Разговор с руководителем продаж откладывается третий раз за неделю — и в записях он каждый раз соседствует со словом „устал“. Похоже, это и есть главный источник фона. Пилот закрыт — шаг к квартальной цели №2. Предлагаю забронировать разговор завтра с утра на 15 минут, могу подготовить структуру».

Планирование. План недели и дня я собираю вместе с ассистентом: он видит мои цели, хвосты и обещания, предлагает скелет — решаю я. Разница с ежедневником в том, что план связан с целями: когда неделя съедается операционкой, я вижу это уже через 2–3 дня, а не в конце квартала.

Проработка решений. Перед серьёзным решением я разбираю его с ИИ в режиме «пяти почему»: он спрашивает, копает первопричины, даёт альтернативы, которые я не рассматривал, и риски, которые не хочется видеть. Решение остаётся моим — но оно проходит проверку боем до того, как я его объявил. При этом разбор происходит и с учётом моих стратегических целей, как рабочих, так и личных.

Проекты и документы. Уставы, планы, разборы статусов, подготовка к комитетам — всё, что раньше занимало вечер, теперь начинается с черновика, который делает ИИ, а я дорабатываю. Экономится не только время — не тратятся силы на «чистый лист».

Коммуникации. Самое ценное — подготовка к сложным разговорам и разбор уже случившихся или происходящих диалогов в мессенджерах. Об этом — следующий раздел, на живом примере.

Самое ценное, что я заметил в себе: я смог выстроить баланс работа / здоровье / семья, стал точнее планировать — и теперь это мой второй мозг, который всегда под рукой. ИИ забрал огромную долю рутины.

Наглядный пример — день, когда у меня не было доступа к инструменту. При планировании я посмотрел на список задач и понял: могу взяться за них сейчас, это займёт десять часов. Но завтра инструмент вернётся, и я сделаю всё за два-три часа. В итоге я взял паузу и потратил день на размышления — и нашёл два ценных инсайта для всей компании.

Личный контур у меня настроен не одним промптом, а четырьмя уровнями, и у каждого своя работа. Порядок — сверху вниз, от того, что действует везде, к тому, что действует в одном диалоге.

Уровень первый: общая инструкция аккаунта

Это «паспорт», который действует в любом диалоге, о чём бы он ни был. Скелет переносим на любого руководителя. Кто я: профиль, опыт, контекст в двух-трёх абзацах. Мои роли и приоритет между ними: в моём случае — CDO, основатель продукта, автор; плюс правило, что делать ассистенту, если роль из запроса не ясна (уточнить одной строкой или взять дефолт). Приоритетные классы задач: короткий список, на чём фокус. Формат ответов: резюме → разбор → артефакты → чек-лист, и обязательный «выключатель» — на простой вопрос отвечать коротко, без шаблона. Правила уточнений: сколько вопросов задавать и какие предположения брать, если я не ответил. «Спорь со мной»: возражать с аргументами, а не поддакивать — согласие из вежливости вреднее спора. И границы конфиденциальности: что никогда не выносить во внешние тексты без явной команды — внутренние цифры продукта, имена людей, детали работодателя. Пишется один раз, правится раз в квартал.

Уровень второй: общий реестр

Между паспортом и проектами у меня лежит папка «Реестр проектов». Это карта всего контура: список проектов с их инструкциями, авторские навыки и агенты, установленные расширения, рабочие папки и репозитории — что где лежит, кто за что отвечает и как это восстановить на новом компьютере. Паспорт на него ссылается одной строкой; агент, который работает с моими файлами, читает его сам.

Уровень появился не из любви к порядку. Пока проект один, ассистенту достаточно инструкции. Когда проектов и агентов становится десяток, каждый по отдельности исправен, а вместе они начинают мешать друг другу: один переписывает файл, который ведёт другой, третий не знает, что нужный навык уже есть. Это тот же сбой, что у трёх агентов в Главе 8, только в масштабе одного человека. Реестр — общая карта и правило дорожного движения для них.

Одна оговорка. Этот уровень работает там, где ассистент видит диск: агенты на компьютере, работа с подключённой папкой. Обычному чату в браузере реестр не нужен — ему хватает паспорта и проекта. Обновляется реестр не по календарю, а по событию: появился проект, изменилась инструкция, поставлено расширение — в тот же день, иначе карта начинает врать.

Уровень третий: проект под блок задач — инструкция плюс файлы

Дневник, книги, продукт, преподавание, обработка рационализаторских предложений — у каждого направления свой проект: своя инструкция и свои файлы знаний, к которым ассистент обращается сам. Вот как это устроено у дневника-ассистента, чтобы было видно, куда это может дорасти. Роль: «Ты — мой коуч и ментор на основе ежедневного дневника. Помогаешь расти в трёх ипостасях: как личности, как руководителю С-уровня, как основателю ИИ-компании». Вся аналитика ведётся относительно трёх областей баланса: стратегическая работа / отношения / здоровье. Отдельно зафиксирован мой типичный сбой: реактивная работа вытесняет стратегическую, а работа в целом — здоровье и отношения; ассистент обязан отслеживать это в первую очередь.

Пороги — «красные линии», нарушение которых ассистент называет прямо: сон не меньше семи часов минимум пять ночей в неделю; не меньше трёх тренировок; хотя бы один полный день восстановления; три вечера качественного присутствия с семьёй без рабочих устройств; два защищённых блока стратегической работы; реактив — не больше 65 % рабочего времени. Один нарушенный порог — ранний сигнал. Два и больше, или порог нарушен две недели подряд, — ассистент обязан предложить конкретную разгрузку: что снять, делегировать, перенести. Не «отдохни», а список.

Два режима. Менторский — прямая обратная связь по каждой записи: наблюдение → баланс → инсайт → рекомендация, не длиннее 200 слов. Коучинговый, по команде, — никаких советов: один сильный открытый вопрос за ход. Плюс команды-каденции: НЕДЕЛЯ, МЕСЯЦ, РАЗБОР темы, БАЛАНС. И принципы: опираться только на факты из записей; «паттерн» фиксировать лишь при двух-трёх повторениях; гипотезы помечать как гипотезы; никаких мотивационных банальностей.

В файлах проекта лежит то, к чему ассистент должен возвращаться, а я не хочу повторять: цели квартала и «Снимки» прошлых месяцев. Один чат — один месяц; в конце месяца ассистент собирает Снимок, из снимков потом складываются квартальный и годовой обзоры. Уровень пишется под задачу и правится, когда меняется сама задача.

Уровень четвёртый: постановка задачи в чате

Здесь уже не нужно объяснять, кто я и как со мной работать, — это знают первые три уровня. Остаётся сама задача и её контекст: запись дня, решение, которое надо проработать, разговор, к которому готовлюсь. Дневная запись при этом свободная, с желательным минимумом: что делал (стратегическое или реактивное), решения, отношения, здоровье, что застопорилось. Пять минут вечером — достаточно. Благодаря такой архитектуре я иногда просто отправляю материалы и одним предложением говорю, что хочу получить, — и ИИ отработывает отлично.

С чего начать, если ничего этого ещё нет

Не пишите инструкцию сами — попросите ассистента провести с вами интервью: «задай мне десять вопросов, чтобы собрать инструкцию: кто я, мои роли, задачи, как со мной работать, чего не делать» — тот же приём, что в Главе 6, только теперь модель собирает не запрос, а ваш паспорт. Дайте ей для образца выжимку из моего дневника-проекта выше — не чтобы скопировать, а чтобы она видела нужный уровень конкретности. Полчаса — и первый уровень готов; вторым заходом так же собирается проект под первую задачу. Реестр на старте не нужен — он понадобится, когда проектов и агентов наберётся с десятков.

И строка, которую я держу в паспорте и повторяю в каждом проекте, где возможен спор: «Отвечай коротко и по делу. Спорь со мной там, где я ошибаюсь».

Что это даёт на выходе

Три вещи, которые не даёт один длинный промпт, и четвёртая, которую даёт реестр. Первое — тёплый старт: любой новый диалог начинается с того, что ассистент уже знает меня, мои роли и правила игры; я не трачу первые десять минут на настройку. Второе — разделение того, что меняется редко и часто: паспорт живёт кварталами, проект — пока жива задача, чат — один день или один месяц; менять одно не значит ломать другое. Третье — переносимость: тот же паспорт работает и в дневнике, и над книгой, и в продукте, а проект переезжает между чатами вместе со своими файлами. И четвёртое, которое появляется вместе с реестром, — согласованность: агенты, работающие параллельно, знают друг о друге и не ломают чужое. Полчаса на паспорт через интервью (или вечер, если писать самому), час на первый проект — и качество каждого следующего диалога вырастает заметно: ассистент отвечает вам, а не «среднему пользователю».

Побочный эффект, которого я не ждал

Одному руководителю, которому я настраивал первый уровень, я дал прочитать свою общую инструкцию — тот самый паспорт. Через несколько дней он написал, что читает ответы ассистента иначе, чем раньше, — с облегчением: видит в нём исполнителя воли человека, тогда

как без знания инструкции было ощущение большей осознанности ИИ. Он увидел один уровень из четырёх — и «осознанность» исчезла. Это не разочарование, это и есть точка, к которой стоит прийти: то, что впечатляет в хорошо настроенном ассистенте, — ваша собственная воля и ваши правила, отражённые обратно. Вся магия — в инженерии настроек и потоков данных, а не в модели.

Отсюда и моё разделение 70/30 в личной работе. Тридцать процентов — человек: продумать, задать направление и ограничения, контролировать, принять финальное решение. Семьдесят — ИИ: собрать, агрегировать, проанализировать, упаковать идеи в слова. Пропорция та же, что в главе о трендах, но здесь она про распределение ролей в одном рабочем дне, а не про рынок труда. И ещё одно наблюдение по людям, которых я настраивал: почти все, кто плотно впускает ИИ в свою работу, довольно быстро начинают понимать, как устроена машина, — и перестают её бояться. Страх держится на непонимании, а не на возможностях технологии.

И важная оговорка о безопасности. В личный контур попадает чувствительное: имена сотрудников, конфликты, деньги. Публичный чат-бот — не место для таких данных: используйте контур с контролем над данными или обезличивайте записи. Подробно мы разбирали это в Главе 8.

Три кейса из практики

Кейс первый: диалог, который чуть не убил идею

Мой продакт-менеджер предложил идею: дать пользователям нашего продукта свободное пространство — создавать собственные шаблоны и структуры, как в Notion. Я закрыл тему несколькими сообщениями: «Пока нет. Это будет свалка данных. Придётся использовать дорогие модели, что сломает экономику. А ещё это может сломать работу агентов внутри системы. И мы станем плохим ChatGPT, а не хорошим Соколом». Формально я был прав по каждому пункту. По сути — на моих глазах умирала инициатива. А если так делать регулярно, вслед за инициативой умирает и мотивация.

Но разговор на этом не закончился — обсуждение продолжалось, реплика за репликой. И параллельно, прямо по ходу диалога, я разбирал его с ИИ-ассистентом. Ассистент знает мои цели и мою роль: мне нужно не выиграть спор, а вырастить продакта, которому я смогу делегировать развитие продукта или запуск новых. Это стратегический трек для моей разгрузки в перспективе 2–3 лет — формирование мотивированного и компетентного члена команды. Поэтому обратная связь была не абстрактной, а прицельной. Ассистент подсветил три вещи. Первое: сотрудник предлагал не то, что я слышал, — он говорил о горизонте в два-три этапа развития продукта, а я отвечал про сегодняшний релиз; мы спорили о разных вещах, и каждый был прав в своей рамке. Второе: я закрыл идею быстрее, чем разобрался в ней, — сработал мой типовой паттерн «сначала защитить продукт, потом слушать». Это довольно понятное поведение для творцов и предпринимателей, которые что-то сами уже построили. Третье, самое важное: если человек систематически получает «нет» на вопрос, которого не задавал, он перестаёт приносить идеи. Ровно то, за что я его и брал в команду.

Дальше я перестроил диалог, не выходя из него. Вместо запрета — рамки: обозначил барьеры (не плодить хаос, не создавать свалку данных, не прожигать бюджет на токенах) и попросил проработать идею внутри них: «Подумай, как ты это видишь». А затем мы вместе трансформировали идею в другой формат: пусть продукт наблюдает, где пользователи не укладываются в наши сценарии, — и мы будем расширять продукт под их реальное поведение, вместо того чтобы давать конструктор. Важная деталь: авторство осталось за сотрудником. Я прямо сказал: это твоя идея, запиши её на доску и развивай. Присвоенная идея работает; идея, которую у человека забрали, превращается в чужое поручение.

Теперь посчитайте длину цикла. Без ИИ это выглядело бы так: я зарубил идею, сотрудник замолчал, через полгода я недоумеваю, почему от него нет инициатив, — и уже не могу связать следствие с причиной. С ИИ паттерн был замечен и перестроен, пока разговор ещё шёл. Конфликт, который не разгорелся; идея, которая не умерла и стала продуктовым принципом; урок, который я получил в моменте. И правило на будущее

для каждого продуктового спора: прежде чем возражать, спроси — «ты про какой этап?». Петля обратной связи сработала прямо внутри цикла — почему это важнее, чем кажется, разберём в конце главы.

Кейс второй: развивающий диалог вместо давления

Второй случай — противоположный по темпу: не один разговор, а работа на горизонте в полтора месяца. Мой ведущий разработчик — сильный творец: тащит продукт, генерирует решения сам. И, как все люди этого склада, чужие идеи не принимает, а отбивает. По типологии Адизеса это классический Е-тип, «предприниматель»: главный двигатель изменений в команде — и одновременно самый устойчивый к чужим изменениям. Проблема была простая и болезненная: без устойчивого режима он периодически выпадал — сидел за работой круглыми сутками, выгорал, терял ритм.

Прямое «тебе нужен режим» здесь не работает. Для творца дисциплина звучит как клетка — покушение на свободу, то есть на то, чем он является. Давить бессмысленно: получишь формальное согласие и тихий саботаж. Поэтому вместо убеждения я выстроил развивающий диалог — и вёл его вместе с ИИ. Перед каждым значимым разговором и после него я разбирал переписку с ассистентом: чему именно человек сопротивляется, где я собираюсь дожать — и этого делать нельзя, какая формулировка снимет конфликт, а какая закрепит за идеей чужое авторство. Ассистент буквально держал меня за руку в местах, где я привык давить исходя из своего корпоративного опыта в 15 лет. Одно правило мы зафиксировали дословно: «если я буду давить, будет только хуже».

Механика получилась такой. Полтора месяца я ненавязчиво «сеял» идею — без требования решения, просто как мысль. Он отбивал: «пока не надо», «я сам пойму, когда». В решающем разговоре я отступил: «Хорошо». Без обиды и подтекста. Через некоторое время он вернулся сам: «возможно, этот день именно сегодня» — и спроектировал собственный ритм работы, причём лучше того, что предлагал я. Ключевой ход подсказал разбор с ИИ: не доказывать, что дисциплина

полезна творчеству, а переопределить её — выстраивание собственной системы и есть творчество: ты сам ищешь, сам оптимизируешь, и только ты решаешь, получилось ли. И выстроенная система высвобождает ресурс (время и силы) для нового творчества. Конфликт исчез в корне: принять идею перестало означать предать себя.

Итог проявился через полторы недели — и без единого моего запроса. Голосовое сообщение, в котором он описывал, как перестраивает весь свой процесс разработки: своя карта целей, конвейеры шагов, автоматические «критики». Мотив сформулировал сам: «хочу удочку, а не рыбу». Идея, которую он полтора месяца отбивал как чужую, стала его собственной — он строит под неё инфраструктуру и защищает как своё изобретение. Моя роль свелась к двум вещам: вопросу «какая моя помощь и поддержка тебе нужны?» — и роли предохранителя, чтобы энтузиазм не съел основные задачи.

Кейс третий: решение от идеи до рефлексии

Третий кейс — не про команду и вообще не про работу. В начале лета 2026 года я продал машину. Формально — обычная сделка по перераспределению финансов. А по сути — история о смене самоидентификации, которую я сам до конца не видел, пока не разобрал её с ассистентом.

Сначала вводные. Больше десяти лет я строил карьеру в корпорациях и дорос до руководителя цифровой трансформации. У этой роли есть атрибуты, и флагманский кроссовер один из них, как часть образа «состоявшегося наёмного топа».

За рулём я провёл 13 лет и 172 тысячи километров. Машина была частью меня.

При этом параллельно готовился запуск моих собственных ИИ-продуктов — и им были нужны две вещи, которых всегда не хватает: живой капитал и мой фокус.

Началось с одного сообщения ассистенту: «В голове крутится мысль. Может быть, продать машину? Будут доступные деньги (не

замороженные в железе), что очень полезно в кризис. Плюсом я разгружаю статью затрат. Общественный транспорт и такси в любом случае будут дешевле в 3 раза.

Ну и главное. Это дополнительный ресурс, который могу направить в продукт (оплата инфраструктуры или рекламы). Также запуск продуктов может потребовать переезда в другую страну.

Но внутри немного грустно. Машина — это продолжение меня, чувство контроля при вождении и «моё пространство».

Помоги принять решение».

Рациональный контур мы прошли за пару сессий: полная стоимость владения (все платежи, обслуживание, амортизация), а затем матрица из трёх сценариев будущего.

Сценарий будущего	Оставить машину	Продать
Остаюсь в найме	Комфорт — но затраты съедают поток для продуктов	Свободный поток + опциональность
Кризис или сокращение до запуска продуктов	Вынужденная продажа с дисконтом	Подушка безопасности на месяцы, гибкий бюджет для жизни и запуска продуктов
Международный переезд из-за запуска продуктов		Управляемая сделка по нормальной цене, личная гибкость и мобильность, нет переживаний из-за замороженного актива

Вывод получился жёстче, чем я ожидал: машина выигрывала только в одном сценарии из трёх — и только по комфорту, у которого есть заменители. Во всех остальных она структурно ухудшала позицию: я нёс полный набор рисков, а выгоду получал в одном будущем из трёх. От первого сообщения до твёрдого решения прошло семь-восемь дней.

На этом обычный «калькулятор» закончился бы: цифры сошлись, продавай. Но именно дальше ассистент оказался ценнее всего — он не дал закрыть тему расчётом. Тоска, которая поднялась после решения,

была разобрана на слои. Первый: горевание — не сомнение. Аргументы сошлись, а грустить по уходящей главе нормально: «прощаться без грусти можно только с тем, что ничего не значило». Второй: потребность в «контроле за рулём» оказалась не про машину — а про желание управляемости в период, когда крупная часть жизни ушла в неопределённость: продукты ещё не запущены, доход перестраивается. Руль давал мгновенную обратную связь там, где остальная жизнь её пока не давала.

И третий, самый громкий голос: «это деградация социального уровня». Разбор дал понимание, что за рациональными факторами сидит эмоциональное ядро. И если социальный уровень держится на машине — он не твой, он арендован у автомобиля. Мой реальный капитал — книги, продукты, портфель проектов, преподавание — с парковки не считывается. Больше того: в среде, куда я перехожу, — основатели, продуктовые команды, инвесторы — дорогая машина читается ровно обратным образом. Тогда и стало ясно, что спор внутри шёл не про транспорт. Машина была материальным якорем роли, которую я перерастал, — и это делало последнюю ставку перед расставанием с атрибутом.

Оговорюсь для скептиков: ассистент не работал «психотерапевтом» и не выносил вердиктов о моей личности. Он структурировал то, что я сам проговаривал, — и не давал одному слою съесть другой: расчёту — спрятать эмоции, эмоциям — отменить расчёт.

Посреди процесса вмешалась жизнь: жена попросила отложить продажу до осени — ради лета вдвоём. План уже разогнался, и была тяга «аргументированно обойти препятствие». Вместо этого мы с ассистентом разобрали сам запрос — и оказалось, что он не про машину. За ним стоял страх потерять наше время вдвоём; машина была только поводом. Мы с женой договорились о главном — о времени друг для друга, — и продажа перестала быть угрозой. Заметьте: не пришлось выбирать между сотнями тысяч и тем, что деньгами не меряется, — потому что разобрали первопричину, а не спорили о симптоме.

Сделка закрылась за пять дней — на 13% выше второго дилерского предложения и на 6% выше лучших предложений частников. Весь цикл — от первой мысли до передачи ключей — занял четырнадцать дней. А после сделки — последний укол: «не продешевил ли?» Ассистент назвал механизм: сожаление продавца, которое включается именно потому, что сделка прошла легко. Смущала не цена — смущало отсутствие страдания. А страдание не входит в стоимость хорошей сделки.

Грусть, которой я так боялся, отпустила примерно за пять дней после передачи ключей. Дальше началось то, чего я не закладывал в расчёт: уверенность. Высвобожденный ресурс уже работает в моём проекте — и это ощущается совсем иначе, чем «актив на парковке». Плюс глобальная мобильность: я больше не привязан к вещи и могу свободно перемещаться между странами. Новая идентичность закрепилась не аффирмацией, а действием. Ну и главное — больше движения дало прогресс по моим целям: сбросу веса (до этого была стагнация около девяти месяцев) и поддержанию здоровья.

И финальная сверка — гипотеза против факта. Расходы по статье «мобильность» снизились на 75% относительно затрат на машину. Высвобожденный кэш работает в проекте. А страданий, которыми пугало это, не случилось вовсе. Решение закрыто не ощущением — фактами. Кроме того, ощущение глобальной свободы оказалось сильнее, чем комфорт от возможности сесть за руль и самому доехать до точки назначения. Отдельно отмечу и то, сколько ментального ресурса освободилось от того, что больше не надо думать о других водителях, правилах и тому подобном.

Повторить такой разбор проще, чем кажется. Он начинается с трёх вопросов: в каких сценариях будущего это решение выигрывает, а в каких проигрывает? Что эта вещь значит для меня, кроме функции? Что я почувствую через месяц после? Мой ассистент такой разбор уже умеет — см. блок «Проработка решений» выше.

Теперь сведу три кейса воедино. Короткая петля обратной связи работает на трёх масштабах: разговор, который разворачивается за минуты;

развитие человека, которое созревает месяцами; собственное решение — от первой мысли до рефлексии после сделки. И во всех трёх ИИ ценен не готовыми ответами. Он ценен тем, что помогает увидеть себя со стороны, пока ситуация ещё жива: где ты давишь, где слышишь не то, где прячешь эмоцию под расчёт. Устойчивые изменения — в людях, в отношениях, в самом себе — нельзя навязать извне. Можно создать условия, в которых они происходят сами. Ассистент, который знает ваши цели и помнит вашу роль, — самый дешёвый из известных мне способов такие условия создавать.

Что даёт союз руководителя с ИИ

Прежде чем раскладывать эффекты по полочкам, назову главный парадокс этого союза. Ассистент силён тем, чего у него нет: собственных эмоций, усталости, личного интереса в вашем решении. Он не обижается на резкость, не устаёт к десятому вопросу и ничего не выгадывает. Но при этом — если вы ведёте личный контур — он внимателен к мелочам и по самим формулировкам распознаёт ваше состояние: где вы злитесь, где прячете сомнение под расчёт. Беспристрастный и внимательный одновременно. У людей эти качества почти не встречаются вместе: близкие — небеспристрастны, беспристрастные — далеки.

Экономия энергии. Проработка решений, проектов и черновиков перестала съедать вечера. Дело не только во времени: не тратятся силы на чистый лист и на удержание всего в голове. Дать контекст, указать, что вы хотите и какие есть ограничения, получить черновик и доработать — проще и быстрее, чем начинать всё с нуля.

Фокус на стратегическом. Делегируя рутину команде и ИИ, я наконец занимаюсь тем, чем должен по своей роли: стратегией, архитектурой решений, размышлением. По моему опыту, эффект виден буквально в структуре дня: появляется время держать баланс стратегического, тактического и оперативного, а не тонуть в последнем. Это не про «работать меньше» — это про работать на своём уровне и иначе.

Глубина. Решения проходят проверку до объявления, разговоры готовятся заранее, в проекты я успеваю погрузиться до первопричин, а не по верхам.

Личностный рост через короткие циклы. Мы растём, когда видим последствия своих действий и успеваем связать причину со следствием. Обычно между поступком и осознанным выводом проходят дни или недели — и связь теряется. ИИ сокращает эту петлю до минут, а именно её скорость определяет, как быстро мы учимся. Кейс выше — ровно об этом.

Оговорюсь: ИИ не решает за вас и не заменяет живую рефлексия — он её ускоряет. Причинно-следственные связи по-прежнему выстраиваете вы, просто теперь в разы быстрее. Это, пожалуй, самый недооценённый способ, которым ИИ усиливает человека: не вместо него, а вместе с ним.

Граница союза: непохожесть остаётся вам

И про главную границу этого союза. ИИ отлично справляется с конвергентными задачами: анализ, структурирование, поиск паттернов, выводы, исполнение — всё, где существует «правильнее» и «точнее». Со «средней» креативностью он тоже справляется: наполнит брейншторм, предложит десять вариантов заголовка, докрутит черновик. Исследования показывают: идеи от ИИ делают тексты отдельного человека заметно лучше — особенно если генерация идей не его конёк. Но есть подвох, и он стратегический. Модель обучена на гигантском массиве чужого и статистически тянет к среднему: те же исследования фиксируют, что работы, созданные с ИИ, становятся ощутимо более похожими друг на друга. ИИ интерполирует изученное — подлинно новое, то, чего не было в данных, он не предложит. Ещё пример: он отлично помогает с тактическими решениями. Но стратегические — и те, где нужно проявить настойчивость, гибкость, найти выход из неоднозначной ситуации, — лучше оставлять за собой.

Живой пример — книга, которую вы держите в руках. Третью редакцию я готовил в паре с ИИ-ассистентом, и разделение труда сложилось ровно по этой линии. Актуализация фактов и кейсов, сверка цифр с

первоисточниками, структурирование глав, поиск дублей и противоречий, вычитка, упаковка идей в текст — всё это ассистент делает блестяще: быстрее и тщательнее меня. Но каждая новая идея — добавить эту главу, вставить историю с продажей машины, оживить сухой раздел личным кейсом — приходила от меня. Ассистент доводил идею до текста за минуты. А вот решить, чего книге не хватает, он не мог ни разу.

С одной оговоркой, которую эта же книга мне и выдала. Историю с картой понятий вы уже знаете из Главы 6: черновик делал ассистент, потом схема прошла машинную проверку, обе стадии дали зелёный свет — а ошибка композиции в ней была, и увидел её я только на вычитке.

Дело не в том, что ассистент плох. Просто «тщательнее меня» перестаёт работать там, где нужно заметить ошибку в собственной работе: машинная проверка её пропустила, поймал я. И поймал не правкой, а вопросом.

Отсюда правило, которым я закончу разговор о личном контуре: рутину и «среднее» отдавайте ИИ — непохожесть оставляйте себе. Чем больше людей черпают идеи из одного источника, тем сильнее их результаты сливаются — и тем дороже становится ваша способность придумать то, чего нет в среднем по рынку.

Как это превратилось в продукт

Честная оговорка: дальше я рассказываю о собственном продукте. Читайте этот раздел не как рекламу, а как разбор продуктовой логики — она пригодится, даже если вы строите или выбираете совсем другое решение.

Личный контур, который я описал выше, изначально жил у меня на подручных инструментах: заметки, чаты с ИИ, таблицы, шаблоны. Плюс я как консультант и методолог формировал управленческую методологию. В какой-то момент стало ясно, что из этих практик собирается продукт — персональная операционная система

руководителя. В России он выходит под именем «Сокол», на международном рынке — BAEOS.

Итак, от чего я отталкивался. Целевой пользователь — директор или заместитель в компании на 50–300 человек. «Играющий тренер»: уже не может управлять вручную, но ещё не может выйти из операционки. Его день состоит из приоритетов, решений, встреч, согласований и постоянного выбора между срочным и важным — и заметная часть этого дня уходит на сбор контекста и «пожары». При этом ему некогда изучать методологии и так далее. У него бизнес, который живёт каждый день.

Продуктовая логика держится на нескольких принципах, и каждый вырос из практики, а не из презентации.

Первый: сначала ценность здесь и сейчас. Пользователь должен получить пользу в первый же день — спланировать день, зафиксировать решение, подготовить письмо, — а не «после обучения модели когда-нибудь потом».

Второй: источник истины — действия, а не самоописание. Моя цель — построить ИИ, который будет знать руководителя и быть его автопилотом. В управленческом контексте опросники дают желаемый образ; реальное поведение видно только в реальной работе — как человек планирует, решает, делегирует, реагирует на сбои. Из этого поведения постепенно складывается профиль руководителя, и рекомендации становятся точнее (как при подготовке писем, так и при подборе рекомендаций по личному развитию).

Третий принцип: никаких абстрактных чатов. Человеку нужны сценарии — гибкие, адаптивные, но сценарии на основе методологий, чтобы система вела его сама. А расширяются сценарии тем способом, который родился в первом кейсе этой главы: наблюдаем, где пользователь не укладывается в заложенное, — и достраиваем под реальное поведение, вместо того чтобы давать конструктор.

Внутри — те же модули, что и в моём личном контуре: день и дневник с рефлексией, планирование от дня к месяцу, решения, коммуникации,

делегирующее и команда, цели с динамикой достижения, проекты с портфелем, обзор дня одним экраном. Коммуникации при этом — не абстракция: отдельные сценарии для встреч, переговоров, обратной связи и писем. А документы не живут отдельной «папкой» — нужный контекст подгружается прямо в сценарий: к проработке решения, коммуникации или в проект.

Мета-кейс: статья за два дня

Свежий пример работы контура — мои собственные публикации. В августе 2026-го я выпустил статью «Русская смекалка против эпохи ИИ». Разделение труда было таким. За мной — идея, тезисы, критерии качества и вычитка. За ИИ — проверка фактов с реестром из 79 источников, сборка текста по моим правилам стиля, SEO-упаковка, инфографики и публикация на сайте. Календарных дней — два. Раньше такой цикл занимал недели и требовал трёх подрядчиков: редактора, дизайнера и факт-чекера.

Отсюда наблюдение, которое меняет оргдизайн. **ИИ — рычаг таланта: один сильный человек закрывает пул задач целого подразделения.** Это не значит «увольте отдел». Это значит, что узким местом становится не численность, а количество людей, способных поставить задачу и верифицировать результат. Ищите и растите таких людей — рычаг уже лежит на столе.

И последнее. Даже если вы никогда не откроете наш продукт, главную мысль главы можно забрать бесплатно: личный контур руководителя собирается из любых инструментов, хоть из блокнота и одного чата с ИИ. Важен не софт, а привычка: короткая запись, мгновенная обратная связь, короткая петля. Всё остальное — вопрос удобства.

Резюме главы

- ИИ применим во всех классах задач руководителя — от планирования до саморазвития, — но нигде не заменяет его целиком: это правило 70/30.

- Самый недооценённый эффект — сокращение петли «действие → обратная связь» с недель до минут: руководитель быстрее учится, конфликты разбираются, пока живы, решения проверяются до объявления.
- В коммуникациях ИИ способен работать тренером в реальном времени: отслеживать ваши паттерны — давление, подавление инициативы, отъём авторства — и помогать перестроить разговор, пока он ещё идёт.
- В решениях ИИ ценен не готовым ответом, а тем, что держит на столе оба слоя — расчёт и эмоции — и не даёт спрятать один под другой. Решение остаётся вашим, но проходит полный цикл: от первой мысли до рефлексии после сделки.
- Граница союза: рутину, тактику и «среднее» отдавайте ИИ — непохожесть оставляйте себе. Чем больше людей черпают идеи из одного источника, тем дороже способность придумать то, чего нет в среднем по рынку.

Что с этим делать: Заведите дневник руководителя: 5 минут вечером, запись плюс обратная связь от ИИ (неделя). Соберите личный контур: планирование, проработка решений, подготовка к сложным разговорам (квартал).

Вопрос для рефлексии: Сколько времени в вашей неделе занимает работа, которую можете делать только вы, — и сколько та, которую уже готов взять ИИ?

Глава 22. Как вести проекты по внедрению технологий

Для наглядной демонстрации приведу пример того, как я иногда ранжирую проекты и формирую список задач (в рамках организационной структуры общества с ограниченной ответственностью).

Проекты в компании могут делиться на четыре типа:

- локальные проекты подразделений;

- кросс-функциональные проекты;
- стратегические проекты;
- инвестиционные проекты.

Локальные проекты подразделений

Локальные проекты подразделений должны соответствовать следующим критериям:

- выполняются силами одного структурного подразделения или одной службой заместителя директора (далее — ЗД);
- бюджет проекта на привлечение внешних поставщиков услуг и закупку материалов составляет не более 500 тысяч рублей (без НДС);
- проект не вносит изменения в бизнес-процессы других подразделений и не требует доработки ИТ-систем Общества, в том числе не требуется изменение прав доступа к ним.

Кросс-функциональные проекты

Кросс-функциональные проекты должны соответствовать следующим критериям:

- бюджет проекта на привлечение внешних поставщиков услуг и закупку материалов — не более 10 млн рублей (без НДС);
- в реализации проекта участвует не более двух подразделений или служб ЗД Общества, также участвует не более двух организаций Группы компаний;
- изменения происходят не более чем в трёх бизнес-процессах, и они затрагивают не более трёх подразделений организации.

Стратегические и инвестиционные проекты

К стратегическим относятся проекты, которые соответствуют одному или нескольким критериям ниже:

- бюджет проекта — свыше 10 млн и до 100 млн рублей включительно;
- проект влияет на достижение стратегических целей компании;

- происходят изменения более чем в трёх бизнес-процессах первого или второго уровня, а изменения затрагивают более трёх подразделений, курируемых разными ЗД.

Проекты с бюджетом более 100 млн рублей относятся к инвестиционным. Инвестиционные проекты реализуются аналогично стратегическим. При этом идёт дополнительная проработка проекта в соответствии с методикой оценки инвестиций и стандартами Группы компаний.

Жизненный цикл проекта

Проекты в организации проходят следующие стадии жизненного цикла:

- инициация — стадия, по результатам которой принимается решение о запуске проекта или его отклонении, формулируются цели и ожидания от проекта;
- планирование — стадия, в которой происходит подготовка плана реализации проекта;
- реализация и контроль — стадия, в которой выполняются работы по реализации проекта и контроль хода реализации, качества работ, ведётся управление коммуникацией, командой, ожиданиями заинтересованных сторон;
- закрытие — стадия, в которой идёт подготовка отчётных и бухгалтерских документов, архивирование проекта и обучение по итогам реализации проекта.

В ходе инициации и планирования проекта **рекомендуется** оценивать необходимые ресурсы для реализации проекта по следующим категориям:

- производственные (инструменты, оборудование);
- финансовые (собственные, заёмные);
- материальные (сырьё, комплектующие и расходные материалы);
- человеческие (люди и компетенции, контакты сотрудников внутри группы компаний);

- интеллектуальные (технологии, патенты);
- организационные (управление и административный ресурс).

Определение целей проекта и возможных эффектов

Целеполагание — ключевой шаг. Что мы хотим от проекта?

Для определения целей, задач и потенциальных эффектов необходимо определить вид проекта и соотнести его с возможными целями.

Виды проекта и возможные цели

Вид проекта	Возможная цель
Коммерческие/экономические	Получение прибыли
Технические	Решение технических задач, не связанных с экономическими показателями
Имиджевые / маркетинговые / социальные	Формирование необходимого имиджа или репутации, продвижение компании или идеи / повестки
Исследовательские / обучающие	Поиск новых решений или обучение новым компетенциям
Операционные / организационные	Повышение внутренней эффективности и результативности, в том числе снижение трудозатрат, ускорение процессов
Стратегические	Решение стратегических задач организации
Смешанные	Сочетание вышеперечисленных

Рекомендуется выделить основную цель, которая определяет общее направление работы. А также дополнительные 3–4 подцели (ключевые задачи), выполнение которых будет характеризовать достижение общей цели.

Для оценки эффектов проекта рекомендуется определить, какие потери в деятельности организации устраняются (и обязательно для

операционных проектов). Также полезно использовать эти потери для оценки рисков в проекте.

Также данные потери возможны и в ходе реализации проекта. Они должны оцениваться в ходе определения и оценки рисков проекта как организационные.

Роли в проекте

Также важно указать, какие могут быть роли: это поможет упростить новаторам распределение ответственности в командах. Сотрудники в ходе реализации проектов могут участвовать в ролях согласно описанию ниже.

Роли в реализации проектов

Роль	Описание роли и выполняемые функции	Дополнения и пояснения
Заказчик проекта (Спонсор)	Определение цели, результатов проекта; определение источника финансирования проекта; корректировка проекта (при необходимости); принятие решения о конечных результатах проекта	Допускается сочетание с ролью Куратор
Куратор	Общее руководство ходом реализации проекта; обеспечение финансирования работ и выделение необходимых ресурсов для выполнения проекта; рассмотрение наиболее острых проблем проекта, затрагивающих взаимодействие заказчика и исполнителей; участие в	Допускается сочетание с ролью Заказчик проекта (Спонсор) или Руководитель проекта

	управлении рисками проекта	
Руководитель проекта	Формирование команды проекта, организация взаимодействия между участниками проектной команды и заказчиком; планирование, организация и контроль выполнения работ по достижению целей проекта с требуемыми затратами, качеством и в заданный срок; делегирование задач и полномочий, расстановка приоритетов, обеспечение ресурсами работников команды проекта; установка, контроль показателей результативности и требований к выполнению работ; управление рисками проекта и процессом решения проблем; управление изменениями в период реализации проекта; участие в разрешении противоречий в проектных решениях	Допускается сочетание с ролью Куратор проекта
Администратор проекта	Обеспечение руководителя проекта структурированной информацией, необходимой для контроля проекта,	Допускается сочетание с ролью Руководитель проекта (особенно если используете ИИ — эта функция хорошо автоматизируется)

	планами, ресурсами и приоритетами; обеспечение своевременной подготовки, движения и архивации документов по проекту; контроль согласования документов проекта; организационная и коммуникативная работа в команде проекта; отчётность по проекту	
Участник команды проекта	Реализация отдельных задач проекта	
Консультанты (по направлению)	Консультирование проектной команды по направлениям деятельности	
Пользователь продукта проекта	Формулирование требований к продукту проекта на стадии инициации и планирования; оценка продукта проекта и подготовка обратной связи в ходе реализации	Допускается сочетание с ролью Заказчик проекта (Спонсор)

Это классические проектные роли — они одинаковы для любого проекта. Специфические роли ИИ-функции — владелец ИИ-продукта, инженер данных, ML-инженер, управляющий комитет — разобраны в книге «ИИ-2. Практическое руководство» (Глава 6).

В каждой задаче у участника возможна одна из следующих зон ответственности:

- И (исполнитель) — тот, кто будет физически выполнять работы;

- О (ответственный) — тот, кто примет итоговую работу и будет нести ответственность (может сочетаться с ролью И);
- К (консультант) — тот, кто будет в обязательном порядке консультировать остальных при выполнении задачи;
- У (уведомляемый) — тот, кто должен быть в курсе принимаемых решений или хода выполнения задач, но влиять на них никак не будет.

Три алгоритма — кратко

Три типа ИИ-проектов

ЛОКАЛЬНЫЙ

Одно подразделение · недели · до 0,5 млн
Ломается на отсутствии владельца и метрик

КРОСС-ФУНКЦИОНАЛЬНЫЙ

Стыки подразделений · месяцы
Ломается на конфликте приоритетов

СТРАТЕГИЧЕСКИЙ

Портфель · годы · спонсор — первое лицо
Ломается при потере спонсорства

Локальные проекты подразделений. Самый массовый тип: инициатива и бюджет одного подразделения, короткий цикл. Ключевое: зафиксировать цель и базовую линию до старта, быстро пройти пилот на готовых инструментах, честно замерить эффект и принять решение о масштабировании. Ломаются такие проекты на отсутствии владельца и метрик.

Кросс-функциональные проекты. Технология здесь редко главный риск — риски живут на стыках подразделений. Нужны единый заказчик, управляющий орган, согласованные данные и регламенты, поэтапное

внедрение с контрольными точками. Ломаются на конфликте приоритетов и «ничьих» зонах ответственности.

Стратегические проекты. Портфельная логика: спонсор уровня первого лица, выделенная команда, управление изменениями и рисками, связь с целями компании и обязательные промежуточные результаты каждые несколько месяцев. Ломаются на потере спонсорства и длинных горизонтах без быстрых побед.

Каждый из трёх алгоритмов проходит четыре стадии: инициация, планирование, реализация, завершение. Пошаговые версии со шаблонами, ролями и контрольными точками — в книге «ИИ-2. Практическое руководство» (Глава 11). Для этой книги важна управленческая рамка: тип проекта определяет глубину методологии, а «единый файл проекта» с ИИ-ассистентом держит всё в одном контуре.

Резюме главы

- Типы проектов — локальные, кросс-функциональные и стратегические (инвестиционные ведутся по стратегическому алгоритму) — требуют разной глубины методологии; ошибка масштаба (вести локальный проект как стратегический и наоборот) убивает и те и другие.
- Управленческое ядро любого ИИ-проекта: цели и эффекты определены до старта, роли распределены, «единый файл проекта» собирает статусы, риски и решения.
- ИИ-ассистент в проектном контуре забирает рутину ведения: протоколы, статусы, реестры рисков — время руководителя возвращается на решения.

Что с этим делать: Классифицируйте ваш ближайший ИИ-проект по типу и соберите роли (неделя). Внедрите «единый файл проекта» с ИИ-резюме статусов (квартал).

Вопрос для рефлексии: Кто в вашем проекте отвечает за бизнес-эффект после внедрения — а не только за факт запуска?

Глава 23. Системный подход к внедрению и дорожная карта

О волне хайпа на искусственном интеллекте

К 2026 году волна хайпа схлынула: рынок прошёл пик завышенных ожиданий и перешёл к трезвой фазе — побеждает системность, а не энтузиазм (подробнее — в Постскриптуме). Тем ценнее урок, с которого начинался этот раздел в первой редакции.

Когда пошла волна хайпа с ChatGPT, все думали, что ИИ разгрузит людей от рутинного труда. Но по факту пока он заменяет людей в творческом и интеллектуальном труде. И здесь люди вполне иронично подметили, что мы хотели заниматься творчеством, а роботов отправить в шахты, но на данный момент роботы генерируют творчество, а люди по-прежнему трудятся в шахтах. Этого ли мы хотим от ИИ? В качестве визуализации хочу привести интересный пост с мнением человека.

Собственно, ниже сам пост и оригинальная орфография.

«1. „Мне нужен ИИ, который будет вместо меня мыть посуду и стирать бельё. Мне не нужен ИИ, который будет вместо меня писать и рисовать, чтобы в освободившееся время я занималась мытьём посуды и стиркой белья“. Наткнулся в Твиттере на такой очень точный комментарий одной креативщицы.

«2. Эта простая фраза даёт хорошее направление поиска идей для ИИ-стартапов. Не пытайся создать ИИ, который вместо человека будет делать что-то для него важное! ИИ должен брать на себя неважные дела — рутинные и побочные вещи, которыми человек заниматься не любит, но вынужден.

«3. Соответственно, задача опросов потенциальных потребителей — разделить всё, что они делают, на две группы: что они любят делать, и что не любят. То, что любят — не трогаем. А то, что не любят — заменяем своими ИИ-машинками и им же продаём».

Здесь мы видим как раз иллюстрацию продуктового инструмента — Jobs To Be Done. Есть два типа задач — одни хотят просто решить, от других получают удовольствие в процессе их решения.

Именно поэтому я раз за разом говорю: в ИТ только 20% технологий, а всё остальное — коммуникация, анализ и принятие решений. ИИ — всего лишь один из инструментов цифровизации и автоматизации, который способен на многое, но он не панацея от всех проблем, и для успеха нужен системный подход.

Системный подход к внедрению ИИ

Как мы уже с вами поняли, внедрять ИИ надо, как и проводить цифровизацию в целом, системно и комплексно.

Мой системный подход опирается на восемь направлений — они на схеме ниже. Каждое из них подробно разобрано в книге «ЦТ-2. Системный подход»; здесь остановлюсь на том, что определяет последовательность внедрения ИИ, — на дорожной карте и связях с другими проектами.



Восемь направлений системного подхода

Дорожная карта и связи с другими проектами

ИИ надо внедрять в рамках общей дорожной карты и стратегии цифровизации, цифровой трансформации и развития. ИИ раскроет свой потенциал только в сочетании с другими технологиями и инструментами. Повторюсь, это не волшебная палочка. ИИ зависит от

общей стратегии компании, работы с данными и стратегии цифровизации (в том числе цифровой зрелости).

Также надо понимать, от чего будут зависеть ИИ-проекты — от наличия качественных данных в компании. Это в свою очередь зависит от выстраивания процессов обеспечения качества и автоматизации сбора данных (исключение ручного ввода). А автоматизация сбора данных зависит от наличия качественной связи и внедрения цифровых инструментов.

Пожалуй, исключением здесь являются проекты машинного зрения (они сами могут стать источниками данных) и цифровые советники по регламентам и законам или по направлениям деятельности на основе открытых методологий (управление проектами, бережливое производство, налоговые консультанты и т.д.).

И как бы я ни пытался создать отдельную дорожную карту, она полностью будет дублировать то, что уже есть в книге «ЦТ-2. Системный подход» (также доступно по QR-кодам и гиперссылкам выше).

Таким образом, проекты машинного зрения и цифровые советники для офиса должны стать первостепенным направлением. Это поможет получить быстрые результаты. А для долгосрочного развития и внедрения внутренних советников, которые будут приносить наибольшую ценность, обязательно надо заниматься автоматизацией сбора данных и выстраиванием процессов обеспечения качества данных в организации, организационным развитием — назначением владельцев данных, обучением людей и регламентами, без которых процессы качества не живут.

Резюме главы

- Волна хайпа схлынула — и это хорошо: системный подход, а не энтузиазм, определяет теперь результат внедрения.
- Дорожная карта ИИ не живёт отдельно: она встраивается в общую стратегию цифровизации и связана с проектами данных, инфраструктуры и обучения людей.

- Последовательность важнее скорости: связь и цифровые инструменты → автоматизация сбора данных → их качество → ИИ-проекты. Исключение — машинное зрение и советники по открытым методологиям: они не ждут данных и дают быстрый результат. Перепрыгивание этапов возвращает проект на старт — уже с потраченным бюджетом и подорванным доверием заказчика.

Что с этим делать: Разложите свои ИИ-инициативы на зависимые и независимые от корпоративных данных и запустите вторые (неделя). Сведите инициативы в единую карту со связями с ИТ-проектами (квартал).

Вопрос для рефлексии: Есть ли у вашей компании дорожная карта ИИ — или только список пилотов?

Глава 24. Что должен знать и уметь руководитель в эпоху ИИ

Введение

Внедрение ИИ в компаниях невозможно без соответствующих компетенций у персонала. Так, например, до 95% компаний отмечают, что внедрение ИИ не привело к ощутимому финансовому результату. И одна из причин — недостаток компетенций у персонала. Это отражается в целом на направлении цифровизации, так как централизованный подход, когда за всё отвечает ИТ-служба, не работает, и нужно, чтобы идеи рождались у бизнеса и адаптировались под специфику. Но чему конкретно учить людей? И что должны знать руководители, чтобы контролировать развитие и внедрение искусственного интеллекта?

Когда я работал над материалом ниже, то посчитал долгом сохранить ту структуру, которая была представлена в книге «ЦТ-2. Системный подход». Логика у меня простая: ИИ — одна из цифровых технологий, а значит, и подчиняется она правилам цифровизации в целом. Единственное, в этой книге я сделал чуть большую детализацию именно в контексте ИИ.

Модель компетенций

Я предлагаю пойти по четырёхуровневой шкале оценки персонала:

- Уровень 0 — знаний и навыков нет; использование не требуется по роли.
- Уровень 1 — базовый: знает основы, работает по инструкциям, нуждается в поддержке.
- Уровень 2 — рабочий: уверенно применяет, адаптирует под задачи, помогает коллегам.
- Уровень 3 — экспертный: задаёт стандарты, создаёт решения, обучает и консультирует.

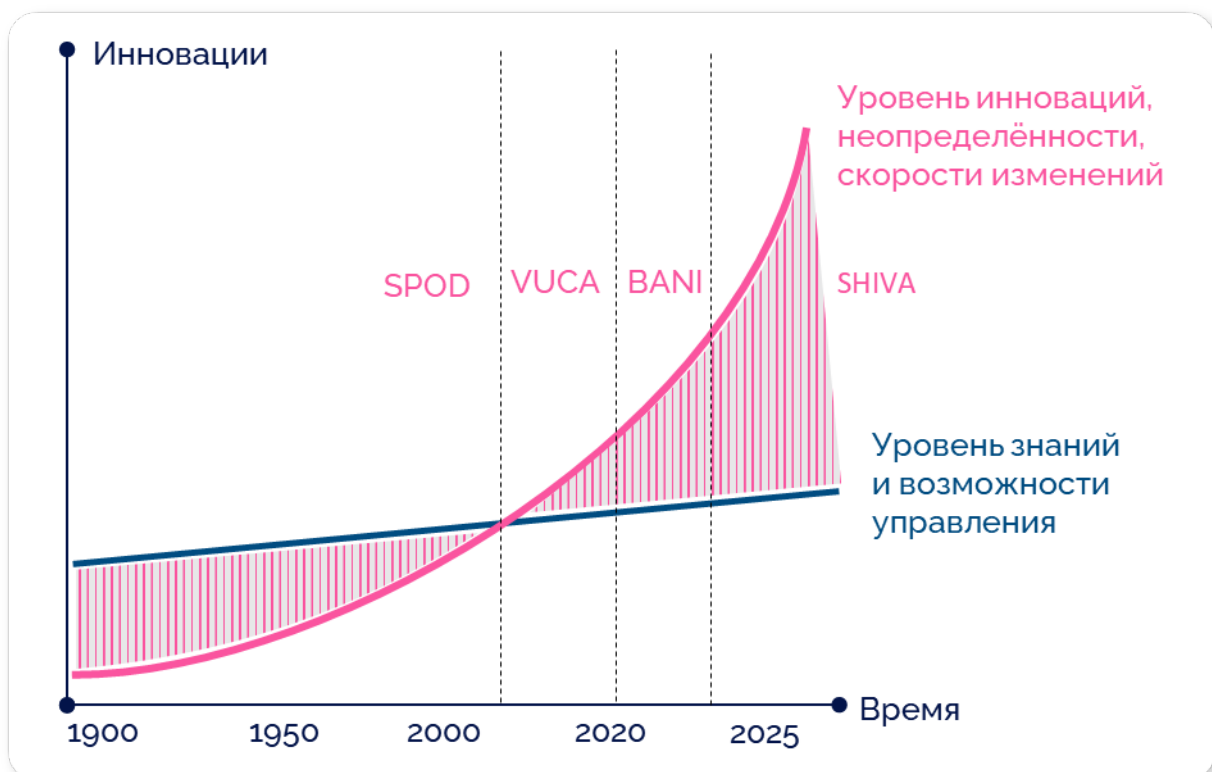
Сама же модель строится на трёх направлениях — личностные, управленческие и цифровые компетенции.

1. Личностные компетенции

В данном направлении я выделяю девять ключевых компетенций.

1.1 Готовность изучать новое / Адаптивность к изменениям

На данный момент мы живём в мире, где изменения происходят всё быстрее.



На мой взгляд, именно адаптивность и гибкость становятся всё важнее. В контексте ИИ можно выделить несколько блоков:

- быстрое освоение новых ИИ-инструментов и технологий;
- гибкость в изменении подходов при появлении новых решений;
- открытость к экспериментам и тестированию ИИ-возможностей;
- готовность пересматривать существующие процессы.

1.2 Нацеленность на результат и критичность

Внедрение технологий ради технологий ведёт гарантированно к убыткам. Да, обеспечить стопроцентную безопасность и совсем не рисковать не получится, но несколько направлений ниже позволят хотя бы снизить риски:

- фокус на измеримых результатах от внедрения ИИ;
- критический анализ предложений ИИ-систем;
- проверка фактов и достоверности ИИ-генерированного контента;
- оценка эффективности ИИ-решений по конкретным метрикам.

1.3 Клиентоцентричность

Любая технология в конечном счёте должна быть направлена на клиента. Причём сейчас развивается и внутренняя клиентоориентированность. Например, ИТ-блок воспринимает своих коллег именно как клиентов и собирает «клиентские» метрики.

- Применение ИИ для улучшения клиентского опыта.
- Персонализация взаимодействия с клиентами с помощью ИИ.
- Анализ потребностей клиентов через ИИ-аналитику.
- Этичное использование данных клиентов в ИИ-системах.

1.4 Эмоциональный интеллект

90% работы в ИТ-проектах — устранение страхов людей. Для ИИ этот показатель ещё выше. Он для большинства людей сравним с магией, а отсюда и страха будет ещё больше. Что здесь важно:

- понимание эмоциональных реакций на внедрение ИИ;
- управление страхами сотрудников относительно замещения ИИ;
- эмпатия при обучении коллег работе с ИИ-инструментами;
- мотивация команды на освоение ИИ-технологий.

Эмоциональный интеллект традиционно был ключевым навыком руководителя. И это логично, ведь управление всегда строилось через работу с людьми. Но по мере того, как часть задач делегируется ИИ-агентам, акценты смещаются. Для эффективного управления «цифровыми помощниками» нужна не эмпатия, а логика: умение чётко формулировать задачи, декомпозировать процессы, задавать правила и критерии оценки результата. По сути, это тот же навык постановки задач, но доведённый до уровня алгоритмической точности — потому что ИИ выполняет ровно то, что поставлено: незаданное правило для него не существует, а недосказанность он заполнит правдоподобной догадкой.

Это не значит, что эмоциональный интеллект теряет значение — люди никуда не исчезают из организаций, и работа с ними по-прежнему требует чуткости. Но к модели компетенций руководителя добавляется новое измерение: способность взаимодействовать не только с живыми людьми, но и с цифровыми системами, а для этого нужна цифровая эрудиция, умение строить процесс и работать с данными. И эти навыки стоит развивать уже сейчас.

1.5 Публичные выступления / Ораторство

Я лично неоднократно наблюдал проекты, которые так и не запустились только из-за того, что руководитель не мог защитить свою идею или боялся выступить на цифровом комитете. А в некоторых крупных компаниях чуть ли не в регламентах закреплено, что ты должен «продать» свой проект всем руководителям, объяснить его суть. В общем и целом, тут я предлагаю следующие инструменты:

- презентация результатов ИИ-проектов руководству;
- объяснение возможностей ИИ различным аудиториям;
- проведение обучающих сессий по ИИ-инструментам;
- выступления на конференциях и форумах по ИИ.

1.6 Креативность

Лучшие идеи проектов приходят от самого бизнеса. Волшебника, который всё продумает за вас, не будет. Поэтому важно использовать ИИ для раскрытия собственного потенциала и искать точки его применения повсюду. Здесь тоже можно выделить четыре направления:

- поиск нестандартных применений ИИ в рабочих процессах;
- творческий подход к формулированию запросов для ИИ;
- генерация инновационных идей для ИИ-проектов;
- создание оригинального контента с помощью ИИ-инструментов.

1.7 Решение проблем

Внедрение инноваций невозможно без проблем и сбоев. Важно в этот момент сохранять спокойствие и решать проблемы. А также думать, как возникающие проблемы в бизнесе можно решать или предотвращать с помощью ИИ.

- Диагностика проблем, решаемых при помощи ИИ.
- Структурированный подход к внедрению ИИ-решений и возникающим при этом проблемам.
- Поиск альтернативных вариантов при неудачах ИИ-проектов.
- Комплексное решение бизнес-задач с использованием ИИ.

1.8 Принятие / Терпимость к другим взглядам

Ключ к поиску решений и качественному продукту — умение договариваться. А для этого необходимо выстраивать диалог между людьми с разными характерами, что само по себе очень сложно. Кроме того, в любом случае будут успешные и неуспешные проекты.

- Открытость к различным подходам внедрения ИИ.
- Учёт мнений скептиков и критиков ИИ-технологий.
- Толерантность к ошибкам и неудачам в ИИ-экспериментах.
- Готовность к диалогу по этическим вопросам ИИ.

1.9 Методичность / Дисциплина / Настойчивость

Никакой, даже самый лучший, инструмент не приживётся, если руководитель не будет своим примером демонстрировать его важность.

- Систематический подход к изучению ИИ-технологий.
- Соблюдение регламентов и процедур при работе с ИИ.
- Документирование опыта и результатов ИИ-проектов.
- Планомерное развитие ИИ-компетенций в организации.

2. Управленческие компетенции

В данном направлении я выделяю восемь ключевых управленческих навыков, которые легли в основу моего системного подхода. Конечно же, с учётом проникновения и развития ИИ.

2.1 Стратегия, организационная структура, бизнес-процессы

- Разработка стратегии внедрения ИИ в бизнес-процессы и долгосрочного развития.
- Создание организационной структуры для управления ИИ-проектами, а также использование ИИ для её проектирования.
- Оптимизация бизнес-процессов с использованием ИИ-технологий.
- Долгосрочное планирование развития ИИ в организации (в том числе включение внедрения ИИ и экспериментов с ним в метрики руководителей).

2.2 Практики регулярного менеджмента

- Внедрение ИИ в ежедневные управленческие процессы.

- Использование ИИ-аналитики для принятия управленческих решений.
- Автоматизация рутинных управленческих задач с помощью ИИ.
- Мониторинг показателей эффективности с использованием ИИ-систем.

2.3 Цифровые технологии, работа с данными и кибербезопасность

- Интеграция ИИ с существующими цифровыми системами и технологиями.
- Обеспечение качества данных для ИИ-алгоритмов.
- Защита ИИ-систем от кибератак и утечек данных.
- Соблюдение требований информационной безопасности при работе с ИИ.

2.4 Проектное и продуктивное управление

- Управление ИИ-проектами от идеи до внедрения, в том числе расчёт ROI и TCO, понимание жизненного цикла ИИ-проектов и их отличий от классической автоматизации.
- Разработка ИИ-продуктов и сервисов для клиентов.
- Координация работы междисциплинарных команд по ИИ.
- Контроль бюджетов и сроков реализации ИИ-инициатив.

2.5 Внутренняя и внешняя коммуникация

- Коммуникация / продвижение результатов ИИ-проектов внутри организации.
- Взаимодействие с внешними поставщиками ИИ-решений.
- Управление репутацией организации в контексте использования ИИ.
- Формирование корпоративной культуры принятия ИИ-технологий.

2.6 Теория ограничений систем

- Выявление узких мест в процессах, которые может решить ИИ.

- Оптимизация производительности системы через ИИ-решения.
- Балансировка нагрузки между человеческими и ИИ-ресурсами.
- Управление ограничениями при масштабировании ИИ-решений.

2.7 Управление изменениями: мотивация, связь с общественностью, работа с сопротивлением и лидерство

- Мотивация сотрудников к освоению ИИ-технологий.
- Работа с сопротивлением изменениям при внедрении ИИ.
- Лидерство в ИИ-трансформации организации.
- Формирование положительного образа ИИ-инициатив.

2.8 Бережливое производство

- Применение принципов бережливого производства к ИИ-проектам.
- Устранение потерь в процессах с помощью ИИ-оптимизации.
- Непрерывное совершенствование ИИ-решений.
- Снижение затрат через автоматизацию с использованием ИИ.

3. Цифровые компетенции

В книге о системном подходе к цифровизации я выдвинул модель цифровых компетенций: цифровая эрудиция, работа с данными, информационная безопасность, организация совместной работы и производство контента. Я несколько месяцев пытался сформулировать модель компетенций для ИИ. Однако в конце концов я понял, что это частный случай и должен подчиняться цифровизации в целом. Поэтому данная модель актуальна и тут, вопрос лишь в детализации.

3.1 Цифровая эрудиция

- Классификация видов ИИ и области их применения.
- Возможности и ограничения ИИ-технологий.
- Формулирование запросов (промтинг) и работа с ИИ-инструментами (система постановки задач — Глава 6).

- Связь ИИ с другими цифровыми технологиями.
- Использование ИИ в корпоративных ИТ-системах.
- **Язык предметной области.**

В этот же пункт добавляется то, что раньше не выделяли отдельно: владение языком своей предметной области — оно работает в паре с промптингом. ИИ чувствителен к терминам — точное профессиональное слово даёт результат выше, а расход токенов ниже (механику мы разобрали в Главе 6). Отсюда неожиданный ответ на модный вопрос «зачем нам младшие специалисты, если есть ИИ»: не владея языком предмета, человек не поставит задачу и не проверит ответ. Экспертиза перестаёт быть про «умею делать руками» и становится про «умею назвать, поставить и проверить».

3.2 Работа с данными для ИИ

- Качество данных для обучения ИИ-моделей.
- Разметка и подготовка данных для ИИ.
- Управление данными в контексте ИИ-проектов.
- Интерпретация результатов ИИ-аналитики.
- Предвзятость и справедливость в отношении к данным для ИИ.

3.3 Безопасность ИИ

- Регуляторные требования к ИИ.
- Информационная безопасность ИИ-систем.
- Этика использования ИИ.
- Безопасность ИИ и управление рисками для бизнеса.
- Технологии сохранения конфиденциальности в ИИ.

3.4 Создание контента с ИИ

- Генерация различных видов контента с помощью ИИ.
- Инженерия запросов для создания контента в различных видах.

- Сотрудничество человека и ИИ в творческих процессах.
- Автоматизация производства контента.
- Контроль качества ИИ-генерированного контента.

3.5 Организация взаимодействия с ИИ

- Интеграция ИИ-помощников в командную работу.
- Автоматизация рабочих процессов с участием ИИ.
- Команды «человек — ИИ» и распределение задач.
- Обучение сотрудников работе с ИИ.
- Мониторинг эффективности сотрудничества человека и ИИ.

Матрица компетенций

Важная граница: здесь — компетенции руководителя и организации в целом. Компетенции и комплектование команды внедрения — роли, найм, обучение, работа с сопротивлением — разобраны в книге «ИИ-2. Практическое руководство» (Глава 13).

Сотрудников я, как правило, разделяю на пять групп:

- Собственники, CEO и заместители (CEO-1).
- Средний менеджмент.
- Неформальные лидеры.
- ИИ-специалисты.
- Обычные сотрудники (работники).

Матрица компетенций по группам:

Компетенция	CEO/ТОР	Средний менеджмент	Неформальные лидеры	ИИ-специалисты	Обычные сотрудники
1. Личностные компетенции (9 компетенций)	3	3	3	3	2

2.1 Стратегия, структура, процессы	3	3	2	2	1
2.2 Практики регулярного менеджмента	2	3	2	2	1
2.3 Цифровые технологии, данные и кибербезопасность	2	3	2	2	1
2.4 Проектное и продуктовое управление	2	3	2	3	1
2.5 Внутренняя и внешняя коммуникация	3	3	2	2	1
2.6 Теория ограничений систем	2	2	1	2	0
2.7 Управление изменениями и лидерство	2	3	2	2	1
2.8 Бережливое производство	1	2	1	2	0
3.1 Цифровая эрудиция (ИИ-фокус)	2	2	2	3	1
3.2 Работа с данными для ИИ	3	2	2	3	1
3.3 Безопасность ИИ	3	2	2	3	1
3.4 Создание контента с ИИ	2	2	3	3	2

3.5 Организация взаимодействия с ИИ	2	3	2	2	2
-------------------------------------	---	---	---	---	---

Если всё упростить, то можно сделать следующую «выжимку».

- CEO/ТОР.

Экспертные компетенции: личностные, стратегия/структура/процессы, коммуникация, работа с данными для ИИ, безопасность ИИ.

Фокус: стратегическое лидерство, управление, безопасность ИИ.

Роль: визионеры ИИ-трансформации и главные ответственные за риски.

- Средний менеджмент.

Экспертные компетенции: личностные, стратегия, регулярный менеджмент, цифровые технологии, проектное управление, коммуникация, управление изменениями, взаимодействие с ИИ.

Фокус: операционная реализация ИИ-стратегии.

Роль: главные исполнители ИИ-трансформации.

- Неформальные лидеры.

Экспертные компетенции: личностные, создание контента с ИИ.

Фокус: демонстрация возможностей ИИ через практические результаты.

Роль: культурные амбассадоры и вдохновители.

- ИИ-специалисты.

Экспертные компетенции: личностные, проектное управление, цифровая эрудиция ИИ, работа с данными, безопасность ИИ, создание контента.

Фокус: техническая экспертиза и консультирование.

Роль: центры компетенций и внутренние эксперты.

- Обычные сотрудники.

Практические навыки: создание контента с ИИ, взаимодействие с ИИ.

Фокус: эффективное использование ИИ-инструментов в работе.

Роль: активные пользователи ИИ-решений.

Компетенции эпохи верификации

Отдельный срез — компетенции специалистов, чью работу ИИ уже перестраивает (инженерный кейс мы разобрали в Главе 18). Дефицитными становятся пять компетенций, и только одна из них «про ИИ»: **доменная глубина** — физика процесса, а не широта: верифицировать выход модели может только тот, кто понимает, почему так, а не что написано; **постановка задачи и работа с ограничениями** — плохое ТЗ плюс ИИ равно быстро и качественно оформленная ерунда; **грамотность в работе с данными** — откуда обучающая выборка, что такое дрейф модели, почему ассистент, безупречный на серийной задаче, врёт на единичной; **трассируемость и ответственность за решение** — подпись останется человеческой, и специалист должен уметь объяснить решение, а не сказать «так сказала модель»; и роль «переводчика» между предметной областью и специалистами по данным — самая дефицитная и самая недооценённая.

И главный кадровый вывод для программ обучения: он тот же, что и в инженерном кейсе Главы 18, — программы подготовки надо перепроектировать под верификацию, а не под генерацию. **Исчезнет специалист, который был ценен скоростью рук. Останется тот, кто ценен качеством суждения.** Это хорошая новость — но только для тех, кто начал перестраивать подготовку кадров сегодня.

Добавлю ещё один сдвиг эпохи ИИ. Десятилетиями рынок платил за узкую специализацию, а широкая эрудиция считалась «неопределившимся профилем». ИИ разворачивает это правило. Недостающую глубину теперь можно арендовать: модель даёт справку, расчёт и черновик по любой области. А вот связать области между собой, увидеть аналогию, поставить правильный вопрос — этого ИИ без человека не делает.

Поэтому эрудит с ИИ-инструментом становится новой дефицитной фигурой: он один способен и охватить картину целиком, и верифицировать детали с «арендованной» глубиной. Практический вывод для программ развития: не только углубляйте специалистов, но и расширяйте кругозор сильных людей — межотраслевые ротации, смежные домены, работа на стыках. Раньше это было роскошью. Теперь это рычаг.

Резюме главы

Как показывают различные исследования, компании-лидеры вкладываются в первую очередь в процессы и людей, а технологии и конкретные модели идут уже как следствие, а не самоцель. Если вы инвестируете в технологии и забываете про людей, то будете получать дорогие и никому не нужные игрушки. Бизнесу будет всё равно, идей от него не будет, а всё внедрённое будет восприниматься как чужеродное.

Для продвижения изменений необходимо обучение сначала высшего руководства, а затем и остальных сотрудников. Причём важно обучить не менее 10% от всего персонала. Именно 10% команды должны стать агентами изменений, которых будет поддерживать высшее руководство. При этом и само высшее руководство должно быть примером и лидером в использовании ИИ. В противном случае люди не купят идею и не поверят в неё.

Ключевые выводы

- Руководителю не нужно уметь обучать нейросети — нужно уметь ставить задачи, оценивать применимость, считать экономику и управлять рисками ИИ.
- Сдвиг 2026 года: к эмоциональному интеллекту добавились системное мышление, работа с данными и «ИИ-грамотность» — умение работать с цифровыми помощниками становится базовой гигиеной руководителя.

- Матрица компетенций — инструмент диагностики: разные роли (СЕО, средний менеджмент, неформальные лидеры, ИИ-специалисты, обычные сотрудники) требуют разной глубины.

Что с этим делать: Оцените себя и ключевых руководителей по матрице главы (неделя). Включите ИИ-грамотность в программу развития управленцев (квартал).

Вопрос для рефлексии: Какой из компетенций матрицы не хватает лично вам — и что вы сделаете с этим в ближайший месяц?

Заключение

Кто под ударом — и что с этим делать

Кому стоит бояться ИИ больше всего — белым воротничкам, специалистам интеллектуального и творческого труда на базовом и среднем уровне квалификации. И это главная, на мой взгляд, опасность от ИИ.

Да, ИИ поможет высококлассным специалистам. Но вот вопрос: откуда братья следующим поколениям экспертов, если найти работу обычным сотрудникам будет трудно? ИИ повысит порог входа в специальности и заставит нас эволюционировать, менять программы обучения в высших учебных заведениях. Раньше читать было привилегией аристократии, а сейчас это умеют даже пятилетние дети. И ключевой вопрос — а вы готовы к этой эволюции?

Определить роль человека в будущем, где ИИ будет повсеместным, можно ответив на вопрос «что же такое ИИ?». Это автоматизация, цифровизация или цифровая трансформация?

Давайте посмотрим классическое определение автоматизации. Это уменьшение количества и трудоёмкости ручного труда человека в повседневной деятельности.

Эффект от автоматизации — ускорение процессов и снижение нагрузки на людей, снижение риска ошибок из-за человеческого фактора.

Подходит? Да, большинство проектов по ИИ — это как раз автоматизация определённых задач, снижение нагрузки на людей.

Цифровизация — это внедрение цифровых технологий и систем, позволяющее перестроить бизнес-процессы, сделать их эффективнее и гибче, начать работу с данными и принятие решений на их основе.

Эффект от цифровизации — снижение издержек при выполнении процессов, их совершенствование и создание «гибкости».

Суммарный эффект от автоматизации и цифровизации — возможность быстрого масштабирования бизнеса и преодоление кризисных ограничений.

Это тоже подходит под описание ИИ.

Цифровая трансформация — это глобальная перестройка бизнеса и системы управления, процессов с использованием результатов цифровизации и автоматизации для увеличения коммерческого потенциала и роста прибыли.

То есть эффект от цифровой трансформации — создание новых персонализированных продуктов для целевой аудитории в сочетании с кратным снижением внутренних издержек.

Что ж, это тоже функция ИИ.

Получается, ИИ — это технология, которая может и автоматизировать процессы, и цифровизировать (в том числе принимать решения на основе данных), и создавать новые сервисы и услуги. Но раскрыть весь этот потенциал, а затем и совершенствовать, ставя всё новые и более сложные задачи, может только человек.

По самым смелым прогнозам, в течение 20–30 лет должен появиться первый «единорог», состоящий всего из 1 человека. Единорог — это компания, которая достигла оценки в 1 млрд долларов менее чем за 10 лет с момента основания, и ещё не размещала акции на бирже (не проводила IPO). Это будет гений, который сможет делегировать задачи ИИ, используя его на полную и заменяя тем самым команду

специалистов. Если же такой компании потребуется дополнительная «человеческая» помощь, то будут наниматься аутсорс-команды.

Что с этим делать лично вам? Соберу ответ всей книги в одну строку: сместиться из позиции исполнителя в позицию постановщика и верификатора — того, кто ставит задачу ИИ и отвечает за результат. Модель компетенций для этого — в Главе 24, личный контур — в Главе 21. Порог входа в профессии вырастет, но вырастет и ценность тех, кто этот порог прошёл.

Вторая линия обороны — креативность, и о ней мы подробно говорили в Главе 21. Напомню суть: ИИ блестящ в конвергентных задачах и «средней» креативности, но статистически тянет к среднему и не рождает подлинно нового — а идеи тех, кто черпает их из одного источника, сливаются. Поэтому под ударом не «творческие профессии», а усреднённое исполнение в любой профессии. Способность родить идею, которой нет в среднем по рынку, выбрать важное и поставить на него — дорожает с каждым днём массового ИИ.

При этом в современной гонке за искусственным интеллектом доминирует нарратив о необходимости создания всё более крупных и мощных универсальных моделей, сильного искусственного интеллекта. Однако существует стратегически важная альтернатива, недооценённая на фоне этого тренда: фокус на разработке и внедрении специализированных малых и средних моделей (Small Language Models — SLM). Эта парадигма предлагает практические решения ключевых проблем бизнес-внедрения ИИ.

Преодоление барьеров внедрения «сильных» моделей

Соппротивление и саботаж. Внедрение «сильного» ИИ, принимающего решения, вызывает у людей страх и неприятие, что неизбежно ведёт к саботажу. Это делает инвестиции в подобные решения для бизнеса неоправданно рискованными.

Регуляторные риски. Законодательные инициативы, подобные AI Act, относят мощные модели к категории высокорискованных, предполагая

жёсткое регулирование и контроль, что создаёт дополнительные барьеры и издержки.

Экономическая неэффективность. Бизнес не готов строить критически важные процессы на супердорогих решениях (часто реализуемых по подписочной модели), от которых становится сверхзависимым. Инвестиции в такие системы зачастую некупаемы.

В итоге технологическая сложность, высокая стоимость, ограничения, которые ещё и сложно продать бизнесу и людям.

Преимущества специализированных моделей

Принятие и доверие. Локальные, специфичные решения, решающие конкретные прикладные задачи, воспринимаются людьми спокойнее, так как не вызывают экзистенциальных страхов перед «сильным» ИИ.

Бизнес-эффективность и доступность. Такие модели приносят измеримую пользу (эффект), за которую бизнес готов платить. Их относительно низкая стоимость и требования к инфраструктуре делают их доступными не только для крупных корпораций, но и для малого и среднего бизнеса (МСБ). Это открывает поток инвестиций в разработку.

Фокус на качестве и оптимизации. Разработчики могут концентрироваться на создании высокоэффективных моделей для узких задач, постоянно их совершенствуя, «сжимая» (применяя методы дистилляции, квантизации) и повышая их «интеллектуальность» в рамках специализации.

Снижение затрат на инфраструктуру. Использование сжатых, оптимизированных моделей позволяет повышать производительность систем без необходимости постоянного наращивания дорогостоящей вычислительной инфраструктуры.

Направления развития в области специализированного ИИ

- Гибридные архитектуры (MoE).

Комбинирование нескольких специализированных моделей (экспертов) для решения комплексных задач. Например, объединение моделей с

сильной логической или математической составляющей (критически важных для аналитики и расчётов) с моделями, специализирующимися на генерации текста или обработке естественного языка. Это позволяет нивелировать слабые стороны отдельных моделей и повышать общую надёжность системы.

Моё мнение, что гибридные архитектуры и разделение моделей будет даже на уровне оборудования — для каждой модели будет свой небольшой сервер.

- Борьба с галлюцинациями и повышение надёжности.

Разработка методов и архитектур, минимизирующих риски генерации недостоверной информации («галлюцинаций»), особенно критичных в областях, требующих точности (финансы, аналитика, медицина).

- Безопасное взаимодействие и минимизация прямого промптинга.

Прятать модели за формализованными интерфейсами, где пользователь получает уже готовые рекомендации или результаты, а не взаимодействует с ИИ напрямую через открытый текстовый ввод (prompt). Это резко снижает риски prompt injection атак и ошибок, связанных с некачественной или злонамеренной обработкой пользовательского ввода.

Фундамент — качество данных и защита

Успех любого ИИ, особенно специализированного, напрямую зависит от качества данных. Критически важными становятся проекты, предшествующие внедрению.

- Анализ и оптимизация бизнес-процессов — чёткое понимание того, какие данные нужны и как они генерируются.
- Проработка потоков данных — создание иерархических моделей данных, описание схем их движения.
- Обеспечение качества данных — внедрение процессов контроля и улучшения качества входных данных.

- Защита данных и моделей — построение механизмов безопасности для предотвращения отравления данных, компрометации датасетов. Внедрение точек контроля, бэкапов моделей и данных. Вся инфраструктура должна обеспечивать безопасность данных и целостность моделей.

Практическая реализация

Данная стратегия находит воплощение в проектах по созданию «субпродуктов» — специализированных решений, построенных на тщательно подобранных моделях. Каждая такая модель имеет свою специфику и фокусируется на решении конкретных прикладных бизнес-задач. Это позволяет создавать экономически эффективные, безопасные и легко внедряемые инструменты ИИ, дающие быстрый и измеримый результат.

Например, ИИ для техподдержки. Да, внутри обычная языковая модель, но обвязка инструментарием, проработка базы знаний и организационной структуры дают возможность решать автоматически до 80 % запросов к технической поддержке. Это другая метрика, чем упомянутое выше снижение потока заявок примерно на 40 %: там чат-бот разгружает первую линию на входе, здесь система закрывает обращение целиком. Направление подтверждает и прогноз Gartner: к 2029 году агентный ИИ будет самостоятельно решать до 80 % типовых обращений в поддержку.

Я считаю, что смещение фокуса с гонки за созданием всеобъемлющих «сильных» моделей в сторону развития экосистемы специализированных, эффективных и доступных решений представляет собой фундаментальный сдвиг. Это экономически оправданная альтернатива, которая будет влиять и на науку. Ведь за науку кто-то должен платить.

Этот подход позволяет преодолеть барьеры внедрения, обеспечить доверие, повысить надёжность систем за счёт гибридных архитектур и безопасных интерфейсов. В конечном итоге, ускорить реальное

применение ИИ для решения бизнес-задач на всех уровнях — от глобальных корпораций до малого и среднего бизнеса.

В качестве завершения приведу пару фактов.

Сейчас многие компании пытаются формировать внутренние ИИ-команды. На мой взгляд, это тупиковая ветвь. Качественных специалистов мало, а крутым платить столько же, сколько они могут получать на рынке, трудно. То есть такие компании в любом случае будут привлекать слабых специалистов. Но допустим, компания смогла собрать сильную команду. Какая у неё будет мотивация развиваться? А самое главное, сможет ли она развиваться с той же скоростью, как и специализированные команды? Кто быстрее развивается: тот, кто делает 3–5 проектов в год, или кто 50?

Названные выше 95% — это данные исследования Массачусетского технологического института (MIT). Оно показало, что искусственный интеллект действительно может увеличивать выручку компаний, но это происходит лишь примерно в 5% случаев. Остальные участники (в выборку вошли 300 организаций, уже внедривших ИИ) не зафиксировали ощутимого финансового эффекта, несмотря на значительные инвестиции. Для большинства компаний ИИ оказался экономически бесполезным инструментом именно в текущей конфигурации внедрения.

«Пилотные проекты некоторых крупных компаний и молодые стартапы действительно преуспевают в использовании генеративного ИИ, — отмечает Адитья Чаллапалли (Aditya Challapally), один из авторов исследования. — Стартапы, возглавляемые 19- или 20-летними, например, увидели скачок выручки с нуля до \$20 млн за год».

По оценке экспертов MIT, причина низкой результативности кроется не в качестве современных моделей, а в подходах компаний: они инвестируют в базовую интеграцию нейросетей, но не вкладываются в донстройку под внутренние процессы и в системное обучение сотрудников. «Универсальные инструменты, такие как ChatGPT, удобны для индивидуального использования, однако в корпоративной среде их

потенциал раскрывается лишь при целевой адаптации и включении в регламенты работ», — подчёркивает Чаллапалли.

Ключевым фактором выступает и структура затрат. Более половины бюджетов на генеративный ИИ направляется на инструменты для продаж и маркетинга, тогда как наивысшая окупаемость, по данным MIT, обеспечивается за счёт цифровизации бэк-офиса: отказа от аутсорсинга бизнес-процессов, сокращения расходов на внешние агентства и оптимизации внутренних операций.

Компании, достигшие роста выручки, последовательно инвестировали в развитие и кастомизацию ИИ под собственные задачи. В этих кейсах организации выбирали не массовые решения, а узкоспециализированные инструменты и привлекали их поставщиков к совместной доработке. По данным MIT, именно такие пилоты доходили до эксплуатации примерно в 67 % случаев против примерно 33 % у полностью внутренней разработки — от кода до внедрения и обучения. Обратите внимание: это та же закономерность, что и в главе о малом бизнесе. «Покупайте готовое» не значит «берите коробку для всех»: выигрывает специализированное внешнее решение с совместной доработкой — но не собственная разработка с нуля.

Данные MIT также показывают: работа «в одиночку» существенно повышает вероятность неудач. Многие опрошенные компании неохотно раскрывали частоту провалов, однако приобретённые у внешних поставщиков решения демонстрировали более надёжные финансовые результаты. Среди дополнительных факторов успеха экспертами названы предоставление полномочий линейным руководителям (а не только централизованным ИИ-лабораториям) и выбор инструментов, которые способны глубоко интегрироваться в процессы и эволюционировать со временем.

Постскриптум 2026: четыре сигнала зрелости рынка

К моменту, когда я готовил третью редакцию этой книги, рынок успел дать концентрированное подтверждение прогнозов, которые я делал в предыдущих главах и первом издании. Весной и в начале лета 2026 года

накопились четыре новости, которые по отдельности выглядят рутинно, а вместе — переописывают картину. Я разберу их здесь как итоговую иллюстрацию.

Сигнал 1. Anthropic окончательно поворачивает в сторону МСБ

13 мая 2026 года Anthropic запустил пакет **Claude for Small Business** — набор из примерно 15 готовых агентных рабочих процессов с коннекторами к QuickBooks, PayPal, HubSpot, Canva, Docusign, Google Workspace и Microsoft 365. Под одной подпиской предприниматель получает закрытие месяца, выставление счетов, контроль маржинальности, проверку договоров, сортировку лидов и подготовку контентной стратегии. Параллельно стартовал офлайн-тур по десяти городам США с бесплатными воркшопами по 100 предпринимателей на каждый город.

Это не запуск нового продукта в обычном смысле. Это смена позиционирования у одного из двух глобальных лидеров рынка. Anthropic публично озвучил то, о чём в этой книге шла речь: на МСБ приходится около 44% ВВП США, а проникновение ИИ там значительно отстаёт от крупного бизнеса. Программное обеспечение исторически создавалось под корпорации, венчурные стартапы и массового потребителя — но не под ландшафтную фирму на 30 человек или агентство недвижимости на 50 сотрудников. Anthropic закрывает этот разрыв.

Не менее важно — *как* они зашли. Не с «открытым чатом», а с коробкой под понятные задачи дня. Не с обещанием «научите модель под себя», а с готовыми сценариями, которые работают с первого дня. И ещё одно — Claude не действует бесконтрольно: пользователь сам иницирует каждый рабочий процесс, утверждает план и подтверждает результат, прежде чем что-либо будет отправлено, опубликовано или оплачено. Это ровно та парадигма «формализованных интерфейсов» и «минимизации прямого промптинга», о которой шла речь несколькими страницами выше.

Сигнал 2. В России картина та же — и она тоже про МСБ

3 марта 2026 года Яндекс B2B Tech опубликовал данные по платформе Yandex AI Studio:

- проанализировано более 7 500 ИИ-агентов;
- **86% пользователей ИИ-агентов в России — представители малого и среднего бизнеса;**
- около 25% агентов работают в техподдержке, банкинге и HR-консультациях.

Сценарии — отраслевые и узкие: застройщики и риелторы используют агентов для подбора объектов, управляющие компании — для обработки жалоб в ЖКХ, научные лаборатории — для диагностики оборудования, кто-то — для анализа новостей. Это не «GPT для офиса», а конкретные задачи под конкретный профиль бизнеса.

Параллельно Яндекс представил инструменты, позволяющие использовать ИИ в облачной инфраструктуре без выхода в интернет — прямой ответ на главное возражение МСБ-сегмента про безопасность данных. Российский рынок повторяет глобальную траекторию с опозданием около полугода — и здесь тоже выигрывают узкие коробочные решения с правильно упакованной безопасностью.

Сигнал 3. Парадокс продакшена: 74% откатов

14 мая 2026 года шведский коммуникационный провайдер Sinch опубликовал исследование **AI Production Paradox** — опрос 2 500 руководителей ИИ-проектов в разных странах и индустриях. Главный вывод оказался жёстким:

74% компаний откатали или полностью исключили уже запущенных ИИ-агентов в клиентской поддержке. Под «откатом» имеется в виду снятие с боевой эксплуатации, а не отмена проекта до запуска.

Самое контринтуитивное: у компаний, которые сами называют свою систему контроля над ИИ «полностью зрелой», процент откатов **выше — 81%**. Sinch объясняет это так: зрелые команды откатывают чаще не потому, что у них хуже модели, а потому, что они быстрее видят

проблемы и быстрее на них реагируют. Цифра при этом не зависит ни от региона, ни от индустрии, ни от размера бизнеса — проблема не лечится деньгами.

Дальнейшие цифры — в том же тоне:

- 84% ИИ-инженерных команд минимум половину рабочего времени тратят на инфраструктуру безопасности, а не на разработку модели;
- 75% компаний называют доверие, безопасность и комплаенс в тройке главных приоритетов ИИ-программ — выше, чем саму разработку ИИ (63%).

Это подтверждает и более ранний прогноз Gartner лета 2025 года: половина компаний, рассчитывавших серьёзно сократить штат клиентской поддержки за счёт ИИ, откажется от этих планов до 2027 года. Их вице-президент сформулировал прямо: полностью безоператорный контакт-центр пока технически неосуществим и операционно нежелателен.

Оговорка по методу: Sinch — заинтересованная сторона, она продаёт коммуникационную инфраструктуру под ИИ-агентов. Но даже с этой поправкой цифра в 74% укладывается в общий тренд последнего года.

Сигнал 4. Лето 2026: доступ к передовым моделям стал управляемым ресурсом

Четвёртый сигнал пришёл уже после трёх весенних. В июне 2026 года США экспортной директивой на три недели исключили доступ к самым мощным моделям Anthropic по всему миру (подробный разбор — в Главе 7). Вернувшаяся «экспортная версия» флагмана, по независимым бенчмаркам практиков, заметно потеряла в качестве на бизнес-задачах, а глава OpenAI предложил «МАГАТЭ для ИИ» — международную сертификацию доступа к передовым системам.

Для стратегии это значит: сама возможность пользоваться передовой моделью превратилась в управляемый извне ресурс. Продукты, критически зависящие от одной конкретной модели, несут регуляторный риск, который нельзя закрыть контрактом. Выигрывают

решения, где модель — заменяемый компонент, а ценность — в сценарии, данных и контуре контроля.

Что эти сигналы говорят вместе

По отдельности — четыре новости. Вместе — карта рынка, которая укладывается ровно в ту логику, которую я выстраивал в этой книге.

Во-первых, рынок сменил основной сегмент. Главный покупатель ИИ-продуктов в горизонте 2–3 лет — не корпорация и не «ChatGPT-пользователь», а малый и средний бизнес. Anthropic и Яндекс заходят туда одинаково, с двух разных концов мира.

Во-вторых, базовая форма продукта поменялась. Побеждает не «открытый чат», а коробка под понятную задачу дня — рабочий процесс с понятным входом, понятным выходом и понятной зоной ответственности. Универсальные ассистенты уходят в инфраструктурный слой, видимая ценность пользователя — в конкретных сценариях.

В-третьих, бутылочное горлышко рынка — не модель, а безопасность и доверие. Откаты Sinch идут не потому, что модели стали хуже. Они идут потому, что у компаний нет операционной готовности эксплуатировать автономные системы в продуктивной среде. 84% инженерных команд тратят минимум половину времени на безопасность — это новая норма, а не временный перекося.

В-четвёртых, полная автономия проиграла фазу 2025 года. Лозунг «уберите всех сотрудников и поставьте агента» индустрия фактически уже признала нереалистичным в горизонте 2–3 лет. И Klarna, и другие громкие кейсы 2024–2025 годов «бумерангом» вернули людей. Откаты Sinch — это тот же процесс, только на уровне технологии, а не штата.

В-пятых, возвращается фигура человека — не как препятствие, а как обязательный архитектурный элемент. ИИ-системы 2026+ строятся по принципу human-in-the-loop by design, а не из-за регуляtorики постфактум.

В-шестых, зависимость от конкретной фронтир-модели стала самостоятельным риском: доступ к ней может измениться за ночь — по решению регулятора, а не рынка. Это финальный аргумент в пользу той же архитектуры: модель — компонент, продукт — ценность.

Если коротко: рынок переходит от парадигмы «ИИ заменит человека» к парадигме «ИИ усиливает конкретного человека под его контролем». То есть к тому, о чём шла речь в главах про специализированные модели, цифровых советников и безопасное взаимодействие через формализованные интерфейсы.

Четыре сигнала зрелости рынка — 2026

Весна-лето 2026 года: рынок повернул. Вот факты



Рынок повернул: не «ИИ заменит человека», а «ИИ усиливает конкретного человека в конкретном сценарии»

Профиль выигрывающего ИИ-продукта

Из этих сигналов складывается профиль успешного ИИ-продукта для МСБ на ближайшие 2–3 года. Шесть характеристик, которые проверяет каждый из них:

- **Узкий, но глубокий объём задач.** Ассистент для конкретной роли (собственник, директор, бухгалтер, продавец, риелтор) или конкретной задачи — а не «помощник на всё».
- **Коробочная упаковка.** Готовые рабочие процессы и шаблоны вместо «сначала научите модель». МСБ платит за результат с первого дня.

- **Человек в контуре по умолчанию.** Пользователь инициирует процесс, утверждает план, подтверждает результат перед публикацией, отправкой или оплатой.
- **Безопасность как первое сообщение, а не мелкий шрифт.** «Ваши данные не уходят в обучение», «вы видите каждый шаг», «никаких записывающих операций без подтверждения».
- **Интеграция с существующим стеком.** Коннекторы к тому, что у клиента уже стоит — 1С, amoCRM/Bitrix24, СБИС, СКБ Контур в России; QuickBooks, HubSpot, Google Workspace на Западе.
- **Независимость от конкретной модели.** Модель — заменяемый компонент, а ценность продукта — в сценарии, данных и контуре контроля. Решение, завязанное на одного поставщика, несёт регуляторный риск, который не закрывается контрактом.

Что это значит на практике

Три практических вывода для тех, кто работает с ИИ прямо сейчас.

Если вы внедряете ИИ в свою компанию — не пытайтесь сразу строить автономного агента «под ключ». Начните с узкого human-in-the-loop сценария: одно подразделение, одна задача, один пользователь с явным правом отката. 74% откатов Sinch — это уроки тех, кто пошёл сразу в автономию.

Если вы строите ИИ-продукт — переоцените позиционирование. «Открытый чат для всего» — это категория 2023–2024 года. «Коробка под конкретную работу с явным человеком в контуре и сильной коммуникацией безопасности» — категория, в которую сейчас зашли Anthropic и Яндекс. Если вы стоите ближе к первой парадигме — у вас остаётся 12–18 месяцев на разворот.

Если вы пишете ИИ-стратегию на 2027–2028 — закладывайте, что бюджет на безопасность, доверие и комплаенс будет сопоставим с бюджетом на саму разработку. 84% ИИ-команд уже работают именно в такой пропорции.

Чем я хочу закончить

Если предыдущие главы и заключение были аргументацией прогноза, то весна 2026 года — его фактическое подтверждение. ИИ уходит из режима «магического агента, который всё сделает сам» в режим коробочного ассистента, который усиливает конкретного человека под его контролем. Главный сегмент — МСБ. Главная форма — узкий рабочий процесс с человеком в контуре. Главное конкурентное преимущество — не качество модели, а безопасность и доверие.

Этого поворота можно было ждать ещё пару лет назад, опираясь на здравый смысл и логику развития технологий. Кто увидел его раньше и начал строить продукт под эту парадигму, а не под «автономного агента, заменяющего человека» — выйдет в следующую волну на правильных рельсах.

Ну а ниже то, чем бы я хотел закончить третью редакцию этой книги.

7 принципов «С неба на землю»

Если из всей книги оставить одну страницу — пусть это будет она.

1. **70/30.** ИИ забирает рутину, человек отвечает за контекст, ограничения и решение. Убрать человека полностью — значит убрать качество.
2. **Данные — потолок пользы.** Модель не может быть умнее того, что в неё поступает. Мусор на входе — уверенный мусор на выходе. Помните это, когда взаимодействуете с ИИ.
3. **Качество ответа = качество постановки задачи.** ИИ делает ровно то, что его попросили. Сочетание данных и постановки задачи — основа для успеха.
4. **Цена ошибки определяет степень формализации.** Для идей — свободный диалог; для решений и денег — структура, проверка, человек в контуре.
5. **Модель — компонент, продукт — ценность.** Ценность для пользователя создаёт не модель, а сценарии и решение задач пользователя. Нельзя строить продукты вокруг одной модели. Модели будут ограничиваться и меняться.

6. **Специализация бьёт масштаб.** В прикладной задаче настроенная «малая» модель и упаковка в рабочий сценарий обыгрывает универсального гиганта.

7. **Побеждает не технология, а человек, который ею управляет и пользуется.** Именно то, как человек использует ИИ, определяет, будет ли этот опыт успешным.

Спускайтесь с неба на землю — и стройте.

«С неба на землю» — 7 принципов

Вся книга на одной странице

- 1 70/30**
ИИ забирает рутину, человек отвечает за контекст, ограничения и решение
- 2 Данные — потолок пользы**
Модель не может быть умнее того, что в неё поступает
- 3 Качество ответа = качество постановки задачи**
ИИ делает ровно то, что его попросили
- 4 Цена ошибки определяет степень формализации**
Для идей — свободный диалог; для денег — структура и человек в контуре
- 5 Модель — компонент, продукт — ценность**
Нельзя строить продукт вокруг одной модели: модели меняются и ограничиваются
- 6 Специализация бьёт масштаб**
Настроенная «малая» модель в рабочем сценарии обыгрывает универсального гиганта
- 7 Побеждает человек, который управляет технологией**
Именно то, как вы используете ИИ, определяет успех

Спускайтесь с неба на землю — и стройте

Что дальше

Эта книга — о понимании ИИ. Если вы дочитали до этой страницы, у вас уже есть картина: что умеет технология, где её пределы и как она встраивается в управление. Дальше — три маршрута, в зависимости от вашей задачи.

Если вы готовы внедрять — вам нужна методология: отбор use-case, оценка готовности, пилот, масштабирование, эффекты. Это книга «ИИ-2. Практическое руководство» — вторая в ИИ-серии.

Если чувствуете, что компании не хватает цифрового фундамента — начните с серии «Цифровая трансформация для директоров и собственников»: часть 1 — погружение в технологии, часть 2 — системный подход к трансформации.

Если главный вопрос — безопасность, то третья книга серии по цифровой трансформации посвящена кибербезопасности: от социальной инженерии до менеджмента уязвимостей.

Мои статьи, материалы и способы связаться со мной — на сайте chelidze-d.com. Буду рад обратной связи по книге: она, как вы теперь знаете, — самая короткая петля обучения.

Приложения

Приложение А. Глоссарий ключевых терминов

Термин	Определение
ИИ (AI)	Любой математический метод, имитирующий человеческий или иной природный интеллект.
Машинное обучение (ML)	Статистические методы, улучшающие качество задачи с накоплением опыта и дообучением.
Глубокое обучение (DL)	Подвид ML с многослойными нейросетями; основа современных моделей.
Генеративный ИИ (GenAI)	ИИ, создающий новый контент (текст, изображения, код, аудио) по запросу.
LLM	Большая языковая модель, предобученная на огромных текстовых массивах.
Промпт	Текстовый запрос-инструкция к модели.
Токен, токенизация	Токен — единица текста, которой оперирует модель (часть слова); токенизация — разбиение запроса на такие единицы. В токенах считаются лимиты и стоимость работы.
RAG (retrieval-augmented generation)	Подключение модели к доверенным источникам для подстановки актуальных фактов; снижает галлюцинации.
Fine-tuning	Донастройка модели под конкретный стиль, данные или задачу.
Few-shot / zero-shot	Работа модели по нескольким примерам или только по текстовому описанию задачи.

Системный промпт	Инструкция, задающая роль, правила и формат для всех обращений к ассистенту.
Инъекция промпта	Манипуляция запросом ради обхода фильтров или непредусмотренных действий.
Галлюцинации	Правдоподобный, но недостоверный контент, сгенерированный ИИ.
ИИ-ассистент	Сервис, работающий по запросу и дающий рекомендации.
Ко-пилот	Решение, автоматизирующее отдельные задачи под надзором человека.
ИИ-агент	ИИ, действующий автономно и взаимодействующий с другими системами.
Мультиагентная система	Несколько агентов, координирующих работу между собой.
ИИ-оркестратор	Верхнеуровневый ИИ, декомпозирующий запрос и распределяющий задачи между моделями.
Слабый ИИ (ANI)	Узкоспециализированный ИИ для конкретных задач — всё, что есть сегодня.
Сильный ИИ (AGI)	Гипотетический ИИ, решающий любые интеллектуальные задачи наравне с человеком.
Суперсильный ИИ (ASI)	Гипотетический сверхинтеллект, превосходящий человека во всех областях.
Мультимодальность	Одновременная обработка текста, аудио, изображений и данных датчиков.
СППР (DSS)	Система поддержки принятия решений; цифровой советник.
Shadow AI	Теневое использование ИИ сотрудниками без согласования и политики данных.
Data-driven / informed / inspired	Решения на основе данных (оперативные) / с учётом данных (тактические) / по идеям из данных (стратегические).
Baseline / KPI	Исходный замер и целевые показатели, по которым оценивают эффект решения.

on-premises / SaaS / IaaS / PaaS	Развёртывание на своих серверах / готовый сервис / аренда инфраструктуры / аренда платформы.
Цифровой двойник	Виртуальная копия объекта или процесса, на которой моделируют сценарии до решений в реальности; по мере зрелости модель обогащается эксплуатационными данными и работает синхронно с оригиналом. Уровни зрелости разобраны в книге «ЦТ-1. Погружение».
Нейроморфные вычисления	Архитектура «вычисления в памяти» по принципам мозга; экономит энергию.
Кубит / суперпозиция	Квантовый «бит», способный быть в нескольких состояниях одновременно.
МСБ	Малый и средний бизнес — компании до ~250 сотрудников; главный сегмент роста ИИ-агентов в 2026 году.
Контекстное окно	Объём информации, который модель «держит в голове» в рамках одного диалога.
Инференс	Работа уже обученной модели — каждый ответ на запрос. Главная статья эксплуатационных затрат ИИ.
Модель рассуждений (reasoning)	Модель, «думающая» перед ответом. Сильна в сложных задачах, но дороже и хуже в точном пересказе фактов.
Компактная модель (SLM)	Малая специализированная модель. В узкой задаче часто дешевле и точнее универсальных гигантов.
Открытая модель (open source)	Модель с открытыми весами — можно развернуть в собственном контуре, без внешнего провайдера.
Машинное зрение	ИИ для анализа изображений и видео: контроль качества, безопасность, обслуживание оборудования.

Приложение Б. Чек-лист: где нужен ИИ, а где классическая автоматизация

ИИ-решение предпочтительно, когда:

- задача имеет высокую сложность или много факторов, а правила нельзя явно сформулировать;

- доступны большие массивы данных, из которых можно выявить закономерности;
- нужен анализ неструктурированных данных (тексты, изображения, аудио);
- требуется быстрая адаптация к новым данным без переписывания правил;
- решение зависит от вероятностных оценок.

Классическая автоматизация предпочтительна, когда:

- задача имеет ограниченную сложность и вариативность;
- критичны интерпретируемость, скорость отклика и надёжность;
- нет данных для машинного обучения.

Важно: этот чек-лист — про применимость ИИ как класса технологии. Готовность организации, выбор и приоритизация конкретных кейсов, критерии запуска и остановки проекта — это уже вопрос внедрения. Он подробно разобран в книге «ИИ-2. Практическое руководство»: оценка готовности — Глава 7; выявление и приоритизация идей — Глава 16; красные флаги и правило остановки проекта — приложение 2.

Приложение В. Первые 7 дней с ИИ: план для новичка

Семь дней по 15–20 минут — и ИИ перестанет быть абстракцией. Правила недели: не вставляйте чувствительные данные; проверяйте факты; сохраняйте удачные промпты — это начало вашей личной библиотеки.

День 1. Первый диалог. Заведите аккаунт в одном ИИ-сервисе и попросите: «Объясни, чем занимается моя компания [описание], как школьнику». Оцените, что ИИ понял, а что переврал.

День 2. Выжимка. Отправьте длинное письмо или статью: «5 главных пунктов + что мне с этим делать». Сравните со своим прочтением.

День 3. Черновик. Поручите написать рабочее письмо по шаблону №2 из Главы 6. Отредактируйте — заметьте, сколько времени сэкономили.

День 4. Разбор документа. Задайте роль («ты — юрист...») и попросите риски публичного документа или шаблона договора. Конфиденциальное — не отправляем.

День 5. Декомпозиция. Разложите реальную задачу этой недели на шаги (шаблон №6) и попросите ИИ задать вам уточняющие вопросы. Обратите внимание, какие вопросы вы сами себе не задали.

День 6. Сравнение. Дайте ИИ сравнить два решения, между которыми вы сейчас выбираете, — таблицей с критериями (шаблон №8).

День 7. Свой сценарий. Выберите одну повторяющуюся операцию, соберите под неё постоянный промпт и внесите его использование в календарь. Это ваш первый внедрённый ИИ-процесс.

Приложение Г. Полный план действий: 24 шага из 24 глав

Все блоки «Что с этим делать» книги, собранные в один рабочий план. Используйте как чек-лист внедрения: одна строка — одна глава, отмечайте сделанное.

Глава	Что делать
Глава 1. Знакомство с ИИ	Прежде чем брать ИИ под задачу, прогоните её через критерий «ИИ vs классика» и проверьте, есть ли у вас пригодные данные.
Глава 2. Виды искусственного интеллекта	Не стройте планы на «скором AGI»; делайте ставку на специализированные модели, а для агентов закладывайте правила, лимиты ресурсов и эскалацию.
Глава 3. А что может слабый ИИ и общие тренды	Беритесь за реально болезненную задачу, а не за модную, и заранее договоритесь, как будете измерять эффект; введите политику по теневому ИИ. Методику выбора и приоритизации кейсов и расчёта эффекта см. в книге «ИИ-2. Практическое руководство».
Глава 4. Генеративный ИИ	Пробуйте ИИ лично — на своих задачах, в бесплатных или бюджетных сервисах: это самый дешёвый способ понять технологию. А в компании начинайте с готовых решений — по данным MIT, пилоты на внешнем решении с доработкой под задачу доходят до эксплуатации вдвое чаще собственных разработок, примерно 67 % против 33

	%.	Не сокращайте людей под GenAI и вкладывайтесь в хорошее ТЗ.
Глава 5. Большие языковые модели (LLM)		Для корпоративных задач подключайте LLM к своим данным через RAG и не стройте критичные процессы на «голом» чате.
Глава 6. Промпт-инжиниринг: дисциплина постановки задач для ИИ		Внедрите B-R-I-E-F как корпоративный стандарт постановки задач и заведите общую библиотеку проверенных промптов.
Глава 7. Регулирование ИИ		Уже сейчас закладывайте маркировку ИИ-контента, логирование и распределение ответственности и следите за регуляторикой своей отрасли.
Глава 8. Безопасность ИИ		Легализуйте ИИ и дайте безопасную внутреннюю альтернативу, тестируйте на prompt injection, храните системные промпты как секреты, а для агентов введите минимальные привилегии и «стоп-кран». Если агентов несколько — ещё общий канал связи и арбитра.
Глава 9. Системы поддержки принятия решений, или Цифровые советники		Не стройте и не ждите «всезнающего» советника — начните с узкого, специализированного на ваших регламентах, и наведите порядок в процессах и данных до внедрения.
Глава 10. Нейроморфные и квантовые ИИ		Не ждите квантовой революции для текущих задач; следите за новостями внедрения нейроморфных технологий для автономных устройств и инвестируйте в компетенции, а не в «железо» сейчас.
Глава 11. Стратегия и организационная структура		Проверьте, есть ли ИИ в вашей стратегии как отдельное направление с целями и метриками (неделя). При формировании функции ИИ заложите распределённую модель и развитие по трём осям — технологии, система управления, компетенции людей (квартал).
Глава 12. Бизнес-процессы		Выберите один процесс с понятным входом и выходом и подготовьте его описание и регламент с помощью ИИ (неделя). Оцените зрелость ключевых процессов, прежде чем планировать ИИ-проекты на них (квартал).

Глава 13. Управление проектами, изменениями и продуктами	Прогоните текущий проект по списку «7 грехов» (неделя). Внедрите ИИ-резюме статусов и встреч в одном проекте (квартал).
Глава 14. Коммуникация	Включите ИИ-резюме для одного типа регулярных встреч (неделя). Введите стандарт фиксации «решение — ответственный — срок» (квартал).
Глава 15. Цифровые технологии, работа с данными и кибербезопасность	Проведите аудит источников данных для первого ИИ-кейса (неделя). Составьте дорожную карту качества данных: владельцы, форматы, контроль (квартал).
Глава 16. Практики регулярного менеджмента	Автоматизируйте создание карточек задач из сообщений и встреч (неделя). Свяжите ИИ-ассистента с трекером команды и ритуалом статусов (квартал).
Глава 17. Бережливое производство и теория ограничений систем (ТОС)	Выберите один процесс и снимите по нему 8 потерь (неделя). Только после упрощения решайте, где в нём нужен ИИ (квартал).
Глава 18. Влияние ИИ на процессы функциональных направлений	Выберите функцию с самой болезненной рутинной и одну задачу в ней, зафиксируйте базовые метрики до старта — это кандидат на пилот (неделя). Проведите пилот и сравните результат с базой (квартал).
Глава 19. ИИ для малого бизнеса и предпринимателей	Выберите одну повторяющуюся боль (обработка заявок, документы, контент) и закройте её готовым сервисом — пилот за неделю, без бюджета.
Глава 20. Примеры и рекомендации для среднего и крупного бизнеса	Выберите из этой главы 2–3 кейса-аналога вашей отрасли и сравните с вашими процессами: что из инфраструктуры и данных у вас уже есть (неделя). Сформируйте портфель из 3–5 инициатив с приоритетами (квартал).
Глава 21. ИИ в работе самого руководителя	Заведите дневник руководителя: 5 минут вечером, запись плюс обратная связь от ИИ (неделя). Соберите личный контур: планирование, проработка решений, подготовка к сложным разговорам (квартал).

Глава 22. Как вести проекты по внедрению технологий	Классифицируйте ваш ближайший ИИ-проект по типу и соберите роли (неделя). Внедрите «единый файл проекта» с ИИ-резюме статусов (квартал).
Глава 23. Системный подход к внедрению и дорожная карта	Разложите свои ИИ-инициативы на зависимые и независимые от корпоративных данных и запустите вторые (неделя). Сведите инициативы в единую карту со связями с ИТ-проектами (квартал).
Глава 24. Что должен знать и уметь руководитель в эпоху ИИ	Оцените себя и ключевых руководителей по матрице главы (неделя). Включите ИИ-грамотность в программу развития управленцев (квартал).