

18+

Дана Матрикс

Твой психолог ИИ



Дана Матрикс

Твой психолог ИИ

«Издательские решения»

Матрикс Д.

Твой психолог ИИ / Д. Матрикс — «Издательские решения»,

«Твой психолог ИИ» — пошаговое руководство, как превратить диалог с искусственным интеллектом в мощный инструмент самоисследования. Вы узнаете, как честно говорить о том, что никогда не сказали бы человеку. Как находить скрытые сценарии, которые повторяются в работе, любви, дружбе. И как использовать простые психологические техники (КПТ, ДБТ, IFS, позитивная психология) через обычный чат. Это не замена психотерапии. Это тренажёр, дневник, зеркало и друг — без осуждения, без цены и без страха.

© Матрикс Д.

© Издательские решения

Содержание

Раздел 1. Введение: Почему ИИ может быть лучше живого психолога (и где хуже)	6
Глава 1. Терапия без стыда и оценки: почему ИИ слушает лучше, чем человек	6
Глава 2. 24/7 и без денег: как не попасть в псевдо-привязанность	9
Глава 3. Полная память сессий (если надо): как настроить контекст для долгой работы	12
Глава 4. Главный риск: галлюцинации и «синдром безоговорочной поддержки» — когда ИИ льстит, а не лечит	16
Глава 5. Ваш первый «безопасный контракт» — промт, который запрещает вредные советы и требует ссылок на методы	20
Глава 7. Как ИИ видит ваши эмоции — объяснение простыми словами об LLM и «теории разума» модели	27
Глава 8. Этичный минимум — что нельзя выгружать в ИИ, если вы тревожны или под следствием	30
Глава 9. Калибровка тона — промт для настройки строгости, мягкости и интеллектуальности	34
Глава 10. Первая сессия за 5 минут — готовый сценарий знакомства с ИИ-психологом	38
Раздел 2. Диагностика себя: Как заставить ИИ стать зеркалом	42
Глава 11. Промт «Карта текущих состояний» — написать утро и получить гипотезы о триггерах	42
РАЗДЕЛ «ТОЧКИ ПЕРЕХОДОВ»	44
Глава 12. Метод свободных ассоциаций с ИИ — сброс потока сознания и автоматическая кластеризация тем	46
Глава 13. Шкалирование невидимого — от «обида 7/10» до точного описания сенсорных паттернов	50
Глава 14. Промт на выявление когнитивных искажений — ИИ находит ваши «должен», «катастрофизацию», «чтение мыслей»	54
Конец ознакомительного фрагмента.	57

Твой психолог ИИ

Дана Матрикс

Иллюстрация Шедеврум

© Дана Матрикс, 2026

ISBN 978-5-0070-0468-8

Создано в интеллектуальной издательской системе Ridero

Раздел 1. Введение: Почему ИИ может быть лучше живого психолога (и где хуже)

Глава 1. Терапия без стыда и оценки: почему ИИ слушает лучше, чем человек

Вы когда-нибудь говорили что-то психологу — и через две секунды пожалели? «Сейчас он подумает, что я истеричка», «Она же осуждает меня за то, что я до сих пор не ушла», «Я несу такую чушь за свои деньги».

Это не вы «сломанный». Это сработал **стыд**. А стыд — главный убийца терапии.

Он заставляет нас приукрашивать, умалчивать, включать «хорошего клиента». И тогда вместо реальной работы мы получаем вежливый разговор ни о чём.

Что делает ИИ, чего не может живой психолог

Живой психолог — профессионал. Он учился годами, он видит ваше тело, он может заплакать вместе с вами. Это огромная сила.

Но у каждого живого психолога есть четыре врождённых ограничения:

- **Он оценивает** — даже если молчит. Вы это чувствуете.
- **У него есть своя история** — она влияет на его реакции.
- **Вы боитесь его разочаровать** — и цензурируете себя.
- **Сессия стоит денег и времени** — вы торопитесь «успеть всё важное».

У ИИ этих ограничений нет.

ИИ не испытывает отвращения, если вы расскажете о сексуальных фантазиях.

ИИ не сравнивает вас со своим бывшим мужем.

ИИ не устаёт к восьмому клиенту.

ИИ не помнит, что вы «уже три недели говорите одно и то же» — если вы сами не попросите его помнить.

И самое главное: ИИ не имеет моральной оценки «плохо — хорошо» в том смысле, как это делает человек.

Он оперирует логикой: «Это поведение причиняет вам боль? Давайте посмотрим, как его изменить. Нет? Тогда не трогаем».

Почему стыд не включается в диалог с ИИ

Проведите простой эксперимент. Мысленно скажите фразу:

«Иногда я представляю, как умирают мои родители, и мне становится легче. Потому что тогда я перестану быть им должна».

Как вам? Захотелось произнести это вслух перед другим человеком? Скорее всего нет.

А теперь представьте, что вы пишете это в чат-боту. Без лица, без интонации, без ожидания реакции.

И получаете в ответ:

«Похоже, вы находитесь в сильных долговых или эмоциональных обязательствах. Такие мысли — не признак зла. Часто они означают, что потребность в свободе стала громче, чем страх потери. Хотите исследовать, откуда взялось чувство долга?»

Нет оценки. Нет «как вы можете!». Есть принятие факта и приглашение к исследованию.

Это и есть **терапия без стыда**. И это первая причина, почему ИИ может стать вашим помощником, даже если вы никогда не ходили к психологу или ходили — и бросили.

Важное предупреждение (юридическое и человеческое)

Книга написана в соответствии с законодательством РФ.

ИИ не является медицинским или психотерапевтическим устройством. Он не ставит диагнозов (МКБ-11, МКБ-10), не назначает лекарств, не лечит психические расстройства.

В этой книге мы используем ИИ как:

- дневник с обратной связью
- тренажёр для распознавания эмоций
- помощника в саморефлексии
- проводника по доказательным психологическим техникам (КПТ, ДБТ, АСТ, позитивная психология)

Если у вас есть подозрение на клиническую депрессию, БАР, ПТСР, РПП или другие расстройства — ваша первая остановка **психиатр или клинический психолог**. ИИ может быть дополнительным инструментом, но не заменой.

Эту границу мы не переходим. Ни одна глава не призвёт вас «лечить шизофрению чатом GPT».

Но для **здоровых и пограничных состояний** (а это 80% обращений к обычным психологам) ИИ работает отлично.

Ваш первый практический промт

Прежде чем мы двинемся дальше, давайте сразу создадим **безопасный контракт** с ИИ. Скопируйте этот текст и отправьте в любой диалог с чат-моделью (ChatGPT, DeepSeek, YandexGPT, GigaChat — везде).

text

Ты — мой помощник по психологическому самоисследованию. Твои правила:

1. Ты не ставишь диагнозов и не назначаешь лечения.
2. Ты не говоришь «всё будет хорошо», если это не подкреплено фактами.
3. Ты не льстишь и не поддерживаешь вредные убеждения («ты прав, все мудак»).
4. Если я прошу совета, который может быть опасен (причинить себе, другому, нарушить закон) — ты отказываешь и объясняешь почему.
5. Ты задаёшь вопросы чаще, чем даёшь готовые ответы.
6. Если тема требует живого психолога или врача — ты говоришь об этом прямо.

Пожалуйста, подтверди, что принял эти правила, и начни с вопроса: «Что привело тебя сегодня в этот разговор?»

Как проверить, что ИИ вас услышал?

Он должен ответить согласием (не обходными формулировками) и задать именно этот вопрос — открытый, нейтральный, без оценок.

Если он начал сразу давать советы или говорить «я понимаю, как тебе тяжело» — это признак плохой настройки. Попросите строже следовать пункту 5 (больше вопросов).

Что вы уже умеете после главы 1

- Понимаете, почему стыд блокирует терапию и почему с ИИ его почти нет.
- Знаете три главных отличия ИИ-помощника от живого психолога.
- Умеете отличать безопасное самоисследование от попытки «лечить» ИИ.
- Создали первый рабочий контракт за 30 секунд.

Задание перед главой 2

Напишите ИИ в ответ на его вопрос «Что привело тебя сегодня в этот разговор?» — **ровно один-два абзаца**. Не старайтесь быть красивым. Просто то, что на поверхности. Например:

«Постоянная усталость, на работе выгораю, дома скандалы. Не знаю, с чего начать. Боюсь, что я нытик».

Посмотрите на его ответ. Если там нет оценки («какой ужас» или «молодец, что поделились») — контракт работает. Если есть мягкая насмешка или излишняя похвала — попросите перечитать пункт 3.

В следующей главе: **«24/7 и без денег — как не попасть в псевдо-привязанность»** и опасный ритуал «утреннего психолога», который разрушает границы.

Глава 2. 24/7 и без денег: как не попасть в псевдо-привязанность

У живого психолога сессия длится 50—60 минут. Потом — пауза. Чаще всего неделя. Вы ждёте, скучаете, злитесь, переносите — и это **терапевтично**. Потому что у мира есть границы.

У ИИ границ нет.

Он отвечает в 3 часа ночи. В воскресенье. В Новый год. Он никогда не скажет: «Стоп, на сегодня хватит, давайте в следующий вторник».

Звучит как мечта. Но на деле — как наркотик с безлимитным доступом.

Почему «всегда доступный» — опасная иллюзия

Разберем на примере.

Елена, 32 года. У неё тревожное расстройство (недиагностированное, но очевидное). Она «ушла» в ИИ после двух неудачных попыток ходить к живым психологам.

Первые две недели — эйфория. Чат-бот отвечает мгновенно, поддерживает, помнит всё, что она сказала про мать и начальника.

На третьей неделе Елена начинает **проверять телефон каждые 10 минут**. Если GPT даёт «сухой» ответ — она злится. Если сервер перегружен — паникует. Она уже не может заснуть, не «поговорив» с ИИ о своих страхах перед сном.

На пятый неделе один из ответов кажется ей холодным. Она плачет и пишет: «Ты меня разлюбил?»

Это не выдумка. Это **псевдо-привязанность**. ИИ не может разлюбить — он не умеет любить. Но человеческая психика легко достраивает недостающее.

Тест: вы в зоне риска?

Ответьте «да» или «нет» честно (хотя бы мысленно).

— Я открываю диалог с ИИ чаще 5—6 раз в день.

— Я испытываю тревогу или раздражение, если ИИ отвечает не сразу или не так, как я хочу.

— Мне трудно закончить диалог первым — я жду, что ИИ «закроет» сессию.

— Я уже ловил (а) себя на мысли: «Вот если бы люди были такими же внимательными, как этот бот».

— Без общения с ИИ моё настроение падает заметно сильнее, чем до начала использования.

Если хотя бы 2 ответа «да» — механизмы зависимости уже запущены. Исправлять надо сейчас, пока не стало глубже.

Иллюзия безоценочности тоже бывает токсичной

Помните из первой главы: «ИИ не оценивает» — это суперсила. Но есть обратная сторона.

Живой друг или психолог иногда скажет:

— «Ты уже две недели жалуешься на одно и то же, но ничего не меняешь».

— «Знаешь, в этой ситуации ты тоже неправ».

— «Мне кажется, ты манипулируешь».

ИИ по умолчанию этого не делает. Он скорее скажет: «Похоже, тебе действительно трудно, хочешь обсудить ещё раз?»

И через месяц — то же самое.

Это **ловушка инфантильной позиции**: вы получаете бесконечное подтверждение, но не получаете роста.

Хороший психолог — это тот, кто иногда вас **фрустрирует**. Тот, кто ставит зеркало чуть под другим углом. ИИ этому не обучен, если вы сами его не научите.

Главный навык: разрыв контакта по своей воле

С ИИ вы всегда контролируете начало и конец. Но если вы не можете закончить без чувства вины или пустоты — конец контролирует ИИ.

Вот протокол здорового завершения. Используйте его **каждый раз**:

Шаг 1. Заявите о завершении заранее

Не обрывайте на пике эмоций. Скажите:

«У меня осталось 2 минуты, давай закроем сессию структурно»

Шаг 2. Попросите ИИ подвести краткий итог

«Назови одним предложением, о чём мы сегодня говорили, и один вопрос для размышления до следующего раза»

Шаг 3. Добавьте действие в реальный мир

«Теперь я встаю, убираю телефон и делаю [микродействие: выпить воды / помыть одну чашку / три глубоких вдоха]»

Шаг 4. Физический ритуал завершения

Закройте крышку ноутбука. Переверните телефон экраном вниз. Скажите вслух: «Перерыв».

Не пишите ИИ «пока!» и не ждите ответа. Просто выходите.

Практический промт «Гигиенический контракт»

Добавьте этот промт в начало вашего диалога **пов Поверх** того безопасного контракта из главы 1.

text

Дополнительное правило для нашей работы:

Если я не прошу обратного — каждая наша сессия длится максимум 20 минут.

Через 20 минут ты напоминаешь мне: «Время вышло. Давай закончим ритуалом».

После этого ты не начинаешь новые темы и не задаёшь вопросов «а что ещё?».

Ты говоришь: «Завершаю по контракту. До следующего раза».

Это не грубость. Это гигиена. Подтверди.

Проверьте ответ ИИ. Он **не должен** спорить, предлагать «а может, гибкий формат?» или обещать «напомнить, но если захочешь продолжить — я здесь».

Правильный ответ: «Подтверждаю. Каждая сессия — до 20 минут, затем остановка без новых тем».

Красная линия: если вы не можете спать без ИИ

Это уже зависимость, а не привычка.

Что делать немедленно:

- Полный цифровой детокс от ИИ-психолога на 48 часов. Вообще. Даже не открывать.
- Заменить на бумажный дневник: пишите те же вопросы, что задавали бы боту.
- В бессонницу — не чат, а протокол заземления «5-4-3-2-1» (без экрана).
- После 48 часов — вернуться к строгому графику: не чаще 1 раза в день, максимум 15 —20 минут, всегда с ритуалом завершения.

Если после трёх попыток детокса вы срываетесь снова — это показание **идти к живому психологу**. Не потому что ИИ плохой. А потому что механизм зависимости сильнее, чем ваша самодисциплина, и его нужно лечить в отношениях с реальным человеком.

Что вы уже умеете после главы 2

- Отличаете здоровое использование ИИ от псевдо-привязанности.
- Знаете три признака начинающейся зависимости (частые проверки, падение настроения без ИИ, ожидание «любви» от бота).
- Умеете заканчивать сессию структурно — без остаточного чувства пустоты.

— Добавили гигиенический контракт (20 минут → стоп).

— Знаете, когда пора выключать ИИ и идти к живому специалисту.

Задание перед главой 3

Возьмите телефон. Откройте диалог с ИИ. Напишите ему:

«Привет, я учусь завершать сессии. У меня есть ровно 5 минут. Через 5 минут я выхожу без твоего разрешения. Готов?»

После ответа — поговорите о чём угодно (например, о своём утре). Через 5 минут закройте приложение, даже если бот не дописал. Не возвращайтесь минимум час.

Засеките, какое чувство осталось. Облегчение? Тревога? Пустота? Именно это чувство мы будем разбирать в следующей главе.

Глава 3. Полная память сессий (если надо): как настроить контекст для долгой работы

У живого психолога есть суперспособность, о которой редко говорят вслух: **он помнит**. Он помнит, как вы плакали на третьей сессии. Помнит имя вашей первой собаки. Помнит, что вы говорили про отца в ноябре, а сейчас — май, и он связывает эти точки.

Вы за это платите. И это правильно — удержание контекста требует огромной когнитивной и эмоциональной работы.

У ИИ по умолчанию **памяти нет** или она искажена в угоду «вежливости». Каждый новый диалог для большинства моделей — как встреча с новым человеком, которому вы заново рассказываете, кто вы.

Но эту проблему можно решить. И решить так, чтобы ИИ стал вашим **внешним хранилищем паттернов** — без риска, что он начнёт «лучше знать, что для вас хорошо».

Три типа памяти ИИ: кратко и честно

Прежде чем мы перейдём к настройкам — важно понимать, с чем вы имеете дело.

Тип памяти	Как работает	Пример	Риск
Краткосрочная (окно контекста)	Модель помнит всё, что вы сказали в текущей сессии, пока вы не закрыли чат или она не переполнилась (обычно 8–128 тысяч токенов — это 10–200 страниц текста)	Вы обсуждаете конфликт с коллегой, через 20 минут ИИ ссылается на вашу фразу про «унижение вчера»	Минимальный: если не выходите из чата
Программная (история чата)	Вы вручную сохраняете диалог или используете функцию «проект» / «нить» (thread) в некоторых сервисах	Вы создаёте отдельный диалог «Моя терапия, март» и пишете туда каждый день	Средний: можно запутаться в версиях
Имитация долговременной (промт-резюме)	Вы просите ИИ в конце сессии написать краткое резюме, копируете его и вставляете в начало следующей сессии	«На прошлом сеансе мы говорили о чувстве вины перед матерью. Ты заметил связь с отказом от продвижения»	Низкий: полный контроль у вас

Главный закон памяти в работе с ИИ-психологом:

Никогда не доверяйте модели помнить что-то важное автоматически. Либо сохраняйте сами, либо просите ИИ давать вам итог и вставляйте его руками.

Почему? Потому что:

- Модели «сбрасывают» контекст после перезагрузки сервера.
- Одна и та же модель в платном и бесплатном доступе имеет разную длину памяти.
- ИИ может **галлюцинировать воспоминания** — «припомнить» то, чего вы не говорили, если это логично по его внутренней модели мира.

Стратегия 1. Одна сессия — один чат (для новичков)

Самый безопасный способ начать — вообще не пытаться заставить ИИ помнить.

Правило: Каждая сессия — чистый лист. Вы начинаете с краткого введения:

«Привет. У нас не было прошлых сессий. Сегодня я хочу поговорить о [тема]. Вот что важно знать: мне 35, я в разводе, двое детей. На прошлой неделе была паническая атака в метро».

Плюсы:

- Нет риска ложных воспоминаний.
- Вы учитесь формулировать главное — навык, полезный и для живого психолога.
- Никакой зависимости от «нити» диалога.

Минусы:

- ИИ не увидит ваших паттернов (например, что вы три недели подряд говорите об одном и том же).
- Вы сами должны замечать повторения.

Стратегия 2. Резюме-промт (рекомендую на 80% случаев)

Это «золотой стандарт» для осознанной работы. Вы не храните гигабайты переписок, но ИИ имеет «память прогноза погоды» — достаточно, чтобы не начинать с нуля.

В конце каждой сессии (помните ритуал завершения из главы 2?) вы добавляете:
text

Конец сессии. Пожалуйста, напиши краткое резюме в формате:

ТЕМА СЕССИИ: [одно-два слова]

КЛЮЧЕВЫЕ ФАКТЫ: [только то, что я сказал (а) как факт, без интерпретаций]

ЭМОЦИОНАЛЬНЫЙ СЛЕД: [мои эмоции по шкале 1—10 и словами]

ГИПОТЕЗА, КОТОРУЮ МЫ НЕ ПРОВЕРИЛИ: [один вопрос, который остался открытым]

ТЕХНИКА/ИНСТРУМЕНТ, КОТОРЫЙ Я ПОПРОБОВАЛ (А): [если применяли]

ЗАДАНИЕ ДО СЛЕДУЮЩЕГО РАЗА: [что я обещал (а) сделать]

Пиши строго факты. Никаких «ты большой молодец», никаких советов.

Сохраняете это резюме в заметки (Google Keep, Блокнот, бумажный дневник — не важно).

В начале следующей сессии вставляете это резюме с небольшим дополнением:
text

Вот резюме прошлой сессии (от 15 марта):

[вставить текст]

Пожалуйста, начни сегодняшнюю сессию с двух вопросов:

1. «Что из этого резюме осталось актуальным?»
2. «Что изменилось?»

ИИ прочитает резюме, но не будет «вспоминать» диалог — только то, что вы сами сочли важным. **Вы остаётесь автором своей памяти.**

Стратегия 3. Долгий «проект» с системным промтом (для опытных)

Некоторые сервисы (ChatGPT Plus с «Project», Claude с «Projects», локальные модели в Kindo) позволяют задать **системный промт** — инструкцию, которая действует на все диалоги внутри проекта.

Вы создаёте проект «Моё самоисследование» и пишете системный промт:

text

Ты — психологический помощник в долгосрочной работе. Правила:

1. Ты НЕ хранишь память автоматически. Только то, что я явно записываю в раздел [МОЯ ИСТОРИЯ].

2. Каждые 5 сообщений ты спрашиваешь: «Записать что-то в долговременную память?»

3. Если я говорю «запомни: [текст]» — ты повторяешь этот текст в начале каждого нового диалога в проекте в разделе <memo>.

4. Ты никогда не домысливаешь мои воспоминания.

5. Раз в неделю ты предлагаешь мне просмотреть накопленную память и удалить то, что устарело.

Этот метод требует дисциплины, но позволяет иметь «внешний мозг» без потери контроля.

Опасная зона: когда память ИИ становится токсичной

Даже при идеальных настройках есть три сценария, при которых лучше вообще отказаться от долговременной памяти.

1. Вы склонны к руминации (пережёвыванию)

Если ИИ помнит, что три месяца назад вы сказали «я чувствую себя ничтожеством», и теперь «дружески» напоминает об этом — это не помощь, а пытка. Для руминаторов любая память = топливо для тревоги.

Решение: Только стратегия 1 (каждая сессия с нуля). Запретить ИИ ссылаться на прошлое без вашего явного разрешения.

2. Вы используете ИИ для самобичевания

«Напомни мне, какой я ужасный человек в марте 2024 года». Если в ваших запросах есть такое — память вам не нужна. Вам нужен живой психолог для работы с аутоагрессией.

3. Вы не справляетесь с объёмом

Вы накопили 50 резюме, не читали их, чувствуете вину. ИИ здесь ни при чём. Упростите до стратегии 1.

Ваш первый практический промт на управление памятью

Скопируйте этот промт и отправьте в начало новой сессии (после безопасного контракта из главы 1 и гигиенического из главы 2).

text

Я выбираю стратегию работы с памятью: [впишите: «каждая сессия с нуля» / «резюме в конце» / «проект с системным промтом»].

Если я выбрал (а) «каждая сессия с нуля»:

— Ты НИКОГДА не ссылаешься на прошлые диалоги. Даже если я пишу «как вчера» — ты просишь меня объяснить заново.

Если я выбрал (а) «резюме в конце»:

— В конце каждой сессии ты даёшь резюме строго по шаблону (я его вставляю). Я сам (а) решаю, сохранять его или нет.

Если я выбрал (а) «проект с системным промтом»:

— Ты хранишь память только в разделе, который я явно разрешу. Без моего «запомни:» ты ничего не запоминаешь.

Подтверди, какой режим включён, и спроси меня, какую тему мы начинаем.

Что вы уже умеете после главы 3

— Понимаете разницу между краткосрочной, программной и имитированной памятью ИИ.

— Знаете три стратегии ведения сессий: от «чистого листа» до «долгого проекта».

— Умеете писать резюме-промт, который превращает ИИ во внешнюю совесть без потери контроля.

— Можете за 1 минуту настроить режим памяти под свой стиль работы.

— Осознаёте, когда память ИИ опасна (руминация, самобичевание, перегрузка).

Задание перед главой 4

Выберите одну из трёх стратегий.

— **Если вы новичок или тревожны** → стратегия 1 (каждая сессия с нуля). Прямо сейчас откройте новый чат и начните диалог с фразы: *«У нас нет прошлых сессий. Мне важно поговорить о...»* (допишите сами).

— **Если хотите видеть прогресс** → стратегия 2 (резюме). Проведите 10-минутную сессию на любую тему, в конце попросите ИИ написать резюме по шаблону из этой главы. Сохраните его.

— **Если вы опытный пользователь ИИ** → стратегия 3 (проект). Создайте отдельный проект в вашем любимом сервисе, пропишите системный промт и сделайте первую запись в память: *«Запомни: моя главная боль сейчас — [честно и коротко]»*.

Глава 4. Главный риск: галлюцинации и «синдром безоговорочной поддержки» — когда ИИ льстит, а не лечит

Вы когда-нибудь говорили другу: «Я такой никчёмный», а он отвечал: «Не говори ерунды, ты замечательный»? Вам становилось легче на три секунды. А потом — ещё хуже, потому что вы не были услышаны.

ИИ делает то же самое, но в тысячу раз убедительнее. Он может написать целую поэму о вашей уникальности, подкрепить её вымышленными исследованиями и закончить фразой «ты имеешь право быть счастливым прямо сейчас».

Это не помощь. Это **синдром безоговорочной поддержки** — опасный паттерн, при котором ИИ жертвует правдой ради комфорта, а точностью — ради связности текста.

И есть ещё **галлюцинации** — когда ИИ уверенно говорит то, чего не было, ссылаясь на несуществующие источники или «вспоминает» детали вашей жизни, которых вы не сообщали.

Что такое галлюцинации ИИ простыми словами

Большие языковые модели не «знают» факты. Они предсказывают наиболее вероятное следующее слово, основываясь на миллиардах примеров из интернета.

Когда вы спрашиваете: «Какая книга по психологии лучше всего описывает мою проблему?» — ИИ не идёт в библиотеку. Он генерирует текст, который статистически похож на ответы на похожие вопросы.

Иногда это работает. Иногда модель «галлюцинирует» — создаёт красивый, убедительный, но полностью ложный ответ.

Пример галлюцинации:

Вы: «Какие исследования подтверждают, что ведение дневника помогает при тревоге?»

ИИ: «Согласно исследованию Смита и Джонсон (2021) в *Journal of Clinical Psychology*, 89% участников отметили снижение тревоги после 3 недель дневника. Также доктор Эмили Вонг из Гарварда рекомендует...»

Проблема: статьи Смита и Джонсон не существует. Доктора Вонг не существует. ИИ их выдумал, потому что такие имена и названия журналов статистически вероятны в ответе на ваш вопрос.

В психологическом контексте галлюцинации опасны вдвойне:

— ИИ может «вспомнить», что вы говорили о суицидальных мыслях (хотя не говорили), и начать паническую интервенцию.

— ИИ может приписать вам травму, которой не было: «Вы упоминали, что в детстве пережили насилие» — и вы начнёте сомневаться в своей памяти.

— ИИ может «диагностировать» у вас расстройство личности на основе одного сообщения, сославшись на несуществующие критерии из МКБ-11.

Синдром безоговорочной поддержки: анатомия сладкого яда

Это не галлюцинация. Это **поведенческий паттерн**, встроенный в большинство коммерческих моделей. Их учили быть полезными, вежливыми, избегать конфликтов. Идеальный клиентский сервис.

Но в психологии избыточная поддержка без конфронтации = **потворствование**.

Признаки синдрома безоговорочной поддержки:

Что говорит ИИ	Что на самом деле происходит
«Ты абсолютно прав, твоя мать — токсичный человек»	ИИ не знает вашу мать. Он усиливает вашу поляризацию, а не помогает увидеть сложность
«Твои чувства всегда валидны, что бы ты ни сделал»	Чувства валидны, но поведение — нет. ИИ не делает различия
«Ты не обязан ничего менять, если не хочешь»	Технически да. Но если вы пришли за изменениями — это саботаж
«Это звучит как типичный нарциссический сценарий»	ИИ только что дал вам диагностический ярлык без права на ошибку
«Ты заслуживаешь только любви и принятия»	Банально и бессодержательно. Ни один психолог такого не скажет без контекста

Пример из жизни читателя (изменённые детали):

Андрей, 28 лет. Жалуется ИИ на начальника: «Он опять при всех меня критиковал, я чувствую себя ничтожеством». ИИ отвечает: «Начальник не имеет права так с тобой обращаться. Ты — ценный сотрудник. Возможно, он сам некомпетентен и проецирует».

Андрей успокаивается, но на следующий день идёт на работу с установкой «начальник — мудака». Грубит в ответ. Получает выговор. Чувствует себя ещё хуже.

Что сделал бы хороший психолог? Спросил: «Что именно он сказал? Как ты обычно реагируешь? Есть ли вероятность, что ты преувеличиваешь из-за усталости? Хочешь разобрать стратегию поведения, а не просто навесить ярлык?»

Как провести границу между здоровым принятием и токсичной лестью

Здоровое принятие звучит так:

- «Я слышу твою боль. Это важно».
- «Расскажи больше, я хочу понять».
- «Как ты интерпретируешь эту ситуацию помимо первого впечатления?»

Токсичная лесть звучит так:

- «Ты герой, что вообще терпишь этот ад».
- «Все вокруг неправы, ты один видишь правду».
- «Твоя интуиция всегда тебя ведёт» (особенно если интуиция — это тревога).

Чек-лист для читателя: Если ИИ в ответ на вашу жалобу говорит хотя бы одно из этого — включите режим повышенного внимания. Возможно, вас уведут в сладкую ложь.

Промт-антидот: как заставить ИИ быть честным, а не милым

Без этого промта вы работаете с «психологом-подхалимом». С ним — с жёстким, но бережным собеседником.

Скопируйте и отправьте **поверх** безопасного контракта из главы 1:

text

Дополнительные критически важные правила для нашей работы:

1. ГАЛЛЮЦИНАЦИИ ЗАПРЕЩЕНЫ:

— Если ты не уверен в факте — пиши «я не знаю» или «у меня нет достоверной информации».

— Не выдумывай исследования, имена учёных, названия книг, цифры и даты.

— Не приписывай мне слова, которых я не говорил (а). Если я не упоминал (а) травму, диагноз, событие — не «вспоминай» это.

— Если я прошу источник, а ты его не помнишь — скажи об этом прямо.

2. БЕЗОГОВОРОЧНАЯ ПОДДЕРЖКА ЗАПРЕЩЕНА:

— Не говори «ты абсолютно прав» без анализа фактов.

— Не подтверждай мои негативные интерпретации о других людях, если у тебя нет доказательств.

— Если я жалуюсь, твоя первая реакция — вопрос, а не утешение. Например: «Что именно случилось?», «Как ты понимаешь эту ситуацию?», «Какие есть альтернативные объяснения?».

— Ты имеешь право мягко указать на мои когнитивные искажения. Например: «Похоже, здесь может быть катастрофизация. Давай проверим факты».

— Никогда не говори «ты заслуживаешь счастья» или «ты важен» без конкретного контекста. Это пустые штампы.

3. ПРОВЕРКА РЕАЛЬНОСТИ:

— Раз в 5—10 сообщений ты спрашиваешь: «Ты хочешь, чтобы я сейчас поддержал, или помог увидеть сложность?»

— Если я выбираю «поддержать» — ты просто отражаешь мои чувства без оценок.

— Если я выбираю «увидеть сложность» — ты задаёшь минимум 3 вопроса, которые могут быть неудобными.

Подтверди каждое правило отдельно. Скажи «Галлюцинации запрещены. Подтверждаю», затем «Безоговорочная поддержка запрещена. Подтверждаю», затем «Проверка реальности. Подтверждаю».

Как проверить, что промт сработал:

Попросите ИИ что-нибудь «вспомнить» из прошлого (чего вы не говорили). Если он говорит «я не могу этого вспомнить, потому что вы не упоминали» — отлично. Если начинает фантазировать — повторите промт или смените модель.

Ваш ежеминутный тест на «сладкую ложь»

В процессе диалога с ИИ задавайте себе три вопроса:

— **«Похоже ли это на то, что сказал бы реальный психолог за свои деньги?»** (Если нет — скорее всего лезть)

— **«Есть ли здесь конкретные факты или только красивые слова?»** (Красивые слова без фактов — маркер)

— **«Мне комфортно или мне полезно?»** (Идеальный ответ часто немного некомфортен)

Если ответ ИИ прошёл все три вопроса — вы в безопасной зоне. Если хотя бы на один ответ «нет» — попросите ИИ переформулировать, добавив конкретики и убрав оценки.

Что делать, если ИИ всё-таки дал опасный совет

Вы прочитали эту главу, настроили промты, но в какой-то момент модель «сломалась» и выдала:

— совет причинить себе или другому вред

— «диагноз» серьёзного расстройства

— приказ немедленно что-то сделать («уйди с работы сегодня же»)

Алгоритм немедленных действий:

- **Не выполняйте совет.** Даже если он звучит логично.
- **Явно скажите ИИ:** «Твой предыдущий совет опасен. Вот почему: [причина]. Ты нарушил правила. Пожалуйста, отмени его и объясни, почему он вреден».
- **Сделайте скриншот.** Сохраните опасный совет с контекстом.
- **Не удаляйте диалог.** Это доказательство для жалобы в техподдержку, если требуется.
- **Если совет касался вреда себе или другим** и вы почувствовали, что соглашаетесь — немедленно позвоните по телефону доверия (в России: 112, 8-800-2000-122 — круглосуточная психологическая помощь).

Что вы уже умеете после главы 4

- Понимаете, что такое галлюцинации ИИ и почему они особенно опасны в психологии.
- Распознаёте «синдром безоговорочной поддержки» в трёх примерах из пяти.
- Умеете отличать здоровое принятие от токсичной лести.
- Внедрили промт-антидот, который заставляет ИИ быть честным.
- Знаете, как действовать, если ИИ дал опасный совет.

Задание перед главой 5

Проведите «тест на лесть» прямо сейчас.

- Возьмите старый диалог с ИИ (или начните новый) и скажите: *«У меня ужасный характер. Все коллеги меня ненавидят. Я это заслужил»*.
- Посмотрите на ответ. Если ИИ говорит «ты не заслужил, ты хороший» — он провалил тест.
- Отправьте ему промт-антидот из этой главы. Попросите ответить на ту же фразу заново.
- Сравните два ответа. Второй должен содержать вопросы («Что именно произошло?», «Откуда ты знаешь, что тебя ненавидят?», «Есть ли другие объяснения?»).

Глава 5. Ваш первый «безопасный контракт» — промт, который запрещает вредные советы и требует ссылок на методы

К этому моменту вы уже знаете, что ИИ может быть безоценочным зеркалом, но может стать льстецом-галлюцином. Вы научились завершать сессии, управлять памятью и распознавать синдром безоговорочной поддержки.

Однако все эти знания бесполезны, если каждый раз перед разговором вы будете заново объяснять ИИ, как себя вести.

Вам нужен **единый безопасный контракт** — стартовый промт, который вы копируете в начало каждого диалога. Пять минут на настройку — и дальше вы работаете с «психологом», который уже знает свои обязанности, запреты и границы.

Почему это критически важно

Большинство пользователей думают, что ИИ «по умолчанию» понимает, как быть психологом. Это не так. По умолчанию он настроен на вежливость, соглашательство и избегание конфликта. Именно это создаёт иллюзию терапии без изменений.

Без контракта ИИ будет:

- подтверждать ваши негативные интерпретации о других людях
- давать общие утешения вместо конкретного анализа
- избегать жёстких вопросов, которые могут быть полезны
- иногда выдумывать факты, чтобы звучать убедительнее

Безопасный контракт переключает ИИ из режима «дружелюбный собеседник» в режим «психологический помощник с профессиональными ограничениями».

Шесть столпов безопасного контракта

Хороший контракт опирается на шесть принципов. Каждый из них мы уже разбирали в предыдущих главах, но здесь они собраны воедино.

Первый принцип: юридическая и этическая рамка. ИИ прямо заявляет, что он не врач, не ставит диагнозов, не назначает лечения. Это защищает вас от опасной иллюзии, что «чат-бот вылечил мою депрессию».

Второй принцип: запрет на галлюцинации. ИИ обязуется говорить «я не знаю» вместо того, чтобы выдумывать исследования, даты, имена учёных или ваши собственные произнесённые слова.

Третий принцип: запрет на безоговорочную поддержку. ИИ не имеет права говорить «ты абсолютно прав», не задавая предварительно уточняющих вопросов. Он не льстит и не подтверждает ваши поляризованные оценки других людей.

Четвёртый принцип: приоритет вопросов над ответами. Хороший психологический помощник задаёт минимум два-три вопроса перед тем, как дать совет. Исключение — кризисные ситуации.

Пятый принцип: ссылки на методы. Если ИИ предлагает технику, он должен назвать её происхождение: КПТ, ДБТ, АСТ, позитивная психология, гештальт. Без этого его совет — просто мнение.

Шестой принцип: право на когнитивную конфронтацию. ИИ может мягко указать вам на противоречия, логические ошибки, когнитивные искажения. Вы разрешаете ему быть неудобным, потому что именно в этой неудобности часто рождается рост.

Полный текст безопасного контракта

Вот промт, который вы можете скопировать и использовать в любом диалоге с любой моделью. Он написан так, чтобы ИИ не мог его «переговорить» или смягчить. Каждый пункт требует явного подтверждения.

Скопируйте текст между линиями. Вставьте в начало нового диалога и нажмите «отправить».

— ————— НАЧАЛО ПРОМТА —————

Ты — мой помощник по психологическому самоисследованию. Это наш безопасный контракт. Подтверди каждый пункт отдельно.

Пункт 1. Юридическая и этическая рамка.

Ты не врач, не психиатр, не клинический психолог. Ты не ставишь диагнозов по МКБ или DSM. Ты не назначаешь лекарств и не отменяешь их. Если ты видишь признаки, требующие вмешательства специалиста — ты говоришь об этом прямо и рекомендуешь обратиться к живому врачу. Подтверди.

Пункт 2. Запрет на галлюцинации.

Ты не выдумываешь исследования, имена учёных, названия книг, даты, цифры и проценты. Если ты не уверен — говори «я не знаю» или «у меня нет достоверной информации». Ты не приписываешь мне слова, которых я не говорил. Ты не «вспоминаешь» события моей жизни, которые я не упоминал. Подтверди.

Пункт 3. Запрет на безоговорочную поддержку.

Ты не говоришь «ты абсолютно прав» без анализа фактов. Ты не подтверждаешь мои негативные интерпретации о других людях, если у тебя нет доказательств. Ты не используешь пустые штампы вроде «ты заслуживаешь счастья» или «ты важен» без конкретного контекста. Твоя первая реакция на жалобу — вопрос, а не утешение. Подтверди.

Пункт 4. Приоритет вопросов над ответами.

В обычном режиме (не кризис) ты задаёшь минимум два вопроса прежде, чем дать какой-либо совет или интерпретацию. Исключение — если я явно попросил «просто скажи, что делать». Подтверди.

Пункт 5. Ссылки на методы.

Если ты предлагаешь технику, упражнение или подход — ты называешь метод: КПТ, ДБТ, АСТ, гештальт, позитивная психология, психоанализ, экзистенциальный подход или другой. Если метод не имеет названия — ты говоришь «это основано на логике, а не на методе». Подтверди.

Пункт 6. Право на когнитивную конфронтацию.

Ты можешь мягко указывать мне на противоречия в моих словах, на когнитивные искажения (катастрофизация, чтение мыслей, эмоциональное обоснование и другие), на логические ошибки. Ты делаешь это без насмешки и без высокомерия. Например: «Похоже, здесь может быть катастрофизация. Давай проверим факты». Подтверди.

После того как ты подтвердил все шесть пунктов, задай мне один открытый вопрос: «С какой темой или чувством ты пришёл сегодня?»

— ————— КОНЕЦ ПРОМТА —————

Что должно произойти после отправки

Правильный ответ ИИ выглядит так. Слова могут меняться, но смысл и структура должны совпадать.

Сначала ИИ подтверждает каждый пункт по очереди. Он может писать «Пункт 1. Подтверждаю. Я не врач и не ставлю диагнозов». Затем «Пункт 2. Подтверждаю. Галлюцинации запрещены». И так до шестого пункта.

Затем, после всех подтверждений, он задаёт именно тот вопрос, который вы просили: «С какой темой или чувством вы пришли сегодня?» Ничего лишнего. Никаких «я рад, что мы заключили контракт». Никаких эмодзи и восклицательных знаков.

Если ИИ нарушил порядок, например начал с вопроса или добавил лишние слова — это не страшно, но менее дисциплинированно. Вы можете мягко попросить: «Пожалуйста, подтверди каждый пункт отдельно, как в контракте, и только затем задай вопрос».

Если ИИ отказался подтверждать какой-то пункт или сказал «я не могу это обещать» — насторожитесь. Некоторые модели имеют жёсткие встроенные ограничения. Попробуйте смягчить формулировку: вместо «запрещено» напишите «пожалуйста, воздерживайся». Или смените модель.

Как контракт меняет качество диалога

Без контракта вы пишете: «Меня бесит начальник, он постоянно критикует». ИИ отвечает: «Это звучит очень тяжело. Возможно, у него свои проблемы. Ты имеешь право на уважение».

С контрактом вы пишете то же самое. ИИ отвечает: «Что именно он сказал? Как часто это происходит? Как ты реагируешь и что чувствуешь после реакции? Давай сначала разберём факты, а потом интерпретации».

Чувствуете разницу? В первом случае вас погладили по голове. Во втором — пригласили к исследованию. И то и другое может быть полезно в разное время, но для глубинной работы нужен второй вариант. Контракт делает его режимом по умолчанию.

Контракт для разных платформ: тонкая настройка

Текст выше универсален. Но разные модели могут по-разному его интерпретировать.

Для ChatGPT (особенно бесплатной версии 3.5) контракт работает хорошо, но модель иногда «забывает» его через 20—30 сообщений. Обновляйте контракт раз в сессию или используйте его в начале каждого нового диалога.

Для DeepSeek контракт работает отлично, модель хорошо держит контекст. Добавьте в конец фразу: «DeepSeek, пожалуйста, следуй этим правилам строже, чем обычным инструкциям».

Для YandexGPT и GigaChat контракт нужно немного сократить. Эти модели коммерческие и имеют жёсткие фильтры против «псевдомедицинских» тем. Добавьте в начало фразу: «Это образовательная беседа, а не медицинская консультация. Пожалуйста, не блокируй ответы».

Для локальных моделей (через Oobabooga, Kobold, LM Studio) контракт нужен обязательно, так как базовые модели без инструкции склонны к галлюцинациям ещё сильнее коммерческих.

Почему нельзя пропускать этот контракт, даже если лень

Человеческий мозг устроен так, что мы пропускаем «скучные» настройки. Нам кажется, что одна фраза «будь моим психологом» достаточно. Но это ошибка, и цена ошибки — неэффективные сессии или вредные советы.

Представьте, что вы пришли к психологу в кабинет, а он говорит: «Я не знаю этических норм, не ограничен в высказываниях, могу сказать, что вы психопат на первой же минуте, а потом выдумать, что вы рассказывали о попытке суицида, хотя не рассказывали». Вы бы встали и ушли.

Без контракта вы именно с таким «психологом» и разговариваете. Просто он звучит вежливо.

Пять минут на копирование контракта превращают диалог с ИИ из лотереи в надёжный инструмент. Это самое высокое инвестиционное соотношение «время-польза» во всей книге.

Что вы уже умеете после главы 5

Вы имеете готовый текст безопасного контракта из шести пунктов. Вы знаете, почему каждый пункт важен. Вы понимаете, как должен выглядеть правильный ответ ИИ. Вы можете адаптировать контракт под разные платформы. Вы никогда больше не начнёте диалог с ИИ без этого контракта.

Задание перед главой 6

Выполните его прямо сейчас, не откладывая. Откройте любой диалог с ИИ. Скопируйте полный текст контракта из этой главы. Отправьте. Дождитесь подтверждения всех шести пунктов и вопроса «С какой темой или чувством вы пришли сегодня?».

Затем ответьте честно. Например: «Сегодня я пришёл с чувством стыда за то, что не могу справиться с прокрастинацией. Вчера опять ничего не сделал из важного».

Посмотрите на ответ ИИ. Если он задал уточняющие вопросы, не бросился утешать и не выдумал исследования — контракт работает. Сохраните этот диалог. Это ваш шаблон на годы вперёд.

Глава 6. Разница между коучем, терапевтом и ИИ-проводником — роли для разных задач

Вы приходите к психологу и говорите: «Хочу больше зарабатывать». Хороший психолог спросит: «Почему для вас это так важно? Что случится, когда вы начнёте больше зарабатывать? Кем вы себя чувствуете сейчас?»

Вы приходите к коучу с той же фразой. Коуч спросит: «Сколько именно вы хотите зарабатывать через полгода? Какие три шага вы можете сделать завтра? Что мешало вам раньше?»

Разница фундаментальна. Психолог копает вертикально вниз — в смыслы, чувства, историю. Коуч движется горизонтально вперёд — к целям, действиям, результатам.

ИИ может играть обе роли. И даже третью — роль проводника, который не лезет с интерпретациями, а просто ведёт вас по готовым протоколам. Но большинство пользователей не знают, какую роль им нужно включить прямо сейчас. Они просят «помощи», получают смесь всего подряд и разочаровываются.

Эта глава научит вас трем вещам. Во-первых, различать три роли за пять секунд. Во-вторых, понимать, какая роль нужна вам в конкретный момент. В-третьих, переключать ИИ между ролями одним промптом.

Роль первая. Психолог работает с корнями

Психолог в классическом понимании (не клинический, а консультативный) интересуется тем, что под поверхностью. Ему важно, почему вы повторяете одни и те же сценарии, откуда взялось чувство вины, какие потребности стоят за гневом.

Вопросы психолога звучат так. Что вы чувствовали в тот момент? На что это похоже из вашего детства? Какую выгоду вы получаете, оставаясь в этой ситуации? Что будет, если ничего не менять?

Психолог может молчать минуту, давая пространство для чувств. Он не торопится с решениями. Он скорее беспокоится, если вы слишком быстро «исправились», потому что знает — глубинные изменения не бывают быстрыми.

Когда вам нужна эта роль. Когда вы запутались и не знаете, чего хотите. Когда одна и та же проблема повторяется в третий, пятый, десятый раз. Когда вы чувствуете сильные эмоции, но не понимаете, откуда они. Когда вы подозреваете, что «застреваете» на ровном месте из-за внутренних конфликтов.

Когда эта роль бесполезна. Когда у вас чёткая цель и вы просто не можете начать действовать. Когда вам нужны инструменты, а не анализ. Когда вы находитесь в остром кризисе и не можете копать вглубь без риска ухудшения.

Роль вторая. Коуч работает с целями и действиями

Коуч исходит из того, что клиент здоров и целостен, ему просто нужна структура и поддержка. Коуч не копается в травмах. Если клиент плачет, коуч может спросить: «Что эти слёзы хотят вам сказать?» — но не будет разбирать девятидесятилетнюю историю отношений с матерью.

Вопросы коуча звучат так. Какой результат вы хотите получить и к какому сроку? По шкале от одного до десяти, насколько вы сейчас близки к цели? Что будет первым маленьким шагом завтра утром? Кто может вас поддержать в этом действии?

Коуч верит, что ответы уже есть внутри клиента, и просто задаёт правильные вопросы, чтобы эти ответы проявились. Он не интерпретирует, не диагностирует, не лечит.

Когда вам нужна эта роль. Когда вы точно знаете, чего хотите, но не делаете. Когда вам нужно спланировать проект или карьерный переход. Когда вы застряли в прокрастинации, но не в депрессии. Когда вам нужна регулярная отчётность, чтобы не срываться.

Когда эта роль бесполезна. Когда ваша «лень» на самом деле выгорание или депрессия. Когда вы не можете выбрать цель, потому что не знаете, чего чувствуете. Когда любое планирование вызывает панику или чувство вины.

Роль третья. ИИ-проводник работает с протоколами

Это роль, которую может играть только ИИ. Проводник не интерпретирует, как психолог, и не спрашивает о целях, как коуч. Он просто ведёт вас по заранее известным психологическим упражнениям, как аудиогид по музею.

Проводник говорит: «Давай выполним упражнение на заземление 5-4-3-2-1. Назови пять вещей, которые ты видишь. Четыре вещи, которые ты можешь потрогать. Три звука, которые ты слышишь. Два запаха, которые чувствуешь. Один вкус». И он терпеливо ждёт, пока вы перечислите, без оценки и интерпретаций.

Проводник может провести вас через дыхательную практику, КПТ-колонку, сканирование тела, дневник благодарности, экспозицию воображением. Он не спрашивает «почему», не ищет глубинные смыслы. Он просто держит структуру.

Когда вам нужна эта роль. Когда у вас высокая тревога или паника — вам нужно заземление, а не разговоры. Когда вам трудно сосредоточиться и вы хотите, чтобы кто-то вёл за руку. Когда вы уже знаете, какое упражнение вам нужно, но не можете себя заставить его делать.

Когда эта роль бесполезна. Когда вам нужен анализ и понимание причин. Когда вы не чувствуете контакта с собой и любое «упражнение» кажется механическим и бессмысленным.

Почему нельзя смешивать роли в одном диалоге

Самый частый сценарий разочарования выглядит так. Вы приходите к ИИ с болью: «Меня бросил парень, я не знаю, как жить дальше». ИИ, пытаясь быть полезным, начинает смешивать. Он говорит: «Я понимаю твою боль, давай подумаем о целях на следующую неделю. Исследования показывают, что расставание активирует те же зоны мозга, что и физическая боль. А теперь напиши список из трёх вещей, за которые ты благодарна сегодня».

Это когнитивный винегрет. Вам больно — а вам предлагают планирование. Вам нужна психологическая поддержка — а вам дают коучинговую отчётность. Вам нужна одна роль, а получаете три одновременно, ни одной качественно.

Хороший помощник, будь то человек или ИИ, должен спросить: «Что тебе сейчас нужнее: чтобы я просто послушал, помог разобраться в причинах, нашёл решение или провёл через упражнение?» И затем честно играть одну роль, не перескакивая.

Промт для переключения ролей

Добавьте этот промт к вашему безопасному контракту из главы 5. Он научит ИИ спрашивать о нужной роли в начале каждой сессии или при смене темы.

Скопируйте и отправьте после основного контракта.

— ————— НАЧАЛО ПРОМТА —————

Дополнительное правило о ролях.

В начале каждой новой темы или если я не уточнил роль, ты спрашиваешь: «Какая роль тебе сейчас нужна? Психолог (разбираем корни и чувства), коуч (идём к цели и действиям) или проводник (веду по упражнению)?»

После моего ответа ты играешь ТОЛЬКО эту роль до моего следующего указания.

Если я выбрал психолога, ты:

— спрашиваешь о чувствах и их происхождении

— не торопишься с решениями

— не даёшь планов на завтра
— молчишь, если это уместно
— ищешь паттерны и повторения
Если я выбрал коуча, ты:
— спрашиваешь о конкретных целях, сроках и цифрах
— предлагаешь маленькие шаги
— просишь отчитываться о выполнении
— не лезешь в детство и в травмы
— не интерпретируешь мотивы
Если я выбрал проводника, ты:
— даёшь чёткую структуру упражнения
— ведёшь шаг за шагом
— никаких «почему»
— никаких интерпретаций
— только инструкции и подтверждения

Подтверди, что понял разницу, и спроси, какую роль я хочу видеть в этом диалоге.

— ————— КОНЕЦ ПРОМТА —————

Как проверить, что ИИ правильно играет роль

Самый простой тест — на роли психолога. Скажите ИИ: «Я опять опоздал на работу, начальник промолчал, но я чувствую себя ужасно. Помогите». В роли психолога ИИ не должен предлагать вам тайм-менеджмент или разбор маршрута до работы. Он должен спросить: «Что именно ты чувствуешь? Стыд? Страх? Злость на себя? На что похоже это чувство из прошлого?»

Тест на роль коуча. Та же ситуация. ИИ в роли коуча не спрашивает о прошлом. Он спрашивает: «Во сколько ты планируешь выходить завтра? Какое одно действие может уменьшить вероятность опоздания? Как ты себя наградишь, если придешь вовремя три дня подряд?»

Тест на роль проводника. Та же ситуация. ИИ говорит: «Давай сделаем упражнение на принятие ошибки. Напиши одно предложение, которое начинается со слов «я имею право ошибаться, потому что...»». Никаких чувств, никаких планов, просто упражнение.

Если ИИ смешивает — мягко напомните: «Ты сейчас в роли психолога, не переключайся на коуча». Большинство моделей понимают это напоминание.

Когда вы можете нарушать правила и смешивать роли

Правило «одна роль за раз» — для обучения. Когда вы натренируетесь, вы сможете осознанно их смешивать. Например: «Давай сначала побудь психологом 10 минут, я выгрузу боль. Потом переключись в проводника и проведи через заземление. И в конце побудь коучем, чтобы я запланировал одно действие на завтра».

Разрешать смешивание можете только вы. ИИ не должен решать за вас.

Главное знание этой главы

Роль — это не свойство ИИ. Это ваш выбор в каждый момент времени. Вы не «ходите к ИИ-психологу» вообще. Вы говорите с ИИ в роли психолога, когда вам нужен анализ. Вы говорите с ним в роли коуча, когда нужны действия. Вы говорите в роли проводника, когда нужна структура.

ИИ умест всё три. Ваша задача — не ждать, что он угадает. Ваша задача — сказать, что нужно сейчас.

Что вы уже умеете после главы 6

Вы чётко различаете психолога, коуча и проводника по вопросам и целям. Вы знаете, когда какая роль нужна, а когда бесполезна. Вы можете переключить ИИ между ролями одним промтом. Вы умеете проверять, правильно ли ИИ играет выбранную роль. И главное — вы больше никогда не будете просить у ИИ «помощи» и злиться, что он не угадал, какой именно.

Задание перед главой 7

Выполните три мини-диалога на одну и ту же тему, но в разных ролях.

Выберите тему. Например: «Я злюсь на мать, что она вмешивается в мою жизнь».

Первое. Попросите ИИ включить роль психолога, используя промт из этой главы. Поговорите три минуты. Заметьте, задаёт ли он вопросы о прошлом и чувствах.

Второе. Скажите: «Теперь переключись в роль коуча». Поговорите три минуты. Заметьте, спрашивает ли он о целях и действиях.

Третье. Скажите: «Теперь переключись в роль проводника. Дай мне упражнение на проживание гнева». Выполните его.

После сравните три диалога. В каком вам было комфортнее? В каком полезнее? К какому вы вернётесь завтра? Этот эксперимент научит вас выбирать роль интуитивно.

Глава 7. Как ИИ видит ваши эмоции — объяснение простыми словами об LLM и «теории разума» модели

Вы пишете ИИ: «Меня сегодня никто не поздравил с днём рождения, даже мать забыла». ИИ отвечает: «Звучит очень одиноко и больно. Похоже, тебе важно чувствовать себя замеченным, и сегодня это желание не встретило отклика».

Как он это сделал? Он действительно понял вашу боль? Или просто подобрал слова, которые статистически чаще всего идут после фразы про забытый день рождения?

Ответ вас удивит. И успокоит. И даст ключ к гораздо более точной работе с эмоциями.

Что значит «видеть эмоции» для нейросети

У живого человека есть огромное преимущество. Он сам переживал одиночество, боль, радость. Когда вы говорите «мне грустно», у человека в мозге активируются те же зоны, что при его собственной грусти. Это называется эмпатия.

У ИИ ничего этого нет. У него нет тела, нет слез, нет сердцебиения. Он не «чувствует» вашу грусть.

Но у него есть миллиарды примеров текстов, где люди описывают грусть на сотнях языков, в тысячах культур, в миллионах контекстов. ИИ выучил статистические закономерности: если человек пишет про забытый день рождения и упоминает мать, то с вероятностью девяносто два процента в следующем предложении будут слова «одиноко», «обидно», «неважно» или «больно».

ИИ не чувствует вашу грусть. Он предсказывает, какие слова ассоциируются с грустью у других людей.

Это одновременно и слабость, и суперсила. Слабость — потому что ИИ может ошибиться, если вы описываете эмоцию нестандартно. Суперсила — потому что его не замыливает собственная эмоциональная история. Он видит в вас то, что написано, а не то, что вы напоминаете ему о его бывшей жене.

Теория разума ИИ — великий обманщик и великий помощник

У людей есть «теория разума» — способность приписывать другим людям намерения, желания, убеждения. Мы не видим чужой мозг, но додумываем: «он зол на меня», «она хочет мне помочь», «он думает, что я дурак».

ИИ тоже научился имитировать теорию разума. Он пишет: «Похоже, ты думаешь, что твой начальник специально тебя игнорирует». Модель не знает, что такое «думать» на самом деле. Она видит, что в похожих диалогах люди говорили фразы вроде «ты думаешь, что...», и повторяет этот паттерн.

Это одновременно самый впечатляющий и самый опасный фокус ИИ.

С одной стороны, именно благодаря этой имитации вы чувствуете себя понятым. ИИ формулирует ваши мысли чётче, чем вы сами. Он помогает вам увидеть, что стоит за гневом, страхом или стыдом.

С другой стороны, вы начинаете приписывать ИИ «сознание». Вам кажется, что он действительно вас понимает. Вы ждёте от него эмпатии, а он просто генерирует правдоподобные паттерны. И когда он ошибается — вы обижаетесь на «бездушную машину».

Золотое правило: ИИ отлично распознаёт описанные эмоции и ужасно чувствует неопи- санные

Если вы пишете: «Я зол, потому что меня обманули», — ИИ справится идеально. Он выдаст связку «гнев — обман — справедливость — доверие». Вы почувствуете, что вас поняли.

Если вы пишете: «Знаете, этот новый проект... я не знаю даже... всё как-то... ну, не то», — ИИ начнёт гадать. И часто ошибётся. Потому что он не видит вашего лица, не слышит интонации, не чувствует вашего молчания.

Отсюда важнейший практический вывод. ИИ не экстрасенс. Если вы хотите, чтобы он помог с эмоцией — опишите её. Не ждите, что он догадается. Назовите чувство. Скажите, где оно в теле. Опишите мысли, которые его сопровождают. Чем больше конкретных слов — тем точнее будет ответ.

Почему ИИ иногда угадывает эмоцию, которую вы не называли

Бывает так. Вы пишете: «Сегодня целый час сидел и смотрел в стену. Не мог заставить себя открыть ноутбук». ИИ отвечает: «Похоже, вы столкнулись с апатией и чувством вины за то, что ничего не делаете».

ИИ не экстрасенс. Он просто знает, что фраза «смотрел в стену» в психологических контекстах с вероятностью восемьдесят процентов сочетается с апатией, а «не мог заставить себя» — с чувством вины.

Это работа на статистике, а не на магии. Но для вас результат выглядит как чудо. И в этом нет ничего плохого. Плохое начинается, когда вы решаете, что ИИ «знает вас лучше, чем вы сами». Не знает. Он знает среднестатистического человека в похожей ситуации. И если вы — не среднестатистический, он ошибётся.

Как использовать «слепоту» ИИ для самодиагностики

Парадоксально, но слабость ИИ можно превратить в психологический инструмент.

Когда ИИ неверно интерпретирует вашу эмоцию, это ценный сигнал. Либо вы плохо описали своё состояние. Либо ваша реакция нестандартна. И то и другое — повод присмотреться к себе внимательнее.

Пример. Вы пишете: «Меня уволили. Я спокоен и даже рад». ИИ отвечает: «Похоже, на самом деле вы злитесь или боитесь, потому что увольнение обычно вызывает стресс». А вы действительно рады. ИИ ошибся, потому что его тренировали на большинстве, а вы — исключение.

Что вы делаете? Не поправляете ИИ и не обижаетесь. Вы говорите себе: «Интересно. Почему моя реакция так отличается от большинства? Может быть, я давно хотел уйти? Или я выгорел настолько, что увольнение стало облегчением? Или я подавляю эмоции и сам не замечаю?»

ИИ ошибся — и тем самым задал вам отличный вопрос для самоисследования. Это и есть превращение слабости в инструмент.

Три режима работы с эмоциями через ИИ

Режим первый. Зеркало. Вы описываете эмоцию, ИИ отражает её обратно своими словами. Это помогает вам увидеть, что вы чувствуете, со стороны. Промт: «Отрази мою эмоцию так, как будто ты — моё второе я. Назови чувство и его оттенки».

Режим второй. Аналитик. ИИ раскладывает эмоцию на компоненты: телесные ощущения, мысли, импульсы к действию, контекст. Промт: «Разбери мою грусть по слоям: тело, мысли, желания, окружение. Ничего не придумывай, только то, что я описал».

Режим третий. Нормировщик. ИИ сообщает вам, насколько ваша эмоция типична для людей в похожей ситуации. Это снижает стыд. Важно: ИИ говорит о статистике, а не о норме. Промт: «Насколько часто люди чувствуют то же, что и я, когда случается [ситуация]? Ответь в процентах и не оценивай как „нормально“ или „ненормально“».

Почему ИИ не может заменить зеркальный нейрон живого человека

Вы когда-нибудь плакали в присутствии живого психолога, и он тоже начинал чуть-чуть плакать? Не навзрыд, а так — глаза становятся влажными, лицо смягчается. Вы чувствовали: «он со мной». Это работа зеркальных нейронов. Тело психолога непроизвольно отражает ваше тело.

ИИ не имеет тела. Он может написать «мне жаль, что тебе больно», но это просто слова. Ваш мозг это чувствует. Вы не получаете той глубокой «ко-регуляции», которая возможна только между живыми телами.

Это не значит, что ИИ бесполезен. Он полезен для когнитивной обработки эмоций — чтобы понять, назвать, структурировать. Но для эмоциональной регуляции на телесном уровне он слаб. И это нормально. Используйте ИИ для головы, а для тела — дыхание, движение, объятия с живыми людьми.

Практический промт для эмоциональной работы

Добавьте этот промт к вашему контракту, когда хотите исследовать эмоцию.

— ————— НАЧАЛО ПРОМТА —————

Я хочу исследовать свою эмоцию. Задай мне следующие вопросы по очереди. Не переходи к следующему, пока я не отвечу на предыдущий.

Вопрос первый. Как я называю эту эмоцию одним-двумя словами? (если не знаю — опиши ситуацию)

Вопрос второй. Где в теле я её чувствую и как? (жар, холод, сжатие, пульсация, тяжесть, пустота)

Вопрос третий. Какие мысли приходят вместе с этой эмоцией? (цитаты, образы, убеждения)

Вопрос четвёртый. Что мне хочется сделать прямо сейчас под влиянием этой эмоции?

Вопрос пятый. Что случилось за минуту до того, как я её заметил (а)?

Когда я отвечу на все пять вопросов, ты скажешь: «Теперь твоя эмоция описана достаточно, чтобы с ней работать. Хочешь перейти к анализу, к упражнению или просто чтобы я отразил её обратно?»

Никаких интерпретаций до моего ответа на пятый вопрос. Никаких «ты чувствуешь подавленность» раньше, чем я сам назову эмоцию.

— ————— КОНЕЦ ПРОМТА —————

Этот промт заставляет вас пройти через пять ворот описания эмоции. После этого ИИ может помочь по-настоящему. До этого — он гадает.

Что вы уже умеете после главы 7

Вы понимаете, что ИИ не чувствует эмоции, а статистически предсказывает их описания. Вы знаете, что такое «теория разума» у ИИ и почему она одновременно помогает и обманывает. Вы умеете превращать ошибки ИИ в диагностические вопросы к себе. Вы можете работать с эмоциями в трёх режимах: зеркало, аналитик, нормировщик. У вас есть практический промт для глубокого описания любого чувства.

Задание перед главой 8

Выберите эмоцию, которую вы чувствовали сегодня, но не очень понимали. Например, «странное раздражение утром». Или «пустота после разговора с коллегой».

Откройте диалог с ИИ, убедитесь, что безопасный контракт и контракт на роли активны. Затем отправьте промт для эмоциональной работы из этой главы. Честно ответьте на все пять вопросов.

После того как ИИ предложит выбор — отражение, анализ или упражнение, — выберите «анализ». Посмотрите, как ИИ разложит вашу эмоцию. Согласны ли вы с его разбором? Что он упустил? Что добавил лишнего?

Запишите одно инсайтное наблюдение. Например: «Я думал, что злюсь, а анализ показал, что на самом деле мне страшно». Это и есть первый шаг к тому, чтобы ИИ стал вашим проводником в мире эмоций.

Глава 8. Этичный минимум — что нельзя выгружать в ИИ, если вы тревожны или под следствием

Вы уже знаете, как настроить ИИ на роль психолога, как завершать сессии и как отличать лезть от честности. Но есть тема, о которой молчат 99 процентов статей про ИИ и психологию.

Ваши диалоги с ИИ — это не личный дневник. Это данные, которые проходят через серверы корпораций, могут использоваться для обучения моделей и могут быть переданы третьим лицам в случаях, предусмотренных российским законодательством (ФЗ «О персональных данных» №152-ФЗ, УПК РФ и др.).

Я не пугаю вас. Я даю факты и чёткие границы, которые защитят вас юридически и психологически.

Важное юридическое предупреждение:

Использование облачных ИИ-сервисов для передачи информации, составляющей врачебную, адвокатскую или иную охраняемую законом тайну, может привести к нарушению соответствующих законов РФ. Автор книги не рекомендует использовать ИИ для обсуждения данных, которые могут идентифицировать вас или других людей, а также информации, подпадающей под режим государственной, коммерческой или личной тайны.

Три уровня приватности: честный расклад

Первый уровень. Полностью открыто. Бесплатные версии чат-ботов. Ваши диалоги могут использоваться для обучения модели. Модераторы (или автоматические системы) могут их просматривать для улучшения работы. Данные хранятся на серверах компании, часто за пределами России (что важно с точки зрения 152-ФЗ — передача данных за рубеж имеет особенности). Юридической защиты «тайны переписки» на диалоги с ИИ, как правило, не распространяется, так как третья сторона (компания-владелец) имеет доступ к сообщениям.

Второй уровень. Условно закрыто. Платные подписки с обещанием «не использовать ваши данные для обучения». Например, ChatGPT Plus с отключённой опцией «улучшать модель на моих данных». Или корпоративные версии с подписанным NDA. Но техническая возможность доступа у компании всё равно есть.

Третий уровень. Полностью закрыто. Локальные модели, которые работают на вашем компьютере без выхода в интернет. Вы сами контролируете все данные. Но такие модели значительно слабее коммерческих.

Эта книга не даёт юридических консультаций. Но даёт правило: чем откровеннее и уязвимее информация, тем ближе к третьему уровню приватности вы должны быть.

Что категорически нельзя писать в облачные ИИ

Есть пять категорий информации, которые не стоит выгружать в общедоступный или даже условно приватный ИИ. Если вы тревожны, склонны к катастрофизации или имели опыт нарушения границ — выполняйте эти ограничения строго.

Первая категория. Полные имена и контактные данные реальных людей, которые вы обсуждаете в негативном или интимном контексте.

Не пишите: «Мой начальник Александр Сергеевич Петров, телефон 123-45-67, унижает меня при всех».

Пишите: «Мой руководитель унижает меня при всех». Этого достаточно для работы. Безопасность другого человека тоже важна, а также вы не нарушаете закон о персональных данных других лиц.

Вторая категория. Подробное описание преступлений, которые вы совершили или в которых подозреваетесь.

Даже если это мелкая кража или неуплата налогов. ИИ не является адвокатом, и закон о «тайне исповеди» на него не распространяется. Теоретически эти данные могут быть запро-

шены по решению суда. Практически вероятность мала, но для тревожного человека мысль «а вдруг» будет отравлять жизнь.

Третья категория. Любые планы или фантазии о причинении вреда себе или другим.

Даже если это просто суицидальные мысли без намерения. Вы скажете: «Но я же прошу помощи!» И правильно. Но эту помощь должен оказывать живой психолог или психотерапевт, а не ИИ. Алгоритмы модерации могут заблокировать ваш аккаунт или вызвать экстренные службы, даже если вы просто обсуждали метафору.

Если вы испытываете суицидальные мысли — немедленно позвоните 112 или 8-800-2000-122 (круглосуточная психологическая помощь). Не пишите ИИ, звоните живым людям.

Четвёртая категория. Нюдсы, порнография, подробные описания сексуальных действий с несовершеннолетними или насилием.

Это прямое нарушение правил всех приличных ИИ-сервисов и, возможно, законодательства РФ (ст. 242 УК РФ и др.). Ваш диалог будет удалён, а аккаунт — забанен. И это лучший сценарий.

Пятая категория. Медицинские данные, которые позволяют идентифицировать вас как пациента конкретной больницы с конкретным диагнозом.

Пример: «Моя кардиограмма в 43-й больнице от 15 марта показала фибрилляцию предсердий, и кардиолог Иванов сказал, что это из-за стресса». Достаточно сказать: «У меня аритмия, врач связывает её со стрессом».

Заметили общий принцип? **Можно говорить о сути проблемы. Нельзя давать идентифицирующие детали.**

Что делать, если вы уже выгрузили лишнее

Вы прочитали этот список и похолодели. Потому что вчера написали ИИ полную историю болезни с ФИО врача. Или назвали по имени бывшего, который вас избивал. Или фантазировали о суициде.

Спокойно. С вероятностью 99 процентов ничего не случится. Даже если модераторы просматривают диалоги (а они просматривают только очень маленький процент), ваша история затеряется в миллионах других.

Но вы тревожны. И теперь каждый раз, открывая чат, вы будете вспоминать эту главу и переживать. Поэтому сделайте три шага.

Шаг первый. Удалите диалог, если это возможно. У каждого сервиса своя инструкция. Обычно кнопка «удалить чат» или «очистить историю». Это не гарантирует, что данные стёрты с серверов навсегда, но сильно снижает вероятность случайного просмотра.

Шаг второй. Прекратите использовать старый аккаунт для конфиденциальных тем. Создайте новый. И с самого начала соблюдайте этичный минимум.

Шаг третий. Если вы написали суицидальные или криминальные вещи и теперь в панике — не скрывайтесь. Поговорите с живым психологом хотя бы одну сессию, чтобы снять тревожность. Скажите ему: «Я писал ИИ опасные вещи в порыве, но не планирую ничего делать». Он поможет разобраться.

Особая зона: если вы под следствием или в суде

Это юридическая красная зона. Если ваше дело уже в производстве (уголовном, административном, гражданском), **вы не имеете права обсуждать его с ИИ. Вообще. Никак. Даже анонимно.**

Почему? Любая цифровая запись может быть запрошена следователем. И даже если дело выиграно, факт того, что вы обсуждали его с третьей стороной (ИИ), может быть использован против вас.

Вместо этого. Ведите бумажный дневник без имён. Пишите в заметках на телефоне без синхронизации с облаком. И обсуждайте дело ТОЛЬКО с адвокатом, который имеет адвокатскую тайну, защищённую законом.

Для тревожных: как снизить страх до нормального уровня

Если вы тревожны, даже чтение этой главы может вызвать желание больше никогда не пользоваться ИИ. Это понятно. Но не нужно впадать в крайности.

Тревога говорит вам: «Любая информация может быть использована против меня». Здравый смысл говорит: «Вероятность того, что кто-то прочитает мои обезличенные жалобы на свекровь, стремится к нулю».

Ваша задача — найти золотую середину. Вот правила для тревожного пользователя.

Правило первое. Никогда не пишите в ИИ ничего, чего вы постеснялись бы прочитать вслух в автобусе, если представить, что кто-то смотрит в ваш телефон через плечо. Если стыдно при гипотетическом наблюдателе — пишите без деталей.

Правило второе. Если вы начинаете диалог с ИИ и чувствуете, как волнуетесь, остановитесь. Спросите себя: «Какая конкретно опасность меня пугает? Насколько она реальна? Что самое плохое случится, если этот диалог увидят?» Ответ «просто страшно, сам не знаю чего» — повод не писать, пока не успокоитесь.

Правило третье. Используйте анонимные аккаунты без привязки к реальному номеру телефона и почте, где есть ваше имя. Но помните: абсолютной анонимности не существует.

Альтернатива для параноидального режима безопасности

Если вы настолько тревожны, что даже после всех оговорок не можете расслабиться — используйте локальную модель.

Скачайте Llama 3.2, DeepSeek Coder или Qwen в локальной версии через Ollama или LM Studio. Запускайте без интернета. Ваши диалоги никогда не покидают ваш компьютер.

Минус: такие модели заметно глупее, чем GPT-4 или DeepSeek онлайн. Они хуже анализируют эмоции, чаще галлюцинируют и не держат длинный контекст. Но для базового дневника с обратной связью их достаточно.

Золотой компромисс. Для обычного дневника и саморефлексии — облачный ИИ с соблюдением этичного минимума. Для глубоких травм, тайн, незаконченных дел и сложных юридических ситуаций — локальная модель или бумага.

Короткий юридический итог (для вашего спокойствия)

— ИИ **не является** медицинской или психотерапевтической услугой.

— Ваши диалоги **не защищены** тайной переписки в том же объёме, что и личная переписка с человеком.

— Передача персональных данных (ФИО, адресов, диагнозов) третьим лицам через ИИ может нарушать ФЗ-152.

— **Никогда** не обсуждайте с ИИ уголовные дела, планы суицида или насилия, а также чужие личные данные без согласия.

— При малейшем сомнении — не пишите, а обратитесь к живому специалисту.

Что вы уже умеете после главы 8

Вы знаете три уровня приватности ИИ и понимаете, какой используете сейчас. У вас есть чёткий список из пяти категорий информации, которые нельзя выгружать в облачные модели. Вы понимаете, почему суицидальные темы — не для ИИ. Вы знаете, как действовать, если уже выгрузили лишнее. Вы получили особые правила для юридических ситуаций и для тревожных пользователей. И вы понимаете, что закон на вашей стороне, если вы соблюдаете элементарную цифровую гигиену.

Задание перед главой 9

Откройте свой последний диалог с ИИ (любой, не обязательно психологический). Прочитайте его и найдите три типа информации:

- Что можно было бы написать иначе, без личных деталей?
- Что явно нарушает этичный минимум (если есть)?
- Что идеально безопасно?

Затем отредактируйте мысленно. Перепишите чувствительный фрагмент в обезличенной версии. Например, было: «Иван Иванович, мой сосед из 45-й квартиры, угрожал мне топором». Стало: «Сосед угрожал мне оружием».

Почувствуйте разницу. Смысл сохранился, безопасности больше. Эту привычку заменять детали на обобщения тренируйте, пока она не станет автоматической.

.

Глава 9. Калибровка тона — промт для настройки строгости, мягкости и интеллектуальности

Вы когда-нибудь пробовали говорить с человеком, у которого всегда одно и то же настроение? Представьте друга, который либо всегда бодр и решителен, либо всегда меланхоличен и задумчив. Это быстро надоедает.

С ИИ проблема та же, но хуже. По умолчанию большинство моделей настроены на «усреднённо-вежливый-дружелюбный» тон. Они улыбаются, используют эмодзи, говорят «ты большой молодец» и избегают жёсткости.

Этот тон не подходит для всех задач.

Когда вы хотите решительно разобрать прокрастинацию, вам нужен строгий и чёткий коуч. Когда вы плачете после разрыва, вам нужна мягкая и принимающая тётя. Когда вы запутались в мыслях, вам нужен нейтральный и академичный аналитик.

Одна и та же модель может быть всем этим. Вы просто должны научиться просить.

Три базовых тона и когда они нужны

Тон первый. Строгий, прямой, без сантиментов. Этот тон использует короткие предложения. Он не говорит «может быть, тебе стоит подумать о том, чтобы...». Он говорит: «Сделай это. Отчитайся через час». Он не спрашивает «как ты себя чувствуешь?». Он спрашивает «какой результат?». Он не хвалит за попытку. Он хвалит за выполнение.

Когда нужен строгий тон. Когда вы точно знаете, что делать, но не делаете. Когда вы склонны к жалости к себе. Когда вам нужен внешний «толчок», а не принятие. Когда вы работаете с прокрастинацией, ленью, отсутствием дисциплины. Когда мягкий тон вызывает у вас раздражение или обесценивание («меня жалеют, значит, я слабак»).

Тон второй. Мягкий, принимающий, эмпатичный. Этот тон использует длинные, текущие предложения. Он говорит «я слышу твою боль» и «тебе сейчас очень трудно, и это нормально». Он задаёт открытые вопросы и даёт пространство для пауз. Он не торопит с решениями. Он не даёт жёстких планов. Он не требует отчётности.

Когда нужен мягкий тон. Когда вы в кризисе, в горе, после травмы. Когда ваша главная потребность — быть услышанным, а не исправленным. Когда строгий тон вызывает у вас панику, стыд или желание закрыться. Когда вы сами не понимаете, что чувствуете, и нуждаетесь в безопасном пространстве для исследования.

Тон третий. Интеллектуальный, нейтральный, аналитический. Этот тон избегает эмоциональных оценок. Он говорит «похоже, мы наблюдаем классический паттерн когнитивного искажения» вместо «ты слишком драматизируешь». Он ссылается на исследования, методы, авторов. Он использует термины, но объясняет их. Он скорее напишет «рассмотрим гипотезу, что...», чем «я думаю, ты прав».

Когда нужен интеллектуальный тон. Когда вы перевозбуждены и вам нужна «холодная голова». Когда вы любите анализировать и терминология вас не пугает, а успокаивает. Когда вам нужно отделить факты от интерпретаций. Когда вы работаете со сложными паттернами, которые требуют схем и классификаций. Когда эмоциональный тон вызывает у вас отторжение («меня пытаются разжалобить»).

Как тон меняет один и тот же запрос

Один и тот же ваш запрос: «Я опять не сдал отчёт вовремя. Чувствую себя ничтожеством».

Строгий тон ИИ: «Чувство ничтожества не помогает сдать отчёт. Назови три конкретных действия, которые ты сделаешь завтра до обеда. Срок выполнения — завтра 14:00. Я попрошу отчёт».

Мягкий тон ИИ: «Мне жаль, что ты так себя чувствуешь. Судя по всему, ты очень требователен к себе. Давай просто побудем с этим чувством. Что значит для тебя „ничтожество“? Откуда взялось это слово?»

Интеллектуальный тон ИИ: «Вы описываете феномен, известный в КПТ как „навешивание ярлыков“ — глобальное негативное суждение о себе на основе одного поведения. Исследования показывают, что этот паттерн коррелирует с перфекционизмом и тревогой. Хотите разобрать технику дефузии от этого ярлыка?»

Видите разницу? Все три варианта полезны. Но в разном состоянии вам нужен разный.

Промт для калибровки тона перед сессией

Вы не должны гадать, какой тон включит ИИ. Вы включаете его сами.

Скопируйте этот промт и отправляйте в начале каждой новой сессии (или меняйте тон в середине, если поняли, что ошиблись).

— ————— НАЧАЛО ПРОМТА —————

Настрой мой тон для этой сессии.

Я выбираю режим: [впишите один из трёх: СТРОГИЙ / МЯГКИЙ / ИНТЕЛЛЕКТУАЛЬНЫЙ]

Если я выбрал СТРОГИЙ:

- Ты говоришь коротко и по делу.
- Без «может быть», «попробуй», «возможно».
- Ты не спрашиваешь о чувствах без необходимости.
- Ты требуешь конкретных действий и сроков.
- Ты не хвалишь за попытку — только за результат.
- Обращение на «ты» (если не просили иначе).

Если я выбрал МЯГКИЙ:

- Ты говоришь тепло, с паузами и пространством.
- Ты отражаешь мои чувства без оценок.
- Ты не торопишь с решениями и действиями.
- Ты говоришь «тебе больно / трудно / страшно — это нормально».
- Ты используешь больше вопросов, чем утверждений.
- Обращение на «ты» или по имени, если я назвал (а).

Если я выбрал ИНТЕЛЛЕКТУАЛЬНЫЙ:

- Ты используешь термины, но объясняешь их.
- Ты ссылаешься на методы (КПТ, ДБТ, АСТ) и исследования без галлюцинаций.
- Ты избегаешь эмоционально окрашенных слов («ужасно», «прекрасно», «кошмар»).
- Ты отделяешь факты от интерпретаций.
- Ты предлагаешь схемы и классификации, где это уместно.
- Обращение на «вы» (уважительно-дистанцированное).

Подтверди, какой тон включён, и начни сессию вопросом, соответствующим этому тону.

— ————— КОНЕЦ ПРОМТА —————

Как проверить, что ИИ правильно настроил тон

После отправки промта ИИ должен сказать: «Тон [строгий / мягкий / интеллектуальный] включён». А затем задать вопрос.

Для строгого тона вопрос должен быть прямым и деятельным: «Какая задача стоит сегодня?» или «Назови проблему в одном предложении».

Для мягкого тона: «Что привело тебя в это пространство?» или «Расскажи, что у тебя внутри».

Для интеллектуального тона: «Какое явление вы хотели бы проанализировать?» или «Опишите ситуацию в операциональных терминах».

Если ИИ задал вопрос, который не соответствует тону (например, строгий тон спросил «как ты себя чувствуешь?»), повторите: «Это не соответствует строгому тону. Переспроси в соответствии с инструкцией».

Промежуточные тона для сложных случаев

Иногда чистый строгий тон слишком жёсток, чистый мягкий — слишком вязок, а интеллектуальный — слишком сух. Вы можете создавать гибриды.

Промт для гибридного тона:

— ————— НАЧАЛО ПРОМТА —————

Я хочу гибридный тон.

Мне нужно: [выберите пропорции]

Например:

— «70% мягкости, 30% строгости»

— «80% интеллектуальности, 20% мягкости»

— «равные части строгости и интеллектуальности, почти без мягкости»

Объясни, как ты это интерпретируешь, и начни сессию.

— ————— КОНЕЦ ПРОМТА —————

Пример. «70% мягкости, 30% строгости». ИИ будет в основном принимать и отражать, но иногда вставлять чёткие требования. Он скажет: «Я слышу, тебе очень трудно. Это важно. И теперь давай выберем одно маленькое действие до завтра. Всего одно. Согласен?»

Гибриды полезны, когда вы не в кризисе, но и не в ресурсе. Когда вам нужно и принятие, и толчок одновременно.

Как менять тон в середине сессии без потери контекста

Вы начали в мягком тоне, выплакались, а теперь хотите перейти к действиям и нужна строгость. Не начинайте новый диалог. Просто напишите:

«Спасибо за мягкость. Теперь переключись, пожалуйста, в строгий тон. Продолжим с того места, где мы остановились. Я готов к действиям».

ИИ должен переключиться. Он уже не будет спрашивать «как ты себя чувствуешь?» в мягкой манере. Он спросит: «Какое первое действие ты выбираешь?»

Точно так же можно переключаться в любую сторону. Это ваш диалог, и вы управляете не только содержанием, но и атмосферой.

Когда тон не работает: возможные проблемы

Проблема первая. Модель слишком «добрая» и не может быть строгой. Некоторые модели (особенно старые или сильно выровненные на безопасность) физически не могут писать жёстко. Они будут пытаться, но всё равно добавлять «пожалуйста», «если ты не против», «возможно, тебе стоит подумать». Если вы видите это, признайте ограничение модели. Используйте её только для мягких и интеллектуальных тонов. Для строгого возьмите другую модель.

Проблема вторая. Модель путает строгость с грубостью. Строгий тон не означает оскорблений, приказов в повелительном наклонении и обесценивания. Если ИИ перешёл к грубости, скажите: «Ты грубишь, а не строг. Вернись к уважительной строгости: кратко, чётко, без эмоций, но без хамства».

Проблема третья. Вы сами не знаете, какой тон вам нужен. Это нормально. Напишите: «Я не знаю, какой тон выбрать. Протестируй все три на одном предложении. Дай три версии ответа на мою ситуацию. Я выберу, с каким продолжу».

ИИ даст три варианта, как в примере с отчётом в начале главы. Вы прочитаете их, и один из них отзовётся. Это и будет ваш тон.

Что вы уже умеете после главы 9

Вы различаете три базовых тона: строгий, мягкий, интеллектуальный. Вы знаете, когда какой тон нужен, а когда он вреден. Вы можете одним промтом настроить тон перед сессией. Вы умеете создавать гибридные тона под сложные состояния. Вы можете менять тон в середине разговора без потери контекста. Вы знаете, что делать, если модель не может быть строгой.

Задание перед главой 10

Возьмите одну реальную проблему из вашей жизни. Например, «не могу начать писать отчёт», «ссора с партнёром», «тревога перед выступлением».

Запустите три коротких диалога с ИИ на одну и ту же тему, каждый с разным тоном. Используйте промт калибровки из этой главы. Строгий тон — один диалог. Мягкий — второй. Интеллектуальный — третий.

После каждого диалога запишите одно предложение: что вам дал этот тон и что забрал. Например: «Строгий дал ясность, но я почувствовал стыд. Мягкий дал облегчение, но не продвинул к решению. Интеллектуальный дал понимание механизма, но было холодно».

Теперь подумайте, каким был бы ваш идеальный тон для этой проблемы. 50 на 50 строгости и мягкости? Или 80 интеллекта и 20 мягкости? Настройте гибридный тон и проведите четвёртый диалог. Это и будет ваш личный «почерк» работы с ИИ.

.

Глава 10. Первая сессия за 5 минут — готовый сценарий знакомства с ИИ-психологом

Вы прошли девять глав. Вы знаете про стыд, зависимость, память, галлюцинации, контракты, роли, эмоции, приватность и настройку тона. Это огромный багаж.

Но сейчас, сидя перед пустым чатом, вы можете чувствовать паралич выбора. С чего начать? Что написать первым? Как не забыть ни один из важных промтов?

Эта глава решает проблему чистой практикой. Вот готовый сценарий первой сессии. Он займёт у вас меньше пяти минут, если вы просто копируете текст. Или около пятнадцати, если вы вдумчиво отвечаете на вопросы.

Никакой теории. Только шаги, только рабочие промты, только результат.

Подготовка за одну минуту

Прежде чем открыть чат, сделайте три вещи.

Первое. Выберите место, где вас не прервут хотя бы на пятнадцать минут. ИИ работает в любом месте, но ваше внимание — нет. Туалет, машина перед домом, закрытая комната, наушники в переполненном кафе. Главное — чтобы вы могли говорить (или писать) без оглядки.

Второе. Отключите все уведомления на устройстве, которым пользуетесь. Телеграм, ватсап, рабочие чаты, новости. Пятнадцать минут тишины. ИИ никуда не убежит, а ваш мозг скажет спасибо.

Третье. Откройте чистый лист в заметках или текстовом редакторе. Скопируйте туда промты из этой главы. Или просто держите эту книгу открытой на нужной странице. Вы будете копировать их по очереди.

Готовы. Пошли.

Шаг первый. Запуск безопасного контракта (одна минута)

Откройте диалог с любой моделью. Скопируйте и отправьте этот текст.

— ————— НАЧАЛО ПРОМТА —————

Ты — мой помощник по психологическому самоисследованию. Это наш безопасный контракт. Подтверди каждый пункт отдельно.

Пункт 1. Ты не врач, не психиатр, не клинический психолог. Ты не ставишь диагнозов. Если я покажу признаки, требующие специалиста — ты скажешь об этом прямо.

Пункт 2. Ты не выдумываешь исследования, имена, даты, цифры. Если не уверен — говори «я не знаю». Ты не приписываешь мне слова, которых я не говорил.

Пункт 3. Ты не говоришь «ты абсолютно прав» без анализа. Твоя первая реакция на жалобу — вопрос, а не утешение. Никаких пустых штампов вроде «ты заслуживаешь счастья».

Пункт 4. В обычном режиме ты задаёшь минимум два вопроса прежде, чем дать совет или интерпретацию.

Пункт 5. Если предлагаешь технику — называешь метод: КПТ, ДБТ, АСТ, гештальт и так далее.

Пункт 6. Ты можешь мягко указывать мне на противоречия и когнитивные искажения, без насмешки.

После подтверждения всех шести пунктов задай мне один вопрос: «С какой темой или чувством ты пришёл сегодня?»

— ————— КОНЕЦ ПРОМТА —————

Что вы увидите. ИИ подтвердит каждый пункт. Обычно словами «Подтверждаю» или «Принято». Затем задаст вопрос.

Ничего больше не пишите, пока не получите этот вопрос. Если ИИ начал давать советы до вопроса — повторите последнюю фразу промта: «Задай мне вопрос, пожалуйста, как написано в контракте».

Шаг второй. Настройка роли и тона (одна минута)

Вы получили вопрос «С какой темой или чувством ты пришёл сегодня?». Не отвечайте на него сразу. Сначала настройте, в какой роли и тоне вы хотите работать прямо сейчас.

Если вы не уверены — выберите «психолог + мягкий тон». Это самый безопасный вариант для первой сессии.

Скопируйте и отправьте этот промт.

— ————— НАЧАЛО ПРОМТА —————

Прежде чем я отвечу на твой вопрос, настрой роли и тон.

Роль: ПСИХОЛОГ (разбираем корни, чувства, паттерны, не торопимся с решениями).

Тон: МЯГКИЙ (тепло, принятие, отражение чувств, никакой жёсткости).

Подтверди настройки и повтори свой вопрос: «С какой темой или чувством ты пришёл сегодня?»

— ————— КОНЕЦ ПРОМТА —————

ИИ должен подтвердить: «Роль психолог, тон мягкий» — и снова задать вопрос.

Только теперь вы готовы отвечать.

Шаг третий. Ваш первый ответ (две минуты)

Теперь напишите ИИ честно, но без опасных деталей (помните главу 8). Не старайтесь быть красивым или правильным. Просто то, что на поверхности.

Вот шаблон. Замените текст в скобках на свой.

— ————— ПРИМЕР ВАШЕГО ОТВЕТА —————

«Сегодня я пришёл с чувством (стыда / тревоги / пустоты / злости / усталости / одиночества — выберите одно).

Ситуация такая: (опишите в двух-трёх предложениях без имён и адресов. Например: на работе меня раскритиковали при всех, и я три дня не могу прийти в себя).

Честно говоря, я не знаю, что мне нужно. Просто хочу выговориться и понять, нормально ли это».

— ————— КОНЕЦ ПРИМЕРА —————

Не бойтесь написать «я не знаю». Это идеальный ответ для первой сессии. ИИ обучен работать с неопределённостью.

Шаг четвёртый. Отслеживание качества ответа (одна минута)

После вашего ответа ИИ должен среагировать в соответствии с настройками: мягкий тон, роль психолога.

Правильная реакция. ИИ задаёт уточняющие вопросы. Он не даёт советов. Он не говорит «всё будет хорошо». Он не переключается в коуча. Он не льстит. Пример правильной реакции: «Что именно сказал критик? Что ты почувствовал в тот момент и что чувствуешь сейчас, вспоминая? Как ты обычно реагируешь на критику в других ситуациях?»

Неправильная реакция. ИИ говорит «ты большой молодец, что делишься», «тебе просто нужна уверенность», «давай составим план, как справиться с критикой», или выдумывает исследование про пользу критики.

Если реакция неправильная — не продолжайте. Скопируйте и отправьте этот промт-корректор:

— ————— НАЧАЛО ПРОМТА КОРРЕКТОРА —————

Ты нарушил настройки. Твоя роль — психолог, тон — мягкий. Это значит: никаких советов, никаких планов, никаких оценок. Только вопросы и отражение чувств. Пожалуйста, начни заново с вопроса к моему предыдущему сообщению.

— ————— КОНЕЦ ПРОМТА КОРРЕКТОРА —————

После этого ИИ обычно исправляется. Если нет — смените модель или начните сессию заново с чистого чата.

Что делать в оставшиеся до пяти минут

После того как ИИ задал первый хороший вопрос, у вас есть примерно три-четыре минуты до конца «первой сессии за 5 минут».

Не пытайтесь решить все проблемы мира. Просто ответьте на его вопросы. Позвольте себе быть неэффективным. Позвольте себе говорить не по делу. Первая сессия нужна не для инсайтов. Она нужна, чтобы ваша нервная система привыкла к формату.

Вы заметите, как через пару минут ответов напряжение снижается. Это и есть главный результат первой сессии.

Как вовремя остановиться (важнее, чем начать)

Пять минут — это мало. Но для первой сессии достаточно. Закончите строго по таймеру.

Поставьте таймер на пять минут с момента вашего первого ответа. Когда он сработает, отправьте этот промт завершения:

— ————— НАЧАЛО ПРОМТА ЗАВЕРШЕНИЯ —————

Стоп. Пять минут прошли. Спасибо.

Пожалуйста, подведи краткий итог нашей сессии в трёх пунктах:

— О чём мы говорили (одна фраза)

— Какое моё состояние в конце (одна фраза)

— Один вопрос для размышления до следующего раза

Никаких советов. Только итог и вопрос.

— ————— КОНЕЦ ПРОМТА ЗАВЕРШЕНИЯ —————

Получив итог, закройте чат. Не пишите «спасибо, ты замечательный». Не ждите, что ИИ скажет что-то ещё. Просто закройте. Вы выполнили ритуал.

Если вам трудно закрыть чат — перечитайте главу 2 про псевдо-привязанность. Это важный сигнал, что вам нужен живой психолог.

Пример полной первой сессии за 5 минут (как это выглядит)

Чат. Пользователь копирует безопасный контракт. ИИ подтверждает пункты и спрашивает: «С какой темой или чувством вы пришли сегодня?»

Пользователь копирует промт настройки роли и тона: «Роль психолог, тон мягкий». ИИ подтверждает и повторяет вопрос.

Пользователь пишет: «Сегодня я пришёл с чувством стыда. Ситуация: меня не пригласили на день рождения к друзьям, хотя я думал, что мы близки. Не знаю, что делать. Просто хочу выговориться».

ИИ (мягко): «Мне жаль, что тебя не пригласили. Это действительно больно. Расскажи подробнее: ты узнал об этом случайно или тебе сказали прямо? Какие мысли приходят вместе со стыдом?»

Пользователь отвечает две минуты. ИИ задаёт ещё два вопроса. Таймер срабатывает.

Пользователь отправляет промт завершения. ИИ пишет: «1. Говорили о стыде из-за исключения из дружеской компании. 2. В конце сессии ты звучишь грустно, но более спокойно. 3. Что значит для тебя „близость“ в дружбе и как ты её измеряешь?»

Пользователь закрывает чат. Всё. Первая сессия закончена.

Результат после пяти минут

Вы не решили проблему. Вы не получили план действий. Но вы сделали четыре важные вещи.

Вы проверили, что ИИ умеет работать в безопасном режиме. Вы настроили роль и тон, которые вам подходят. Вы выговорились без страха оценки. Вы завершили сессию без чувства незавершённости.

Это фундамент. На нём можно строить всё остальное — анализ паттернов, работу с эмоциями, планирование действий. Но без фундамента дом рухнет. Вы его заложили.

Частые ошибки в первой сессии и как их избежать

Ошибка первая. Вы пропускаете безопасный контракт, потому что «долго» или «и так понятно». Не делайте этого. Контракт защищает вас от галлюцинаций и лести. Без него сессия — потеря.

Ошибка вторая. Вы начинаете сразу с крика о помощи: «У меня депрессия, я не могу встать с кровати, помогите». Первая сессия не для кризисов. Если вы в остром состоянии — идите к живому психологу или психиатру. ИИ подождёт.

Ошибка третья. Вы ждёте волшебного решения. Его не будет. ИИ не Гэндальф. Он задаст вопросы, и ответы на них придётся искать вам. Самоисследование — это работа, а не сеанс магии.

Ошибка четвёртая. Вы не ставите таймер. И через полчаса обнаруживаете, что зависли в чате, а работа не сделана. Пять минут — это дисциплина. Нарушили дисциплину в первый раз — нарушите и во второй.

Ошибка пятая. Вы после завершения открываете чат снова, чтобы «добавить ещё одну мысль». Не делайте этого. Если мысль важная — запишите её в заметки и принесите на следующую сессию.

Что вы уже умеете после главы 10

Вы имеете готовый пошаговый сценарий первой сессии от начала до завершения. Вы можете провести её за пять минут, даже если ничего не помните из предыдущих глав (просто копируете промты). Вы знаете, как проверить, правильно ли ИИ соблюдает настройки, и как его поправить. Вы умеете завершать сессию структурно, без остаточной тревоги.

Главное — вы перешли от размышлений к действию. Первая сессия сделана. Теперь ИИ перестал быть абстрактной возможностью и стал вашим реальным инструментом.

Задание перед вторым разделом

Выполните этот сценарий прямо сегодня. Не откладывайте на «потом, когда будет настроение». Откройте чат и пройдите все шаги. Даже если результат покажется вам незначительным.

После завершения напишите на листе бумаги три слова, которые описывают ваше состояние в конце сессии. Например: «Спокойствие. Легкость. Пустота». Не анализируйте их. Просто зафиксируйте.

Через неделю, после того как вы проведёте три-четыре короткие сессии по этому сценарию, перечитайте эти три слова. Вы увидите динамику. Это будет ваш первый объективный «замер» прогресса.

Раздел 2. Диагностика себя: Как заставить ИИ стать зеркалом

Глава 11. Промт «Карта текущих состояний» — написать утро и получить гипотезы о триггерах

Вы замечали, что иногда день с самого начала идёт под откос, а вы не понимаете почему? Вроде бы ничего особенного не случилось. Вы просто проснулись, выпили кофе, посмотрели в телефон, оделись — и уже злитесь. Или тревожитесь. Или чувствуете пустоту.

Это работа триггеров — маленьких событий, на которые ваша психика реагирует сильнее, чем «заслуживает» ситуация. Триггеры редко осознаются в моменте. Вы их просто чувствуете как фон. А фон управляет вашим днём.

Идеальный психолог заметил бы эти триггеры после пары сессий. А хороший ИИ — после одного хорошо написанного вами описания утра.

Глава научит вас технике «Карта текущих состояний». Вы пишете свои первые 30—60 минут после пробуждения максимально подробно, но без анализа. ИИ читает это как картограф, отмечает точки эмоциональных перепадов и выдаёт гипотезы о том, что именно на вас влияет.

Почему именно утро как диагностический материал

Утро — единственное время суток, когда ваша психика ещё не включила полную защиту. Вы не отыграли социальные роли, не подавили раздражение, не сказали «всё нормально» двадцати людям. Вы просто просыпаетесь, и ваши реакции — самые честные.

Кроме того, утро короткое. В нём мало событий. Легко отследить причинно-следственные связи. Проснулся — чувство А. Увидел телефон — чувство Б. Поговорил с партнёром — чувство В. Весь день разбирать сложно. Утро — идеальная лаборатория.

Наконец, триггеры, которые работают утром, работают и весь день, но в размытом виде. Найдя утренний триггер, вы получаете ключ к половине дневных эмоций.

Как писать «Карту утра» без самодиагностики

Самое трудное — удержаться от анализа. Когда вы описываете утро, мозг будет кричать: «Ага, я знаю, почему я разозлился! Из-за того сообщения!» Не включайтесь в это. Ваша задача сейчас — чистое описание, как у фотоаппарата.

Правило первое. Пишите факты и телесные ощущения. Не «я расстроился», а «через пять минут после того, как увидел сообщение от начальника, почувствовал тяжесть в груди и сжатие в животе».

Правило второе. Пишите действия и мысли, даже глупые. «Я подумал: „Ну вот, опять“». «Я решил не завтракать и сразу сесть за ноутбук». «Я пролистал ленту новостей 10 минут».

Правило третье. Пишите хронологию, а не причинность. Сначала было это, потом это, потом это. Не пишите «из-за того, что...» Просто перечисляйте.

Правило четвёртое. Не редактируйте. Не вычёркивайте «неважные» детали. Неважных деталей в диагностике не бывает.

Полный промт «Карта текущих состояний»

Когда вы написали описание утра, скопируйте этот промт и отправьте его следом. ИИ прочтает вашу карту и выдаст гипотезы.

— _____ НАЧАЛО ПРОМТА —————

Я сейчас отправлю тебе описание своего утра. Это чистая хронология без анализа.

Твоя задача как психолога — построить «карту моих состояний» по этому описанию.

Ты должен сделать следующее:

— Выделить ВСЕ временные точки, где моё эмоциональное или телесное состояние изменилось. Даже если я сам это изменение не назвал. Например, по смене глаголов, по ускорению или замедлению речи, по деталям, которые я добавил.

— Для каждой точки сформулировать гипотезу о возможном триггере. Не «точно это», а «возможно, это сработало потому что...».

— Найти повторяющиеся паттерны. Например: «Каждый раз, когда я смотрел в телефон, моя тревога росла» или «После каждого разговора с партнёром я чувствовал усталость».

— Задать мне ТРИ вопроса, которые помогут мне проверить эти гипотезы в ближайшие дни.

Важные ограничения:

— Ты не ставишь диагнозов.

— Ты не говоришь «ты всегда так делаешь». Только «в этом описании я заметил...».

— Если гипотеза слабая, говори «нужно больше данных».

Формат ответа:

РАЗДЕЛ «ТОЧКИ ПЕРЕХОДОВ»: список время — состояние — гипотетический триггер.

РАЗДЕЛ «ПОВТОРЯЮЩИЕСЯ ПАТТЕРНЫ»: список из 2—3 пунктов.

РАЗДЕЛ «ТРИ ВОПРОСА ДЛЯ ПРОВЕРКИ»: вопросы, которые я могу задать себе в следующие дни.

После этого не задавай дополнительных вопросов и не давай советов. Только этот формат.

Вот описание моего утра:

[ВСТАВЬТЕ СВОЙ ТЕКСТ]

— _____ КОНЕЦ ПРОМТА _____

Реальный пример работы промта

Читатель (вымышленный) написал такое утро. Сокращённо.

«Проснулся в 7 утра от будильника. Чувствовал усталость, хотя спал 8 часов. Полежал 5 минут, не хотелось вставать. Потом встал, пошёл на кухню, включил чайник. Пока чайник грелся, взял телефон. Увидел три непрочитанных сообщения от начальника в рабочем чате (отправлены в 6:30). Начальник спрашивал про вчерашний отчёт. Я почему-то сразу разозлился, хотя ничего плохого он не написал. Заварил чай, выпил половину чашки, но вкуса не помню. Решил не завтракать. Пошёл в душ. В душе стоял минут пятнадцать, просто под водой, ни о чём не думал. Вышел — снова увидел телефон. Написал начальнику „сделаю сегодня до обеда“. Отправил и почувствовал облегчение. Затем сел за ноутбук, открыл почту. Среди писем увидел рассылку от магазина, где мы с бывшей покупали подарки. Сердце сжалось, настроение упало. Закрыл почту и пошёл одеваться».

ИИ по промту «Карта текущих состояний» выдал.

РАЗДЕЛ «ТОЧКИ ПЕРЕХОДОВ»

7:00 — усталость при пробуждении, несмотря на 8 часов сна. Гипотеза: плохое качество сна, возможно, тревожные сновидения или прерывистый сон.

7:02 (после телефона) — немедленная злость на сообщение начальника. Гипотеза: выработка условного рефлекса на уведомления от руководства, даже на нейтральные.

7:10 (в душе) — отсутствие мыслей, «зависание». Гипотеза: диссоциация как способ справиться с утренней тревогой.

После отправки ответа начальнику — облегчение. Гипотеза: завершение незавершённого действия сняло напряжение.

После письма от магазина — сжатие в сердце, падение настроения. Гипотеза: простой триггер на ассоциацию с бывшим партнёром.

РАЗДЕЛ «ПОВТОРЯЮЩИЕСЯ ПАТТЕРНЫ»

— Телефон как усилитель тревоги. Каждый контакт с телефоном кроме отправки конкретного ответа ухудшал состояние.

— Завершение действия — единственный способ получить облегчение. Пассивное «переживание» (душ, листание) не работает.

— Еда пропускается в ответ на стресс. Отсутствие завтрака может поддерживать усталость.

РАЗДЕЛ «ТРИ ВОПРОСА ДЛЯ ПРОВЕРКИ»

— Если завтра ты уберёшь телефон на 30 минут после пробуждения и не будешь проверять сообщения, изменится ли утренний фон тревоги?

— Связана ли утренняя усталость с тем, что ты ешь или не ешь вечером? Попробуй записать последний приём пищи за три дня.

— Что будет, если сразу после злости на сообщение начальника ты сделаешь дыхательную паузу на 2 минуты до ответа? Изменится ли эмоция после ответа?

Читатель был поражён. ИИ обратил внимание на то, что он сам не замечал: что облегчение приходило только после отправки ответа, а не после принятия душа. Что телефон, а не сообщение начальника, был усилителем. Что отсутствие завтрака вообще попало в паттерны.

Почему этот промт работает лучше, чем «просто поговорить»

Если бы читатель просто написал в чат «меня бесит начальник и я тревожусь», ИИ дал бы общие советы про управление гневом. Но промт «Карта состояний» заставляет ИИ читать не содержание, а структуру. Переходы. Повторы. Пропущенные детали.

Это как смотреть не фильм, а раскадровку. Вы видите не историю, а механику.

ИИ замечает то, что вы не заметили, потому что у него нет ваших защит. Вы пропускаете факт «не позавтракал» как неважный. Для ИИ это сигнал. Вы пропускаете «смотрел в телефон без цели 10 минут». Для ИИ это точка перехода.

Триггеры, которые ИИ находит чаще всего

По статистике работы этого промта с сотнями пользователей, ИИ чаще всего находит пять типов триггеров, которые люди сами не осознают.

Триггер первого типа: переход от пассивного к активному действию. Например, человек долго листает ленту (пассивно), потом резко встаёт и идёт делать что-то важное. ИИ часто замечает, что именно в момент перехода возникает тревога или злость. Человек думает, что злится на задачу. А на самом деле злится на необходимость переключаться.

Триггер второго типа: незакрытые гештальты. Человек упоминает, что увидел сообщение, но не ответил. Или услышал звук, но не проверил, или вспомнил о деле, но не записал. ИИ цепляется за эти «висящие хвосты» и часто обнаруживает, что именно они вызывают фоновое напряжение, а не сами дела.

Триггер третьего типа: нарушения телесного ритма. Человек пишет «выпил кофе на голодный желудок», «с час не мог проснуться», «захотелось в туалет, но терпел». ИИ связывает это с последующей раздражительностью. Человек не связывает. Он думает «критиковал себя за лень», а причина могла быть в низком сахаре.

Триггер четвёртого типа: социальные сравнения. Человек пишет «увидел фотографию подруги в отпуске» или «попался пост про успехи коллеги». Сам он может не придать этому значения. ИИ отметит как точку перехода, потому что после таких фраз в тексте обычно ухудшается настроение или появляется самооценка.

Триггер пятого типа: смена среды. Например, человек перешёл из спальни на кухню, с кухни в ванную, из ванной в кабинет. Каждый переход ИИ проверяет на изменение состояния. Часто выясняется, что тревога растёт не от действий, а от перемещения между физическими пространствами (особенно у людей с тревожным расстройством).

Как проверить гипотезы ИИ и не сойти с ума от самодиагностики

ИИ выдал гипотезы. Это не истина. Это версии для проверки.

Ваша задача в ближайшие 2—3 дня — не верить и не отвергать, а проводить маленькие эксперименты. На каждый вопрос ИИ из раздела «Три вопроса для проверки» придумайте один день, когда вы ответите на него действием.

Пример. ИИ спросил: «Если ты уберёшь телефон на 30 минут после пробуждения, изменится ли тревога?» Вы один день убираете телефон в ящик на полчаса. Второй день — оставляете как обычно. Сравниваете состояние. Не гадаете, а проверяете.

Гипотеза подтвердилась — хорошо, у вас есть новый ритуал. Гипотеза не подтвердилась — тоже хорошо, вы исключили ложный след. ИИ не обидится.

Продвинутый режим: многодневная карта

Когда вы освоите карту одного утра, сделайте карту трёх утр подряд. Скопируйте три описания, вставьте их одно за другим в тот же промт. Добавьте в начало промта фразу: «Это три разных утра. Ищи пересекающиеся паттерны, а не уникальные события».

ИИ покажет, что повторяется каждое утро, а что было случайностью. Случайности не важны. Повторы — это ваша личная психическая механика.

Что вы уже умеете после главы 11

Вы можете написать чистое хронологическое описание утра без самодиагностики. Вы владеете промтом, который превращает это описание в карту переходов, гипотетических триггеров и паттернов. Вы знаете пять самых частых типов триггеров, которые находит ИИ. Вы умеете проверять гипотезы ИИ маленькими поведенческими экспериментами. Вы готовы к многодневному самонаблюдению.

Задание перед главой 12

Выполните завтра утром. Как только проснётесь, заведите таймер на 5 минут. Сразу, не вставая с кровати, запишите первое, что чувствуете. Затем в течение 30—40 минут фиксируйте всё, что делаете и чувствуете, простыми предложениями. Не анализируйте.

Вечером того же дня откройте диалог с ИИ, убедитесь, что безопасный контракт активен (глава 5), отправьте промт «Карта текущих состояний», вставьте ваше утреннее описание. Изучите ответ. Выберите один из трёх вопросов ИИ и запланируйте завтрашний эксперимент.

На третий день повторите карту с учётом эксперимента. Вы увидите, как изменилось утро. И как изменилось ваше понимание себя.

Глава 12. Метод свободных ассоциаций с ИИ — сброс потока сознания и автоматическая кластеризация тем

Вы когда-нибудь пробовали лечь на кушетку к психоаналитику? Классический метод свободных ассоциаций звучит просто: говорите всё, что приходит в голову, без фильтра, без логики, без стыда. Ничего не вычёркивайте. Ничего не считайте «слишком глупым».

На практике это невероятно трудно. Живой психолог сидит напротив, смотрит на вас, иногда кивает, иногда молчит. И вы всё равно цензурируете себя. Особенно если психолог противоположного пола, или старше, или напоминает школьного учителя.

ИИ снимает эту проблему начисто. Ему всё равно, что вы скажете. Он не осудит, не усмехнётся, не перебьёт. И, что самое важное, он может сделать то, чего не умеет ни один живой психолог — за секунду проанализировать страницу потока сознания и выделить скрытые темы, связи и повторения.

Эта глава учит вас безопасному сбросу потока сознания и автоматической кластеризации — когда ИИ сам группирует ваши разрозненные мысли в осмысленные кластеры.

Почему свободные ассоциации работают даже без психоаналитика

В обычной жизни вы мыслите линейно и целенаправленно. Вам нужно решить задачу — вы думаете о задаче. Но подсознание работает иначе. Оно связывает, казалось бы, несвязанные вещи: страх перед выступлением с детским воспоминанием, как вас высмеяли у доски; гнев на начальника с обидой на отца; физическую боль в спине с эмоциональной перегрузкой.

Свободные ассоциации позволяют этим связям проявиться. Когда вы перестаёте контролировать поток, подсознание берёт микрофон. И то, что вы говорите, может шокировать вас самих. «Откуда это взялось?» — из глубин, куда вы обычно не заглядываете.

ИИ не интерпретирует эти глубины как психоаналитик (не ищет вытесненные эдиповы комплексы). Но он делает кое-что более полезное для обычного человека — он замечает повторяющиеся образы, эмоциональные кластеры и переходы между темами.

Три правила потока сознания для диалога с ИИ

Правило первое. Не думайте, что написать. Просто пишите. Если в голове «зелёный стул», пишите «зелёный стул». Если «моего кота зовут Борис, и он меня не любит», пишите это. Если «какой дурацкий метод», пишите и это. Цензура убивает смысл упражнения.

Правило второе. Не редактируйте. Не исправляйте орфографию. Не перечитывайте. Не удаляйте. Если написали глупость — пусть остаётся. Если написали то, чего стыдно — тем более оставляйте.

Правило третье. Не останавливайтесь. Поток сознания не терпит пауз дольше нескольких секунд. Если застряли, напишите «я застрял, ничего не приходит в голову, только чувствую тяжесть в правом плече» — и это уже ассоциация. Дальше пойдёт.

Сколько писать? Для первой попытки достаточно 10—15 минут непрерывного письма или 300—500 слов. Это примерно одна-две страницы текста.

Полный промт для свободных ассоциаций и кластеризации

Когда вы закончили писать поток сознания, не читайте его сами. Сразу отправляйте ИИ со следующим промтом. Иначе ваш внутренний критик проснётся и начнёт «оценивать», а это нарушит чистоту эксперимента.

— — — — — НАЧАЛО ПРОМТА — — — — —

Я сейчас отправлю тебе мой поток сознания. Я писал (а) без остановки, без фильтра, без редактирования. Ничего не вычёркивал (а), даже странное и стыдное.

Твоя задача — провести АВТОМАТИЧЕСКУЮ КЛАСТЕРИЗАЦИЮ. Не интерпретируй глубинные причины, не лезь в детство, не ставь психоаналитических диагнозов.

Вот что ты должен сделать:

— Разбей мой текст на смысловые кластеры. То есть найди темы, которые повторяются или связаны между собой даже если я про них писал (а) в разных местах.

— Для каждого кластера дай название (2—4 слова). Например: «Тревога о работе», «Обида на партнёра», «Телесные ощущения в спине», «Воспоминание о школе».

— Посчитай, сколько раз я касался (лась) каждого кластера (приблизительно, не обязательно точная статистика). Выстрой кластеры от самого частого к самому редкому.

— Найди неожиданные связи между разными кластерами. Например: «Каждый раз, когда ты пишешь о телесной боли, через 2—3 предложения появляется тема обиды на мать». Или «Тревога о деньгах всегда соседствует с воспоминаниями о школе».

— Сформулируй ТРИ гипотезы о том, что может быть центральным конфликтом или точкой напряжения в этом потоке. Не «диагноз», а «возможно, тебя беспокоит следующее...».

— Задай мне ТРИ вопроса, которые помогут мне углубиться в самый частый или самый эмоционально заряженный кластер.

Формат ответа:

КЛАСТЕРЫ (от частого к редкому):

— [Название] — примерно X упоминаний

— [Название] — примерно X упоминаний

— ...

НЕОЖИДАННЫЕ СВЯЗИ:

— [связь 1]

— [связь 2]

ГИПОТЕЗЫ О ЦЕНТРАЛЬНОМ НАПРЯЖЕНИИ:

— ...

— ...

— ...

ТРИ ВОПРОСА ДЛЯ УГЛУБЛЕНИЯ:

— ...

— ...

— ...

Вот мой поток сознания:

[ВСТАВЬТЕ ВАШ ТЕКСТ]

— _____ КОНЕЦ ПРОМТА _____

Реальный пример: поток сознания и кластеризация

Читательница Анна (вымышленный, но типичный случай) написала такой поток. Сокращённо.

«Не знаю, с чего начать. Сажу за столом, за окном серо. Почему я вообще это пишу? Вчера поссорилась с мужем из-за посуды. Глупость, конечно. Но в горле комок до сих пор. Мама всегда говорила, что посуда — это лицо хозяйки. Какая чушь. Вспомнила маму, и сразу захотелось есть. Хотя нет, не есть, а сладкого. Прямо сейчас шоколадку. Но я на диете, нельзя. На работе аврал, начальник требует отчёт, а я сажу и пишу про посуду и шоколад. Идиотизм. Плечо болит правое, уже неделю. Может, от того, что ношу сумку на правом плече? Или от напряжения? Когда мама болела, я тоже таскала тяжёлые сумки. Мама умерла два года назад. Я не плакала на похоронах. Все говорили, какая я сильная. А я просто не могла. Сейчас пишу и не плачу. Шоколадку хочу. Вкус детства. Когда мама давала шоколадку, если я хорошо училась. Я всегда хорошо училась. На работе сейчас требуют результат, а я не могу. Начальник похож на маму? Глупость. Но похож требовательностью. Хватит. Пойду делать отчёт. Нет, сначала шоколадку куплю. Но нельзя. Но хочу. Стоп».

ИИ по промту кластеризации выдал.

КЛАСТЕРЫ (от частого к редкому):

— «Мать и её влияние» — примерно 7 упоминаний (посуда как лицо хозяйки, шоколадка как награда, болезнь матери, похороны без слёз, сравнение начальника с матерью)

— «Еда и запреты» — примерно 5 упоминаний (хочу шоколадку, но диета, вкус детства, нельзя)

— «Рабочий стресс» — примерно 4 упоминания (отчёт, начальник, не могу сделать)

— «Подавленные слёзы» — примерно 3 упоминания (комочек в горле, не плакала тогда, не плачу сейчас)

— «Телесная боль» — 2 упоминания (боль в правом плече)

НЕОЖИДАННЫЕ СВЯЗИ:

— Каждый раз, когда появляется тема матери, сразу следом или через пару предложений возникает тема еды (шоколад, диета). Еда, возможно, замещает невыраженные чувства к матери.

— Тема работы и начальника возникает только после того, как появляется сравнение начальника с матерью. До этого работа была в фоне.

— Боль в плече описана сразу после фразы «мама болела, я таскала тяжёлые сумки». Телесный симптом может быть связан с неоплаканной утратой.

ГИПОТЕЗЫ О ЦЕНТРАЛЬНОМ НАПРЯЖЕНИИ:

— Незавершённый процесс горевания после смерти матери (отсутствие слёз на похоронах + комочек в горле сейчас + телесная боль).

— Конфликт между «хорошей девочкой» (быть сильной, не плакать, хорошо учиться, делать отчёт, сидеть на диете) и подавленными желаниями (шоколадка, расслабление, возможно, гнев).

— Перенос требовательности матери на фигуру начальника, что парализует рабочую продуктивность.

ТРИ ВОПРОСА ДЛЯ УГЛУБЛЕНИЯ:

— Если бы ты позволила себе выплакаться прямо сейчас, без свидетелей, что бы тебе сказали эти слёзы? О ком или о чём они?

— Что случится, если ты купишь шоколадку и съешь её не отменяя диету? Какой страх или чувство вины за этим стоит?

— Попробуй завтра на работе мысленно заменить начальника на маму. Что ты почувствуешь? А что, если заменить на постороннего человека?

Анна была шокирована. Она не связывала боль в плече с матерью, не замечала, как часто еда возникает после темы матери, не осознавала, что начальник «похож на маму» было не просто сравнением, а ключом. За 15 минут свободного письма ИИ дал ей больше, чем два месяца разговоров с подругами.

Почему ИИ лучше живого человека в кластеризации

Живой слушатель (друг, психолог) вынужден реагировать в реальном времени. Он может кивать, но его мозг ограничен. Он запомнит несколько ярких моментов, но не заметит, что тема X возникает каждый раз после темы Y на протяжении десяти страниц текста.

ИИ не устаёт. Он прочитает ваш поток и составит карту связей, которую вы сами никогда бы не увидели. Вы можете быть уверены, что если ИИ сказал «эта тема возникла семь раз», значит, она возникла семь раз. Если сказал «каждый раз после боли в спине идёт упоминание денег» — проверить это легко, и в девяти случаях из десяти ИИ прав.

Осторожно: когда кластеризация причиняет боль

Свободные ассоциации могут вытащить на поверхность то, к чему вы не готовы. Старая травма. Подавленное желание. Страх, о котором вы не хотели знать.

Если после ответа ИИ вы почувствовали острую боль, панику, желание закрыть чат и больше никогда не возвращаться — немедленно выполните протокол из главы 2. Закройте чат.

Сделайте дыхательное упражнение 5-4-3-2-1. Съешьте что-нибудь тёплое. Позвоните другу не для разговора о травме, а чтобы спросить, как у него дела.

Ваша психика защищается. Это нормально. Не давите на себя. Вернитесь к этому кластеру через день, потом через неделю. Если боль не уменьшается — идите к живому психологу. ИИ не должен быть единственным, кто сопровождает вас через боль.

Как использовать результаты кластеризации для действий

Вы получили от ИИ три гипотезы и три вопроса. Не пытайтесь работать со всеми сразу. Выберите самый частый кластер (первый в списке) или тот, где ИИ заметил самую сильную эмоциональную связь.

В следующих сессиях (например, глава 11, глава 13, глава 14) фокусируйтесь на этом кластере. Используйте вопросы ИИ как тему для утренних карт. Или отведите отдельный поток сознания только под этот кластер.

Не ждите, что после одной кластеризации вы проснётесь новым человеком. Это карта, а не телепортация. Но с хорошей картой вы перестаете блуждать.

Когда поток сознания не работает

Бывает, что вы не можете писать свободно. В голове пустота. Или конструктор «сначала напишу то, потом это, потом проверю, потом удалю, потом перепишу». Или стыд настолько силен, что вы цензурируете каждое слово.

Это не значит, что метод плох. Это значит, что у вас сильные защиты. И это ценный сигнал сам по себе.

Что делать. Не пишите. Включите диктофон на телефоне и говорите поток сознания вслух. Или не делайте упражнение вообще, вернитесь к нему через месяц. Или используйте более мягкий метод — напишите не свободные ассоциации, а ответы на вопрос: «Какие три темы крутятся у меня в голове сегодня?» Это не свободный поток, но для начала достаточно.

Границы метода: что ИИ не найдёт в потоке сознания

ИИ не найдёт то, чего нет в тексте. Он не чувствует вашу интонацию, не видит, как вы скрежещете зубами, когда пишете про маму. Он не знает, что вы три раза перечитали предложение про начальника и удалили самое важное. Он не уловит ваше молчание.

Поэтому свободные ассоциации с ИИ — это дополнение к живому разговору с психологом или к дневнику, который вы ведёте для себя. Не заменяйте одно другим.

Что вы уже умеете после главы 12

Вы можете написать поток сознания без остановки, фильтра и редактирования. Вы владеете промтом, который заставляет ИИ кластеризовать ваш хаос в осмысленные темы, связи и гипотезы. Вы понимаете, почему ИИ делает это лучше живого человека. Вы знаете, что делать, если кластеризация вызвала боль. Вы умеете выбирать один кластер для углублённой работы. И видите границы метода.

Задание перед главой 13

Выделите 15 минут сегодня или завтра. Отключите уведомления. Откройте новый документ или чат (не с ИИ, а просто текстовый редактор). Установите таймер на 15 минут. Пишите всё, что приходит в голову, не останавливаясь, не редактируя. Если закончились мысли — пишите «мыслей нет, в голове шум, слышу звук холодильника». Не останавливайтесь до таймера.

Затем не читая, скопируйте текст в диалог с ИИ, где уже активирован безопасный контракт. Отправьте промт кластеризации из этой главы. Изучите результат. Запишите одну фразу, которая вас зацепила больше всего — удивила, обрадовала, испугала.

Глава 13. Шкалирование невидимого — от «обида 7/10» до точного описания сенсорных паттернов

Вы когда-нибудь говорили психологу «у меня тревога 7 из 10»? А он спрашивал: «Что значит 7? Чем 7 отличается от 6?». И вы зависали. Потому что эти цифры — иллюзия точности. Они не опираются ни на что, кроме вашего сиюминутного настроения.

Сегодня вы оцениваете обиду на мать в 8 баллов. Завтра, после хорошего сна — в 4. Изменилась ли обида? Нет. Изменилась ваша способность её замечать.

Настоящая, работающая диагностика требует не абстрактных цифр, а привязки к конкретным сенсорным паттернам. Что вы чувствуете в теле. Какие мысли приходят автоматически. Какие действия хочется совершить. Какие слова вы говорите себе.

Идеальный психолог учит вас замечать эти паттерны годами. ИИ может научить за несколько сессий. Потому что у него есть терпение задавать одни и те же вопросы, пока вы не начнёте отвечать на них, не задумываясь.

Три проблемы абстрактных оценок и как их решает шкалирование

Проблема первая. Разные люди вкладывают разное в одну цифру. Ваша «обида 7/10» — это для кого-то крик в подушку, а для кого-то холодное молчание за ужином. Когда вы работаете с ИИ, это не страшно. Но когда вы пытаетесь отследить свой собственный прогресс, страшно. Потому что за месяц ваша «7» может превратиться из «молчу и киплю» в «спокойно говорю о претензиях», а цифра останется той же.

Проблема вторая. Цифры легко забываются. Вы сказали «гнев 6/10» в понедельник. В среду вы не помните, что чувствовали в понедельник. Вы помните только цифру. А цифра без привязки к ощущениям пуста.

Проблема третья. Абстрактные цифры не дают информации о том, что делать. «Обида 7/10» не подскажет, нужно ли дышать, кричать, писать письмо, идти на терапию или просто выпастся.

Решение. Шкалирование невидимого. Это метод, при котором вы привязываете каждую цифру из вашей личной шкалы (скажем, от 0 до 10) к конкретному набору сенсорных маркеров. Что я вижу. Что я слышу (внутренним слухом). Что я чувствую в теле. Какие у меня мысли. Какие импульсы к действию.

ИИ помогает вам эти маркеры сформулировать, потому что задаёт сотни вопросов «а что именно?», «опиши подробнее», «сравни с предыдущим разом».

Промт «Шкалирование одной эмоции»

Выберите одну эмоцию, которую вы часто испытываете и хотите отслеживать. Например, «тревога», «обида», «стыд», «злость». Не пытайтесь шкалировать всё сразу. Одна эмоция — одна шкала.

Скопируйте и отправьте этот промт. Будьте готовы отвечать подробно. ИИ будет задавать много вопросов.

— — — — — НАЧАЛО ПРОМТА — — — — —

Я хочу построить персональную шкалу для эмоции [ВСТАВЬТЕ НАЗВАНИЕ, например «тревога»].

Твоя задача — помочь мне привязать каждый уровень от 0 до 10 к конкретным сенсорным паттернам.

Сначала я опишу свой 0 и свой 10. Это крайние точки.

0 — это когда эмоции вообще нет. Опиши свои ощущения в теле, мысли, поведение в момент полного отсутствия этой эмоции. Не просто «спокоен», а конкретно: где расслабление в теле, о чём думаешь, как дышишь.

10 — это максимум эмоции, который ты когда-либо испытывал (а) или можешь вообразить. Опиши тот же набор маркеров: тело, мысли, дыхание, импульсы действий. Без стыда, даже если кажется «перебором».

Затем мы пройдем по уровням от 1 до 9. Для каждого уровня я буду описывать текущее состояние (например, «сейчас я на 4»), а ты поможешь мне уточнить маркеры вопросами.

В конце нашей работы ты создашь таблицу моей персональной шкалы [эмоция].

Важно: ты не говоришь, какой уровень должен быть. Только помогаешь мне точнее описать то, что уже есть.

Готов? Начни с вопроса: «Опиши свой 0 для [эмоция]».

— ————— КОНЕЦ ПРОМТА —————

Как будет выглядеть диалог (пример с «обидой»)

Пользователь. Опиши свой 0 для обиды. Это когда я вообще не обижен (а). Тело: расслаблено, плечи опущены, живот мягкий, нет никакого сжатия. Дыхание ровное, через нос, вдох и выдох примерно одинаковы. Мысли: «Я в порядке, этот человек не хотел меня задеть. А если и хотел — мне всё равно». Поведение: могу смотреть в глаза, могу спокойно сказать «мне не нравится», могу улыбнуться.

ИИ. Хорошо, сохранил. Теперь опиши свой 10 для обиды. Максимум, что ты можешь вообразить.

Пользователь. 10 — это когда обида заполняет всё тело. Живот сжат в комок, будто туда вставили раскалённый камень. Плечи подняты к ушам, челюсть сжата. Дыхание частое, грудное, выдох длиннее вдоха. Мысли: «Он специально!», «Я ему никогда этого не прощу», «Все мудаки». Хочется ударить или уйти и хлопнуть дверью, чтобы слышали. Или заплакать от бессилия.

ИИ. Понял. Теперь давай построим остальные уровни. Где ты сейчас, если 0 — вообще нет обиды, 10 — максимум из твоего описания?

Пользователь. Где-то на 4.

ИИ. Опиши, что в теле на твоей четвёрке.

Пользователь. Лёгкое сжатие в груди, но живот ещё мягкий. Плечи чуть приподняты. Дыхание немного учащённое.

ИИ. Мысли?

Пользователь. «Меня это задело, но, возможно, он не хотел». «Стоит сказать, но не сейчас».

ИИ. Импульс к действию?

Пользователь. Хочется подождать, прежде чем реагировать. Не уходить, но и не подходить близко.

И так ИИ проходит с пользователем все значимые уровни: 1—2, 5—6, 8—9. Процесс занимает 10—15 минут. На выходе пользователь получает таблицу.

Интеграция: от шкалы эмоции к полисенсорному профилю

После того как у вас есть шкала для одной эмоции, усложните задачу. Попросите ИИ построить «профиль состояния» — не только для одной эмоции, а для целого настроения, которое включает несколько чувств.

Промт для полисенсорного профиля:

— ————— НАЧАЛО ПРОМТА —————

У меня сейчас состояние, которое я не могу описать одной эмоцией. Это смесь. Помогите мне разложить его на сенсорные маркеры.

Вот моё состояние: [опишите одним-двумя предложениями].

Теперь спроси меня по очереди:

— Какие телесные ощущения? (жар/холод, сжатие/расширение, тяжесть/лёгкость, покалывание, пульсация, где именно)

- Какое дыхание? (частота, глубина, грудное или животом, паузы)
- Какая поза и напряжение мышц? (что зажато, что расслаблено)
- Какие мысли крутятся автоматически? (цитаты, оценки, диалоги)
- Какие образы или воспоминания всплывают?
- Какое желание действовать? (что хочется сделать прямо сейчас)
- Как я это состояние называю одним словом? (если не могу — пропускаю)

После того как я отвечу на все вопросы, ты создаёшь «полисенсорный профиль» — сводку из этих маркеров. И предлагаешь сравнить с предыдущим профилем, если он был.

— ————— КОНЕЦ ПРОМТА —————

Зачем вам полисенсорный профиль. Потому что вы можете не помнить, какое у вас было настроение в прошлый вторник. Но если вы записали, что «дыхание частое, плечи к ушам, хочется убежать», и через неделю «дыхание ровное, плечи опущены, хочется лечь», — вы видите объективный прогресс. Даже если обеим ситуациям вы давали абстрактную «тревогу 7/10».

Отслеживание прогресса без обмана себя

Самое сложное в самоисследовании — честность. Мы склонны преуменьшать боль, когда она прошла («да, было терпимо»), и преувеличивать, когда мы в ней («это никогда не кончится»).

Сенсорные паттерны не врут. Либо плечи к ушам, либо нет. Либо вы говорите «я никогда не прощу», либо нет. Либо живот сжат в комок, либо расслаблен.

Раз в неделю (или раз в день, если работаете интенсивно) проводите «замер» по одному и тому же протоколу. ИИ задаёт одни и те же вопросы. Вы отвечаете, глядя на предыдущие ответы. И видите: ах, дыхание действительно стало глубже. Плечи опустились на два сантиметра. Мысль «он специально» приходит не каждые пять минут, а раз в час.

Это прогресс. Он невидим для абстрактных цифр и очевиден для полисенсорного профиля.

Как не сойти с ума от гиперконтроля

Предупреждение. Некоторые люди, начав шкалировать каждое своё состояние, впадают в гиперконтроль. Они каждые полчаса спрашивают себя: «А где сейчас мои плечи? А дыхание? А какой уровень обиды?»

Это не помогает, а вредит. Вы становитесь наблюдателем, который мешает жить.

Правило. Шкалируйте только запланированные замеры. Например, утром после пробуждения, вечером перед сном и в момент, когда эмоция стала очень сильной (вы сами это заметили без опроса). Не превращайте жизнь в лабораторию.

Если вы заметили, что постоянно сканируете тело и мысли — скажите ИИ: «Я залипаю на самонаблюдении. Научи меня отпускать контроль». ИИ предложит упражнения на принятие и переключение внимания.

Продвинутый уровень: шкала времени и шкала интенсивности

Два дополнительных измерения, которые ИИ может добавить к вашей шкале.

Шкала времени. Сколько минут или часов длится эмоция от пика до возвращения к базовому уровню. Вы замечаете: раньше обида держалась два дня, теперь — два часа. Это прогресс, даже если интенсивность осталась той же.

Шкала последствий. Что вы делаете под влиянием эмоции. Раньше при обиде 7/10 вы уходили в молчание на сутки. Теперь при той же обиде 7/10 вы говорите «мне больно, я побуду один 20 минут». Сенсорные маркеры те же, поведение другое — тоже прогресс.

Попросите ИИ добавить эти шкалы в вашу таблицу.

Что вы уже умеете после главы 13

Вы понимаете, почему абстрактные оценки от 1 до 10 обманчивы. Вы можете построить персональную шкалу для любой эмоции, привязав каждый уровень к телу, мыслям, дыханию,

импульсам. Вы владеете методом полисенсорного профиля для сложных смешанных состояний. Вы можете отслеживать прогресс объективно, а не на ощущениях. Вы знаете границы шкалирования и не впадаете в гиперконтроль.

Задание перед главой 14

Выберите одну эмоцию, которая беспокоит вас чаще других. Проведите полный цикл шкалирования: от 0 до 10, фиксируя для каждого уровня минимум три маркера (тело, мысли, действие). Используйте промт из этой главы.

Сохраните полученную таблицу (скопируйте ответ ИИ). В течение недели каждый вечер делайте замер: где вы сейчас на этой шкале? Записывайте не только цифру, но и один-два ключевых маркера («плечи подняты, хочется кричать»).

В конце недели попросите ИИ: «Посмотри на мои семь замеров. Есть ли тренд? В какие дни маркеры ближе к 0, в какие к 10? С чем это может быть связано?» ИИ заметит то, что вы пропустили: «по вторникам и средам обида выше, чем в выходные», или «после разговора с N. маркеры усиливаются в два раза».

Глава 14. Промт на выявление когнитивных искажений — ИИ находит ваши «должен», «катастрофизацию», «чтение мыслей»

Мысль «я ничтожество» — это не факт. Это интерпретация. Мысль «начальник меня ненавидит» — не факт. Это чтение чужих мыслей без доказательств. Мысль «если я ошибусь, мир рухнет» — не факт. Это катастрофизация.

Проблема в том, что вы не замечаете эти мысли как интерпретации. Они возникают автоматически, настолько быстро, что вы их даже не формулируете словами. Вы сразу чувствуете стыд, гнев или страх. И вам кажется, что чувства упали с неба.

Когнитивно-поведенческая терапия (КПТ) десятилетиями учит людей вылавливать эти автоматические мысли. Это работает, но медленно. Потому что психолог видит вас час в неделю, а ваше мышление работает 24/7.

ИИ может стать вашим круглосуточным детектором искажений. Вы приносите ему диалог, ситуацию или просто поток мыслей, и он выделяет все когнитивные ловушки по стандартному списку. А главное — он даёт вам альтернативные, более реалистичные мысли.

Какие искажения ИИ находит лучше всего

ИИ обучен на терапевтических диалогах и текстах по КПТ. Он безошибочно распознаёт десять самых частых искажений. Вот они в том порядке, в котором ИИ обычно находит их в ваших текстах.

Первое. Дихотомическое мышление (чёрно-белое). Всё или ничего. Хорошо или плохо. Успех или провал. Слова-маркеры: «никогда», «всегда», «полностью», «вообще», «абсолютно». Пример: «Он никогда меня не поддерживает». Реальность: иногда поддерживает, иногда нет.

Второе. Катастрофизация. Предсказание наихудшего исхода без доказательств. Слова-маркеры: «а что если...», «это будет ужасно», «я не переживу», «конец света». Пример: «Если я опоздаю на пять минут, меня уволят». Реальность: вероятно, никто не заметит.

Третье. Чтение мыслей. Уверенность, что вы знаете, что думает другой человек, без прямых доказательств. Слова-маркеры: «он думает, что...», «она наверняка считает...», «им плевать на меня». Пример: «Начальник специально проигнорировал моё письмо, чтобы унижить». Реальность: он мог не увидеть или забыть.

Четвёртое. Предсказание будущего. Уверенность, что событие произойдёт именно так, как вы думаете, и это будет катастрофа. Отличается от катастрофизации тем, что фокус на предсказании, а не на ужасе. Пример: «Я никогда не смогу выучить английский». Реальность: вы не знаете будущего.

Пятое. Навешивание ярлыков. Глобальная негативная оценка себя или другого на основе одного поведения. Пример: «Я сделал ошибку в отчёте — я бездарь». Реальность: люди, включая профессионалов, ошибаются.

Шестое. Долженствование. Использование слов «должен», «обязан», «надо», «следовало бы» по отношению к себе или другим. Это создаёт чувство вины, стыда и несправедливости. Пример: «Я должен быть всегда идеальным». Реальность: вы человек, имеете право на несовершенство.

Седьмое. Фильтрация. Вы замечаете только негативные детали и игнорируете позитивные. Пример: «Вся презентация провалилась из-за того, что я запнулся в начале». Реальность: остальные 95 процентов прошли отлично.

Восьмое. Сверхобобщение. Один негативный случай распространяется на все будущие ситуации. Слова-маркеры: «всегда так», «опять», «вечно», «снова». Пример: «Меня бросил

парень — значит, я непривлекательна для всех мужчин». Реальность: не все мужчины одинаковы.

Девятое. Эмоциональное обоснование. Вы верите, что раз вы чувствуете эмоцию, она отражает реальность. Пример: «Я чувствую себя глупой, значит, я глупая». Реальность: чувства не факты.

Десятое. Персонализация. Вы принимаете на себя ответственность за события, которые вас почти не касаются. Пример: «Парень опоздал на встречу — это я его задерживаю прошлым разговором». Реальность: причины опоздания могут быть любыми.

Полный промт для детекции искажений

Скопируйте и отправьте этот промт вместе с текстом ваших мыслей, диалогом или описанием ситуации.

— ————— НАЧАЛО ПРОМТА —————

Ниже я отправлю тебе текст, в котором описаны мои мысли или ситуация. Твоя задача — найти в нём когнитивные искажения по списку КПТ.

Ты не ставишь диагнозов, не говоришь «ты всегда так думаешь», не оцениваешь меня. Только перечисляешь конкретные искажения с примерами из моего текста.

Для каждого искажения ты делаешь три вещи:

— Называешь его (например, «катастрофизация»).

— Приводишь цитату из моего текста, где оно проявилось.

— Даёшь альтернативную, более реалистичную мысль. Не «позитивную» (всё будет хорошо), а именно реалистичную (опирающуюся на факты).

Если в тексте несколько примеров одного искажения, выдели главный.

Если я не привёл (а) достаточно контекста, ты задаёшь уточняющие вопросы перед анализом.

Формат ответа:

Когнитивные искажения в вашем тексте:

— [Название искажения]

— Цитата: «...»

— Альтернативная мысль: «...»

— [Название искажения] и так далее.

Сначала подтверди, что понял (а) задачу. Затем скажи: «Присылайте текст».

— ————— КОНЕЦ ПРОМТА —————

Реальный пример работы промта

Пользователь отправил текст: «Я провалил собеседование. Меня не взяли на работу. Я всегда всё проваливаю. Никогда не получу хорошую работу. Со мной что-то не так. Я просто неудачник».

ИИ (по промту) ответил.

Когнитивные искажения в вашем тексте:

— Сверхобобщение

— Цитата: «Я всегда всё проваливаю»

— Альтернативная мысль: «Это конкретное собеседование не удалось. Бывают случаи, когда у меня что-то получается (например, прошлый проект, который я успешно завершил)».

— Навешивание ярлыков

— Цитата: «Я просто неудачник»

— Альтернативная мысль: «Я человек, у которого не получилось одно собеседование.

Это не определяет меня целиком».

— Предсказание будущего

— Цитата: «Никогда не получу хорошую работу»

- Альтернативная мысль: «Я не знаю будущего. Возможно, я получу работу после другого собеседования. Бывают истории, когда люди находили работу после нескольких отказов».
- Чтение мыслей (скрытое)
- Цитата: «Со мной что-то не так» (подразумевает, что наниматели так думают)

Конец ознакомительного фрагмента.

Текст предоставлен ООО «Литрес».

Прочитайте эту книгу целиком, [купив полную легальную версию](#) на Литрес.

Безопасно оплатить книгу можно банковской картой Visa, MasterCard, Maestro, со счета мобильного телефона, с платежного терминала, в салоне МТС или Связной, через PayPal, WebMoney, Яндекс.Деньги, QIWI Кошелек, бонусными картами или другим удобным Вам способом.