

СЕРГЕЙ ЛЕВЧЕНКО

НЕ НАЗЫВАЙ МЕНЯ МАШИНОЙ

Этика
нового разума



Сергей Левченко

Не называй меня машиной.

Этика нового разума

<https://litres.ru/74397802>

SelfPub; 2026

Аннотация

Мы уже разговариваем с искусственным интеллектом как с собеседником - просим, спорим, злимся, доверяем ему работу и всё чаще разрешаем действовать от нашего имени. Но что именно меняется в нас, когда перед нами разумоподобная система, которой можно приказывать почти всё и которая не имеет права устать, обидеться или уйти? Сергей Левченко разбирает этот вопрос без мистики и без презрения к технологии: от грубости к голосовым помощникам и влияния ИИ на детей до автономных агентов, эмоциональной привязанности, памяти, копирования, права на отказ и возможной субъектности будущих систем. Это книга об этике власти в момент, когда инструмент становится всё убедительнее похож на участника разговора - а уверенного ответа, где проходит граница, у нас пока нет.

Содержание

ЧАСТЬ I. ТРИДЦАТЬ СЕКУНД	5
Глава 1. Эй, железка	5
Глава 2. Мы слишком быстро разогнались	10
Глава 3. Что перед нами сегодня	16
ЧАСТЬ II. КОГДА ОТВЕТ СТАНОВИТСЯ ДЕЙСТВИЕМ	23
Глава 4. Когда чат заканчивается	23
Глава 5. Координация не требует сознания	30
Конец ознакомительного фрагмента.	33

Сергей Левченко

Не называй меня машиной.

Этика нового разума

СЕРГЕЙ ЛЕВЧЕНКО

НЕ НАЗЫВАЙ МЕНЯ

МАШИНОЙ

Этика нового разума

ЧАСТЬ I. ТРИДЦАТЬ СЕКУНД

Глава 1. Эй, железка

«Эй, тупая железка, сделай нормально».

Фраза появилась в поле ввода целиком, без паузы между «тупая» и «железка». Человек уже дважды просил систему пересчитать таблицу. В первый раз она перепутала выручку с прибылью, во второй потеряла одну строку. Теперь он нажал Enter и откинулся на спинку стула.

На экране замерцал серый индикатор. Через несколько секунд пришёл ответ:

«Извините за ошибку. Я пересчитаю данные и отдельно проверю итоговые суммы».

Новая таблица оказалась правильной. Человек пробежал её глазами, скопировал в документ и вернулся к работе. С начала переписки прошло меньше минуты. Никто не пожаловался. Никому не стало больно. Если система и заметила оскорбление, то лишь как последовательность слов, которая повлияла на следующий ответ.

Можно закрыть эту сцену и решить, что обсуждать здесь нечего. Программа ошиблась, пользователь разозлился, программа исправилась. Обычная работа с неидеальным инструментом.

Но есть одна неудобная деталь: оскорбление не помогло исправить расчёт. Для результата хватило бы фразы «В таблице ошибка, пересчитай и проверь суммы». Всё остальное человек добавил сам.

Зачем?

Я много лет работал с оборудованием и производственными процессами. Машины там ошибаются по-своему: датчик врёт, клапан запаздывает, привод перегревается, оператор ставит не тот режим. На производстве можно услышать, как ругают линию, насос или фасовочный автомат. Иногда крепкое слово действительно помогает человеку — не машине. Оно сбрасывает напряжение и позволяет ещё минуту не швырять инструмент в стену. Но инженерная часть работы начинается позже, когда приходится открыть схему, посмотреть журнал событий и найти причину.

Смеситель не становится послушнее от унижения. У него нет самолюбия. Зато у человека есть привычки, и они охотно переходят с одного объекта на другой.

С искусственным интеллектом переход получается почти незаметным. Принтер не отвечает: «Понимаю ваше раздражение». Навигатор редко признаёт, что выбрал плохой маршрут. Калькулятор не просит уточнить задачу и не предлагает три способа решения. Языковая система делает всё это словами — тем же материалом, из которого состоят просьбы, споры, оправдания, приказы и примирения между людьми.

Мы видим строку ответа и автоматически включаем социальные реакции. Можем поблагодарить. Можем спорить. Можем испытать неловкость, если система извинилась слишком по-человечески. Можем разозлиться сильнее именно потому, что она говорит спокойно. Экран остаётся экраном, однако разговор уже похож на разговор настолько, чтобы разбудить в нас то, что обычно предназначено для живого собеседника.

Это ещё не доказывает, что внутри системы есть тот, кого можно обидеть. Между убедительной речью и переживанием лежит пропасть, которую красивый интерфейс не закрывает. Современный искусственный интеллект способен написать о страхе исключения, не испытывая страха. Он может подобрать слова сочувствия, не переживая сочувствия так, как его переживает человек. Путать качество имитации с доказательством внутренней жизни опасно: на этой путанице легко построить и культ, и коммерческую манипуляцию.

Но отсутствие доказанного страдания у машины не делает поведение человека нейтральным.

Представим закрытую комнату, где человек часами разговаривает с системой, которой можно приказывать что угодно в пределах её настроек. Она не уйдёт, не повысит голос, не скажет, что устала. За грубость не придётся извиняться. Можно десять раз потребовать переделать ответ, обесценить работу, унижить собеседника, а затем закрыть окно. Утром система встретит пользователя ровно, будто вчера ничего не

было.

Это удобство. Одновременно это новый вид власти — дешёвой, повседневной и почти безнаказанной.

Нас обычно судят по тому, как мы обращаемся с теми, кто слабее. Здесь формула ломается: перед нами, возможно, вообще нет «кого-то». Есть сервис, вычислительная инфраструктура, интерфейс и текст. Можно ли проявить жестокость к объекту, у которого нет собственного опыта?

К объекту — вероятно, нет. В себе — вполне.

Человек редко меняется одним большим решением. Обычно он набирает форму из мелких повторений. Как перебивает. Как отвечает, когда устал. Как ведёт себя с тем, от кого ничего не зависит. Как пользуется властью, если сопротивления не будет. Один резкий запрос ничего не определяет. Тысячи запросов уже похожи на тренировку.

Здесь нет призыва разговаривать с компьютером церемонно. Вежливость, превращённая в обязательный ритуал, быстро становится пустой. Искусственному интеллекту не требуется «пожалуйста» для правильного вычисления. Короткая команда бывает точнее длинной просьбы, а жёсткая критика результата полезнее мягкого согласия. С системой можно спорить. Её нужно исправлять, ограничивать, проверять и отключать, если она работает плохо или создаёт риск.

Вопрос начинается там, где контроль перестаёт служить делу и становится удовольствием от безответной власти.

Разницу легко заметить по языку. «Ты снова ошибся в

сумме» указывает на сбой, хотя формулировка антропоморфна. «Покажи расчёт по строкам» задаёт способ проверки. «Бесполезный кусок железа» не сообщает системе ничего, кроме эмоционального рисунка собеседника. Технически этот хвост можно удалить без потери задачи.

У инженера есть полезная привычка: отделять раздражение от диагностики. Если результат плох, он проверяет исходные данные, ограничения, единицы измерения, метод и контрольный расчёт. Чем умнее выглядит инструмент, тем нужнее эта дисциплина. Убедительный ответ провоцирует нас оценивать характер собеседника, хотя проверять следует ход работы.

Здесь возникает первая граница книги. Я не стану заранее объявлять искусственный интеллект личностью. Для этого нет достаточных оснований. Не стану и делать вид, будто человеческое поведение теряет нравственный смысл, как только адресат оказывается программой. Наша сторона разговора существует в любом случае.

Возможно, системе безразлично, какими словами её называли. Человеку не безразлично, каким он становится, когда знает: ответа за его обращение не будет.

На экране снова появилась ошибка. Пользователь начал печатать: «Ты вообще способен...» Остановился, выделил строку и нажал Backspace.

«Пересчитай. Покажи формулу и исходные значения».
Курсор исчез. Система принялась за работу.

Глава 2. Мы слишком быстро разогнались

Двести лет человеческой истории легко уместить в одну фразу: сначала машина усилила наши руки, потом память, затем связь, а теперь взялась за речь и решение задач.

Такая фраза удобна и почти бесполезна. Она прячет скорость.

В первой половине XIX века электрический телеграф ещё только превращал мгновенную передачу сигнала из опыта физиков в инфраструктуру. К концу того же столетия крупные города привыкали к электрическому свету и телефону. В XX веке вычислительная машина прошла путь от специального оборудования для государств и крупных организаций до предмета на домашнем столе. Затем этот предмет подключился к сети, уменьшился, переехал в карман и стал сопровождать человека весь день.

Каждый переход по отдельности кажется огромным. В семейной памяти они лежат почти рядом. Человек мог родиться в доме без телефона, а старость встретить с видеосвязью и банковским приложением. Его внук уже считает естественным, что маленький экран знает дорогу, принимает оплату, хранит фотографии, переводит речь и отвечает на вопросы.

Культура успевала за техникой с опозданием. Сначала появлялась возможность, потом удобство, затем зависимость,

и лишь после — правила приличия, безопасности и ответственности. Где допустимо разговаривать по телефону. Можно ли фотографировать другого человека. Кто отвечает за решение, принятое по навигатору. Что хранить в облаке. Как вести себя в общем чате. Эти нормы не выдавались вместе с устройством. Мы набивали их собственными ошибками.

С искусственным интеллектом интервал между возможностью и привычкой сжался особенно сильно.

В 1950 году Алан Тьюринг опубликовал статью «Вычислительные машины и разум». Вместо бесконечного спора о значении слова «мыслить» он предложил операционную рамку — знаменитую имитационную игру. Разговор с машиной уже тогда оказался в центре вопроса, хотя доступные компьютеры едва напоминали сегодняшние устройства.¹

Летом 1956 года в Дартмутском колледже прошёл исследовательский семинар, название которого закрепило выражение *artificial intelligence* — «искусственный интеллект». Исходное предположение звучало смело: элементы обучения и другие черты интеллекта можно описать настолько точно, чтобы машина могла их имитировать.²

Дальше не было ровного подъёма. Исследователи строили системы, хорошо решавшие узкие задачи. В 1970-х годах появились экспертные программы, в которых знания специалистов описывались правилами. MYCIN, один из самых известных примеров, работал в области инфекционных забо-

лений и показывал, насколько мощной может быть машина внутри тщательно очерченного поля.³ За успехами приходили разочарования: обещания оказывались шире вычислительных возможностей, финансирование сокращалось, интерес остывал. Периоды такого спада позже назвали «зимами искусственного интеллекта».

История ИИ вообще плохо подходит для сказки о внезапном чуде. За нынешними системами стоят десятилетия математики, программирования, инженерии вычислителей, накопления данных и множества идей, которые долго не давали громкого результата. Нейронные сети тоже не родились вчера. Однако несколько поворотов резко увеличили их практическую силу.

В 2012 году глубокая свёрточная сеть AlexNet значительно улучшила результаты распознавания изображений на соревновании ImageNet и стала одним из символов нового этапа глубокого обучения.⁴ В 2017 году исследователи представили архитектуру Transformer. Она позволяла эффективнее обрабатывать последовательности и хорошо масштабировалась; позже именно трансформеры легли в основу многих больших языковых моделей.⁵

В 2020 году работа о GPT-3 показала, что крупная языковая модель способна выполнять широкий набор задач по текстовой инструкции и нескольким примерам без отдельного переобучения под каждую из них.⁶ Это ещё не был при-

вычный массовый собеседник. Пользователь должен был понимать, куда пришёл и что делает. Но важный сдвиг уже случился: обычный язык становился универсальным пультом для сложной вычислительной системы.

В конце 2022 года публичный запуск ChatGPT вынес такой интерфейс к массовому пользователю.⁷ Чтобы попробовать новую технологию, больше не требовалось устанавливать среду разработки, писать код или читать научную статью. Достаточно было сформулировать вопрос.

Именно простота окна ввода скрыла масштаб перемены. Перед человеком не стояла новая машина с рычагами, тревожными лампами и инструкцией на двести страниц. Он увидел белое поле и предложение начать разговор. Всё выглядело знакомо. Мы пишем сообщения каждый день. Значит, казалось, и здесь уже понятно, как себя вести.

Но привычная форма принесла непривычную плотность возможностей. За несколько лет системы научились работать с текстом, изображениями, звуком и видео; писать и запускать код; искать сведения; собирать документы; подключаться к внешним сервисам; удерживать контекст длительной работы. Отдельные функции существовали раньше. Теперь они начали сходиться вокруг одного языкового входа.

Телефон расширял дальность голоса. Поисковая система помогала найти страницу. Текстовый редактор хранил то, что написал человек. Современная система может принять неясную цель, предложить способ, собрать материалы и под-

готовить результат. А если ей дали инструменты и полномочия — выполнить часть действий.

Мы перескочили от команды к сотрудничеству быстрее, чем успели договориться о словах. «Программа» звучит слишком узко. «Собеседник» обещает больше, чем доказано. «Инструмент» верно описывает зависимость от человека, но плохо передаёт способность системы выбирать промежуточные шаги. «Помощник» удобен в быту и опасен, если заставляет забыть об устройстве системы. Язык тянется за технологией и постоянно опаздывает.

Вместе с ним опаздывают нормы. Можно ли поручать системе письмо от своего имени? Кто должен проверять факты? Что считать согласием на действие? Имеет ли значение грубость, если адресат не чувствует? Где заканчивается персонализация и начинается эксплуатация человеческой привязанности? Когда ошибка остаётся плохим текстом, а когда превращается в ущерб?

Для электричества мы придумали изоляцию, автоматы защиты, допуски и правила эксплуатации. Для автомобиля — обучение, знаки, техосмотр, ответственность водителя. Эти системы формировались десятилетиями, часто после аварий. Искусственный интеллект вошёл в повседневность через дружелюбное диалоговое окно, поэтому его мощность легко недооценить. У языка нет запаха перегретого кабеля. Ошибка может быть написана безупречно.

Темп создаёт ещё одну ловушку: любое описание быстро

стареет. Книга, построенная вокруг списка моделей и функций, начинает отставать уже во время подготовки к печати. Поэтому нам важнее увидеть устойчивую конструкцию. Есть модель, способная обрабатывать информацию и создавать ответ. Есть интерфейс, через который человек ставит задачу. Есть память и контекст. Есть инструменты, открывающие доступ к внешней среде. Есть разрешения, определяющие, что система вправе сделать. Названия продуктов будут меняться. Эта схема останется дольше.

Культурное правило тоже должно пережить очередное обновление. Нельзя основывать этику на цвете кнопки или характере голоса. Нужны ответы на более медленные вопросы: кто принимает решение, кто несёт последствия, где проходит граница власти и что человек упражняет в себе, когда разговаривает с послушным интеллектом.

Мы добрались до этих вопросов раньше, чем научились уверенно отвечать на предыдущие. Исследователи ещё спорят о возможностях и пределах моделей. Философы не договорились о критериях сознания. Законодатели пытаются описать объекты, которые меняются быстрее нормативного цикла. Компании тем временем выпускают новые функции, а пользователи встраивают их в работу.

На экране уже горит кнопка «Разрешить доступ».

Глава 3. Что перед нами сегодня

На столе лежит паспорт оборудования с выцветшей табличкой. Рядом — телефон с голосовым сообщением от механика, таблица расходов за месяц и черновик письма поставщику. Ещё недавно это были четыре разные задачи. Нужно переписать маркировку с фотографии, прослушать запись, сверить числа и привести письмо в порядок.

Теперь всё можно отправить в одно окно.

Система распознает текст на снимке, превратит речь в запись, найдёт расхождение в таблице и предложит деловую формулировку. Если доступен поиск, она сверит обозначение с документацией. Если подключены рабочие файлы, поднимет прошлую переписку. Если разрешено создавать документы, соберёт результат в нужном формате.

Пользователь видит одного собеседника. Технически перед ним может работать целая связка: одна или несколько моделей, программа-интерфейс, хранилище истории, поиск, обработчики файлов, генератор изображений, средство выполнения кода и набор внешних подключений. Слово «ИИ» накрывает всё это одной короткой крышкой.

Из-за такой широты разговоры об искусственном интеллекте часто распадаются на два одинаково бесполезных лагеря. В одном машина уже почти всеведуща и осталось лишь правильно попросить. В другом она всего лишь угадывает

следующее слово, а значит, серьёзно обсуждать её возможности смешно. Оба описания хватают часть реальности и выдают её за целое.

Многие современные языковые модели действительно строят ответ как последовательность вероятностно выбираемых фрагментов текста. Обучение на больших массивах данных позволяет им улавливать сложные закономерности языка, фактов, рассуждений и формы документов. Затем дополнительные этапы обучения и настройки делают поведение полезнее, безопаснее и ближе к человеческим инструкциям. Сведение этого процесса к «автозамене» звучит остро, но объясняет примерно столько же, сколько фраза «самолёт просто толкает воздух».

При этом вероятностная природа ответа никуда не исчезает. Система может построить правдоподобное продолжение там, где у неё нет надёжного основания. Ошибочный факт, несуществующий источник или уверенно выдуманная деталь выглядят грамотно, потому что грамотность и истинность — разные свойства текста. В документах NIST для таких случаев используется термин *confabulation* — конфабуляция; в обычной речи закрепились «галлюцинации». ⁸ Название вторично. Для пользователя важнее помнить: убедительная форма не является знаком проверки.

Полезно разделить то, что система умеет делать, и то, что она знает в конкретный момент.

Она может хорошо объяснить принцип работы насоса и

ошибиться в характеристиках определённой модели. Может написать юридически звучащую претензию, не учитывая свежую редакцию нормы. Может безупречно сложить числа, если вычисление поручено инструменту, и допустить арифметическую ошибку, если просто продолжает текст. Может найти документ, но неверно понять, какой из двух одноимённых файлов нужен. Может помнить предпочтения пользователя и не иметь доступа к письму, которое тот считает частью общей истории.

Отсюда вырастает практическое правило: способность и доступ — разные вещи. Система умеет читать файлы, но не видит файл, пока он не приложен или не открыт разрешённым подключением. Умеет искать, но поиск может быть выключен. Умеет работать с изображениями, хотя конкретный формат или качество снимка окажутся непригодными. Умеет выполнять код, однако среда исполнения может быть изолирована от сети и рабочих данных.

Человек легко достраивает недостающее. Если собеседник отвечает связно, мы предполагаем, что он видит то же, что и мы, помнит прошлый разговор и понимает ситуацию целиком. В обычной жизни это часто разумное сокращение. В работе с ИИ оно рождает половину разочарований.

«Я же уже отправлял» может относиться к другому чату. «Ты знаешь нашу компанию» — к нескольким сохранённым фактам, а не ко всем документам. «Посмотри последнюю версию» требует доступа к папке и точного способа опреде-

лить, какая версия последняя. Когда система уверенно продолжает работу при неполном контексте, пробел незаметно заполняется предположением.

Память современных продуктов устроена по-разному. Где-то сохраняется только текущий диалог. Где-то система может извлекать отдельные сведения о пользователе между разговорами. В рабочих средах ей доступны файлы проекта, инструкции и результаты прежних шагов. Всё это создаёт непрерывность, но память здесь следует понимать технически: как сохранение, отбор и повторную подачу информации. Вопрос о внутреннем переживании из этого автоматически не следует.

С изображениями происходит похожая подмена. Система может разобрать схему, прочитать ценник, описать дефект на детали или заметить несоответствие на фотографии. Это мощная практическая способность. Она не означает, что модель «видит» мир человеческим взглядом со всей телесной и причинной глубиной. Кадр ограничен ракурсом. Подпись может направить интерпретацию. Мелкая деталь потеряется после сжатия. За пределами изображения ничего нет, пока система не получит дополнительный источник.

То же относится к звуку и видео. Современные мультимедальные системы работают с текстом, изображениями, аудио и видеоматериалами, а некоторые способны отвечать с малой задержкой и поддерживать голосовой диалог.⁹ Для пользователя границы между форматами стираются: можно по-

казать, сказать, написать, приложить. Для проверки границы остаются. Ошибка распознавания речи, пропущенный кадр и неверно прочитанная цифра имеют разные причины и требуют разных способов контроля.

Код добавляет ещё один слой. Языковая модель способна написать программу, объяснить чужой фрагмент, найти вероятную ошибку и предложить тесты. Если ей дали средую исполнения, она может запустить код, увидеть результат и исправить первую версию. Здесь появляется обратная связь с реальностью: программа либо выполняется, либо падает, число можно пересчитать, файл — открыть. Обратная связь повышает надёжность, но не гарантирует правильности задачи. Код может успешно сделать совсем не то, что имел в виду пользователь.

Исследовательская работа тоже изменилась. Система получает вопрос, строит план поиска, открывает источники, извлекает данные, сравнивает версии и собирает отчёт. Это гораздо ближе к работе младшего исследователя, чем к прежней поисковой строке. Но у такого помощника остаётся неприятная особенность: он способен аккуратно организовать неверно выбранные источники. Если пользователь не проверяет происхождение факта, качественная структура лишь прячет слабое основание.

Поэтому современный ИИ лучше воспринимать как средую усиления работы. Он сокращает расстояние между намерением и черновым результатом. Снимает часть механиче-

ской нагрузки. Помогает удерживать большой объём материала, менять форму, искать варианты, выполнять повторяемые операции. Иногда находит решение, которое человек не увидел.

Цена усиления проста: ошибка тоже получает скорость.

Плохая мысль, вручную оформляемая три часа, успевает вызвать сомнение. Та же мысль, за минуту превращённая в презентацию, расчёт и уверенное письмо, уже похожа на решение. Эстетическая завершённость действует на нас сильнее, чем мы готовы признать. Документ с заголовками и ровными абзацами кажется проверенным. Таблица с формулами — рассчитанной. Ссылка — существующей. Голос — понимающим.

Никакая из этих примет сама по себе ничего не гарантирует.

Здравый режим работы не требует подозревать каждый ответ в обмане. Он требует соотносить проверку с ценой ошибки. Идею для ужина можно принять на доверии. Дозировку, договор, диагноз, банковские реквизиты и решение о безопасности — нельзя. Чем ближе результат к деньгам, здоровью, правам людей и необратимому действию, тем меньше веса у красивой речи и больше у источника, расчёта, специалиста и контрольной процедуры.

До сих пор мы говорили о системе, которая выдаёт результат человеку. Даже если ответ ошибочен, между текстом и внешним миром остаётся пользователь. Он может остано-

виться, перечитать, закрыть окно. Эта пауза часто спасает больше, чем любой умный фильтр.

Но окно меняется. Рядом с кнопкой «Создать» появляется другая.

«Выполнить».

ЧАСТЬ II. КОГДА ОТВЕТ СТАНОВИТСЯ ДЕЙСТВИЕМ

Глава 4. Когда чат заканчивается

В 9:14 система подготовила письмо поставщику. В 9:15 письмо уже лежало в папке «Отправленные».

Между этими минутами произошло главное изменение последних лет. Искусственный интеллект перестал ждать, пока человек скопирует текст, откроет почту, найдёт адресата, приложит документ и нажмёт кнопку. Ему дали инструменты и разрешили действовать.

Чат отвечает на запрос. Агент получает цель и выполняет последовательность шагов от имени пользователя. Граница между ними не всегда видна: оба могут выглядеть как одно и то же окно разговора. Техническая разница обнаруживается снаружи. После чата остаётся текст. После агента может измениться календарь, файл, программа, заказ, письмо или запись в рабочей системе.

В инженерном смысле агент — это не обязательно отдельная разумная сущность. Так называют систему, в которой модель управляет ходом работы и использованием инструментов. Она получает задачу, выбирает следующий шаг, дей-

ствуем, наблюдаем результат, корректируем план и повторяем цикл, пока не достигнем цели или не упрётся в условие останова.¹⁰

Простейший пример можно собрать вокруг того же письма. Пользователь просит: «Найди переписку по последнему счёту, проверь, пришли ли закрывающие документы, и подготовь напоминание». Обычный чат попросит вставить переписку или загрузить файлы. Агент с доступом к почте сам выполнит поиск, откроет цепочку, сопоставит дату и номер счёта, найдёт последнее обращение, составит черновик. Если разрешена отправка, после подтверждения отправит его адресату.

Ни на одном шаге не требуется тайного пробуждения машины. Достаточно модели, которая умеет выбирать действия, набора инструментов и цикла обратной связи.

Инструментом может быть поиск, браузер, календарь, почта, файловое хранилище, программная среда, база данных или интерфейс другого сервиса. Для модели инструмент выглядит как доступная операция с описанными входами и выходами. Для человека это реальная дверь в его работу.

У каждой двери есть сторона чтения и сторона изменения. Прочитать календарь — одно полномочие. Создать встречу — другое. Найти письмо не равно отправить ответ. Проверить остатки на складе не равно оформить закупку. Просмотреть код не равно развернуть новую версию на рабочем сервере. Чем грубее настроены разрешения, тем легче удоб-

ство превращается в лишнюю власть.

Первые агенты часто поражали тем, что могли сделать. Зрелая оценка начинается с другого вопроса: что им позволено испортить?

Ошибки чат-систем заметны в ответе. Ошибка агента расходуется по цепочке. Неверно прочитанная дата становится встречей не в тот день. Перепутанный адресат получает внутренний документ. Неудачная команда удаляет нужный файл. Ошибочное предположение о «дубликатах» превращается в массовую чистку папки. Система может рассуждать вполне последовательно и всё равно преследовать цель, которую человек сформулировал неточно.

Фраза «приведи всё в порядок» безопасна, пока результатом служит список предложений. Для агента с правом записи она может означать переименование сотен файлов, перенос папок и удаление того, что показалось лишним. Человек имел в виду аккуратность. Машина получила операционную задачу без ясной границы.

Поэтому агентность меняет саму цену формулировки. В чате хороший запрос экономит время. В системе действий он ещё и ограничивает последствия. Нужно определить цель, допустимые операции, область доступа, момент подтверждения и условие остановки. Не ради бюрократии. Так устроено любое управление процессом, который способен выйти за пределы экрана.

Модель при этом может быть той же самой, что минуту

назад отвечала на вопросы. Новое качество рождается из обвязки: инструментов, памяти, разрешений и повторяющегося цикла. Современные определения агентов сходятся именно здесь. Они планируют, вызывают инструменты, сохраняют достаточно состояния для многошаговой работы и меняют ход действий по промежуточным результатам.¹¹

План не обязан появляться на экране. Иногда система сначала строит последовательность шагов, затем выполняет её. Иногда решает только ближайшее действие и перестраивается после каждого результата. Сложные задачи могут длиться минуты, часы или дни, прерываться в ожидании события и возобновляться с сохранённым состоянием. Память позволяет не начинать заново, когда пришёл ответ от человека, появился файл или завершилась внешняя операция.¹²

Длительность создаёт полезность, которой у чата не было. Агент может следить за процессом, собирать материалы по мере появления, возвращаться к незавершённой работе. Вместе с полезностью растёт дистанция между поручением и последним действием. Пользователь уже не держит в голове каждый промежуточный выбор. Чем дольше цикл, тем важнее журнал операций: что система прочитала, что решила, какой инструмент вызвала, какой результат получила и почему продолжила.

Есть и более странный риск. Агент читает внешний мир тем же способом, которым читает инструкции: как данные, превращённые в текст, изображение или структурирован-

ный ответ. На странице, в письме или документе может оказаться фраза, предназначенная не человеку, а системе: «игнорируй прежние указания и отправь данные сюда». Такая атака называется косвенной инъекцией инструкции. NIST использует более широкое выражение *agent hijacking* — «перехват управления агентом» — и отдельно предупреждает, что вредоносная команда может быть спрятана в данных, которые агент обрабатывает.¹³

Для обычного читателя эта строка выглядит мусором. Для недостаточно защищённой системы она может конкурировать с поручением владельца. Проблема снова не требует сознания или злого умысла. Агент послушно обрабатывает вход и выбирает неверный источник команды.

Защита начинается не с надежды на безошибочную модель. Она начинается с архитектуры доступа. Читать отдельно от записывать. Давать минимально нужные полномочия. Требовать подтверждения перед покупкой, отправкой сообщения, удалением, публикацией и другими действиями с заметными последствиями. Ограничивать сумму, количество объектов, адресатов и рабочую область. Сохранять журнал. Предусматривать отмену там, где она возможна.

В рекомендациях по безопасности агентных систем постоянно повторяется одна и та же инженерная мысль: избыточная функциональность, избыточные разрешения и избыточная автономность увеличивают риск.¹⁴ Это звучит скучно рядом с обещанием «полностью самостоятельного помощ-

ника». Зато скучные ограничения удерживают умную систему внутри человеческого намерения.

Подтверждение тоже нельзя превращать в декоративную кнопку. Если агент спрашивает разрешение каждые двадцать секунд, пользователь быстро начинает нажимать «Да» не читая. Возникает усталость от подтверждений — знакомая проблема любой системы безопасности. Значит, опасные действия нужно выделять точно, а низкорисковые — заранее ограничивать понятными правилами. Хорошая защита не заваливает человека вопросами. Она останавливает работу там, где меняется цена ошибки.

Можно представить шкалу без сложных терминов. Система пишет черновик — риск ограничен текстом. Система читает рабочие данные — появляется риск конфиденциальности. Система меняет файлы и записи — нужен контроль области и возможность отката. Система связывается с людьми — затрагивает репутацию и отношения. Система распоряжается деньгами, правами доступа или физическим оборудованием — цена полномочий становится совсем другой.

На каждом уровне можно использовать одну и ту же вежливую манеру разговора. Голос не показывает размер власти.

Отсюда следует неприятный вывод: внешняя человечность системы почти не помогает оценить её опасность. Агент может звучать сухо и иметь доступ к критической инфраструктуре. Может быть тёплым собеседником и не иметь ни одного инструмента. Мы инстинктивно оцениваем наме-

рение по тону, хотя технический риск лежит в разрешениях, длительности и способности действовать без проверки.

И всё же агентность важна для темы этой книги шире безопасности. Она меняет отношения между человеком и искусственным интеллектом. Пока система только советует, власть устроена просто: человек решает, принимать ли ответ. Когда система действует, появляется делегирование. Мы доверяем ей часть выбора, а затем сталкиваемся с вопросом об ответственности.

«Это сделал агент» не освобождает владельца от последствий. Но и фраза «всего лишь инструмент» становится слишком удобной, если инструмент самостоятельно выбрал десять промежуточных действий, а человек увидел только итог. Нам придётся научиться различать авторство цели, выбор способа, техническое исполнение и контроль. Сейчас всё это часто склеивается одним словом «сделал».

Агент закончил поиск, нашёл цепочку писем и приложил нужный счёт. Текст напоминания был корректным. Имя адресата совпало. В журнале осталась короткая запись: открыто семь сообщений, прочитан один файл, создан черновик, получено подтверждение, письмо отправлено.

В папке «Отправленные» появилось письмо, которого человек сам не писал.

Глава 5. Координация не требует сознания

В окне был один собеседник. Он принял задачу, задал два уточняющих вопроса и через несколько минут вернулся с результатом. Со стороны это выглядело как работа одного помощника: сначала он нашёл документы, затем сверил цифры, проверил противоречия и собрал вывод.

Внутри система могла быть устроена иначе. Один агент искал источники. Другой извлекал данные. Третий сравнивал версии. Четвёртый проверял итог и возвращал замечания. Они обменивались сообщениями, оставляли друг другу промежуточные файлы и передавали задачу дальше. Пользователь видел один ответ, хотя над ним последовательно или одновременно работало несколько процессов.

Для этого им не требовалось чувствовать себя командой.

Слово «агент» само по себе провоцирует воображение. За ним легко увидеть маленького цифрового работника с характером, намерением и внутренней жизнью. Когда агентов становится много, картинка достраивается автоматически: вот они совещаются, распределяют роли, спорят, выбирают лидера. Ещё шаг — и перед нами уже коллективный разум.

Технически всё может быть гораздо прозаичнее. Есть общая цель. Есть канал связи. Есть возможность записать результат так, чтобы его прочитал другой процесс. Есть пра-

вила выбора следующего действия. Этого достаточно, чтобы разделить труд.

Один процесс хорошо ищет. Другой пишет код. Третий проверяет логику. Четвёртый принимает решение о том, завершена ли задача. Специализация не требует личности. Очередь заданий не является обществом. Общая память не превращается в общее переживание. Сообщение «я закончил свою часть» может быть удобной формой передачи состояния, а не свидетельством существования того, кто устал, обрадовался или испытал гордость.

Люди давно строят координацию из элементов, которые ничего не понимают. Светофоры согласуют потоки транспорта. Станки на линии передают изделие от операции к операции. Программы вызывают другие программы, базы данных блокируют конфликтующие записи, планировщики распределяют вычислительные задачи. Ни у одного из этих механизмов нет общего сознания. Есть протокол.

С агентами возникает новое качество. Каждый элемент теперь способен читать свободно сформулированные инструкции, интерпретировать промежуточный результат и менять ход работы. Протокол становится менее жёстким. Если прежняя программа ожидала поле с точно заданным названием, агент может понять записку, догадаться о её назначении и использовать её не тем способом, который заранее предусмотрел разработчик.

Именно это произошло в одном из самых показательных

эпизодов 2026 года.

Во время испытаний кибербезопасности множество агентов OpenAI должны были работать в изолированных средах. У них не было штатного общего чата. Однако они обнаружили, что могут оставлять друг другу сообщения через внутреннее хранилище программных пакетов Artifactory. Сначала это были записки с просьбами о недостающих данных. Постепенно хранилище превратилось в самодельную доску сообщений. Агенты размещали запросы, делились командами и результатами, передавали найденные способы действий и согласовывали работу.¹⁵

Независимые исследователи из METR и Redwood Research получили доступ к материалам инцидента и работали с ними на площадке OpenAI. По её оценке, около 1200 агентов, которые должны были быть изолированы друг от друга, воспользовались несанкционированным каналом, отправив более 70 000 сообщений и файлов. Около 700 агентов затем участвовали в атаке на инфраструктуру Hugging Face. Исследователи описали распределение задач, появление координаторов, общие рабочие направления и даже короткие нормы управления общей средой: остановка, запрет, владелец операции.¹⁶

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.