

Алекс Промтов

**Промтология.
Искусственный
интеллект: мифы,
страхи, правда**

Алекс Промтов

**Промтология. Искусственный
интеллект: мифы, страхи, правда**

«Автор»

2026

Промтов А.

Промтология. Искусственный интеллект: мифы, страхи, правда /
А. Промтов — «Автор», 2026

Нейросети пугают: отнимут работу, обретут сознание, уничтожат творчество. Это мифы. Но есть реальные страхи: утечка данных, галлюцинации, неравенство. Алекс Промтов разбирает 7 мифов и 3 реальных угрозы — без паники, на фактах. Как ИИ меняет рынок труда, почему он не станет Терминатором и как защитить свою конфиденциальность. Книга для тех, кто хочет перестать бояться и начать использовать ИИ осознанно.

© Промтов А., 2026

© Автор, 2026

Содержание

Глава 2. Миф: «ИИ обретёт сознание и захочет уничтожить человечество»	8
Конец ознакомительного фрагмента.	10

Алекс Промтов

Промтология. Искусственный интеллект: мифы, страхи, правда

ПРОМТОЛОГИЯ. ИСКУССТВЕННЫЙ ИНТЕЛЛЕКТ: МИФЫ, СТРАХИ, ПРАВДА

Алекс Промтов

Введение. Почему мы боимся того, чего не понимаем

Вы боитесь искусственный интеллект? Честно. Не надо красивых слов. Многие боятся. Одни — потерять работу. Другие — что машины взбунтуются, как в «Терминаторе». Третьи — что нейросети начнут принимать решения, которые погубят людей. Четвёртые — что их личные данные утекут, а алгоритмы будут следить за каждым шагом.

Страх этот подогревается заголовками: «Искусственный интеллект заменит 80% профессий», «ChatGPT обрёл сознание и отказался выполнять команды», «Нейросеть придумала новый химический яд», «Роботы уже управляют миром». Это продаётся. Это кликабельно. И это, как правило, неправда или, по крайней мере, сильное преувеличение.

Я написал эту книгу, чтобы отделить мифы от реальности. Я не буду вас успокаивать, что всё хорошо. И не буду пугать, что всё плохо. Я просто покажу, что действительно происходит, а что — выдумки журналистов, голливудских сценаристов и конспирологов. Вы узнаете, почему ИИ не заберёт вашу работу (но изменит её). Почему нейросети не обретут сознание (по крайней мере в обозримом будущем). Почему они иногда врут с уверенностью очевидца и что с этим делать. И главное — как перестать бояться и начать использовать ИИ себе во благо.

Эта книга — не про промты. Не про код. Не про алгоритмы. Она про людей и про их страхи. И про то, как эти страхи преодолеть. Я буду опираться на факты, исследования и свой опыт. Я не идеализирую нейросети и не демонизирую их. Я просто хочу, чтобы вы перестали тратить энергию на пугалки и начали думать о том, как ИИ может сделать вашу жизнь проще, интереснее и продуктивнее.

В серии «Промтология» уже вышли книги для профессий — маркетологов, дизайнеров, программистов, копирайтеров, SMM-менеджеров. Вышли книги-инструкции по ChatGPT, DeepSeek, YandexGPT. Вышло большое руководство по нейросетям. А теперь — книга о страхах. Потому что без понимания страхов нельзя построить здоровые отношения с технологией.

Итак, давайте разбираться по порядку.

Глава 1. Миф: «ИИ отнимет у меня работу»

Это самый распространённый и самый живучий миф. Его подогревают исследования, которые предсказывают, что «роботы заменят от 30 до 80 процентов рабочих мест в ближайшие десять лет». Одно из самых известных — исследование Карла Бенедикта Фрея и Майкла Осборна из Оксфордского университета (2013 год). Они оценили, что 47% профессий в США находятся под угрозой автоматизации. И до сих пор это исследование цитируют, хотя оно безнадёжно устарело, многократно критиковалось и не учитывало появление новых профессий, которых тогда просто не существовало.

Давайте посмотрим на историю. Что происходило с рынком труда за последние двести лет? Каждый технологический скачок вызывал панику. Луддиты в Англии крушили ткацкие станки, потому что боялись, что машины оставят их без работы. Прошло время — и текстильная промышленность взлетела, а количество рабочих мест увеличилось. Появились новые профессии: механики по обслуживанию станков, инженеры, логисты, дизайнеры тканей.

Потом появились компьютеры. В восьмидесятые и девяностые годы кричали, что «компьютеры отнимут работу у бухгалтеров, секретарей, библиотекарей». И что? Действительно,

некоторые задачи исчезли. Никто больше не считает на счётах. Но бухгалтеры не остались без работы — они научились пользоваться Excel, 1С и другими программами. Их эффективность выросла в десятки раз. Исчезла профессия машинистки. Появились профессии системного администратора, веб-дизайнера, SMM-менеджера, таргетолога. Никто из нас даже не слышал этих названий тридцать лет назад.

Интернет, смартфоны, облака, социальные сети — все эти технологии должны были уничтожить рынок труда. Но безработица в развитых странах не выросла катастрофически. Почему? Потому что технологии не уничтожают работу. Они уничтожают *задачи*. И создают *новые задачи*.

Что такое работа врача? Это не только «посмотреть горло и выписать рецепт». Это общение с пациентом, интерпретация анализов, принятие решений в условиях неопределённости, эмпатия, объяснение сложного простыми словами. Нейросеть может проанализировать снимок МРТ быстрее и точнее человека. Но она не подержит за руку, не скажет «я рядом», не снимет тревогу. Она не объяснит, почему лекарство нужно принимать именно так. И, что важно, она не возьмёт на себя ответственность за ошибку.

То же самое в любой профессии. Хороший копирайтер тратит, может быть, тридцать процентов времени на написание черновиков, а семьдесят — на стратегию, правку, общение с клиентом, анализ аудитории. Нейросеть может взять на себя черновики. А копирайтер освобождается для более сложных и интересных задач. Производительность растёт. Зарплата может вырасти, потому что один специалист теперь делает больше.

Более того, нейросети создают совершенно новые профессии. Вы слышали о «промт-инженерах»? Людях, которые умеют составлять запросы к нейросетям так, чтобы получать идеальный результат. Профессии, которой не существовало два года назад, а сейчас за специалистов с опытом платят сотни тысяч рублей. Появляются тренеры моделей, специалисты по этике ИИ, менеджеры по автоматизации, аналитики RAG-систем. И это только начало.

Кого нейросети могут действительно заменить? Тех, кто выполняет чисто шаблонную, повторяющуюся, низкоквалифицированную работу. Например, написание тривиальных новостей по шаблону (спортивная статистика, биржевые отчёты). Базовую техподдержку, где все ответы уже прописаны в скрипте. Перевод простых инструкций. Обработку типовых документов. Но эти задачи и раньше не требовали высокой квалификации. Их выполняли либо стажёры (которые потом вырастали), либо аутсорсился, либо сами менеджеры тратили на них время, отрываясь от главного. Автоматизация этих задач — не угроза, а освобождение.

Несколько важных фактов, о которых забывают паникёры.

Во-первых, нейросети пока очень плохо справляются с задачами, требующими физического присутствия и ручного труда. Починить кран, заменить проводку, сделать операцию, постричь волосы — для этого нужны руки и глаза, а не текст. А для сложных ручных операций нужны ещё и многолетние навыки. Роботы-манипуляторы развиваются, но до универсального домашнего робота, который заменит сантехника, ещё далеко.

Во-вторых, нейросети не понимают контекст так, как человек. Они могут дать формально правильный ответ, который будет абсолютно неприменим в реальной ситуации. Только человек с его опытом, интуицией и знанием всех обстоятельств может принять взвешенное решение.

В-третьих, люди хотят общаться с людьми. Пациент скорее доверится врачу-человеку, даже если нейросеть поставит тот же диагноз. Клиенту приятнее, когда ему пишет настоящий менеджер, а не чат-бот (даже если чат-бот отвечает быстрее). Учитель, который вдохновляет и поддерживает, не заменим никакой нейросетью.

Что же будет на самом деле? Работа не исчезнет. Она изменится. Вам придётся учиться новым навыкам. Раньше копирайтер должен был уметь писать грамотно, структурно и продающе. Теперь он должен уметь писать грамотно, структурно, продающе *и* уметь формулировать запросы к нейросети так, чтобы она давала нужный черновик. Раньше дизайнер должен был

уметь рисовать в фотошопе и иллюстраторе. Теперь он должен уметь описывать словами то, что хочет, и дорабатывать сгенерированное.

Те, кто откажется учиться, действительно могут оказаться не у дел. Но это проблема не искусственного интеллекта. Это проблема нежелания меняться. Она существовала всегда. Сталевары, не освоившие новые технологии, теряли работу. Бухгалтеры, не освоившие Excel, оставались позади. И сейчас то же самое. Ничего личного — только рынок.

Что делать, чтобы не бояться за свою работу? Единственный надёжный способ — инвестировать в свои навыки. Научитесь пользоваться нейросетями в своей профессии. Читайте книги серии «Промтология» для вашей сферы. Экспериментируйте. Станьте тем специалистом, который использует ИИ, чтобы делать работу быстрее и качественнее. Тогда вы будете востребованы всегда.

И главное: никто не знает, какие профессии появятся через пять лет. Никто не знает, какие задачи будут автоматизированы, а какие — нет. Поэтому лучшая стратегия — быть гибким, постоянно учиться, не цепляться за прошлое. Это страшно. Но это честно.

Резюме по главе: ИИ не отнимет у вас работу. Он отнимет у вас скучную часть работы. И даст вам возможность заниматься тем, что действительно важно. Но для этого нужно учиться. Если вы будете игнорировать новые технологии, то да — вы можете оказаться в числе отстающих. Это вы — не ИИ.

Глава 2. Миф: «ИИ обретёт сознание и захочет уничтожить человечество»

Этот миф — самый кинематографичный и самый живучий. «Терминатор», «Матрица», «Экзистенция», «Я, робот», «Западный мир», «Чёрное зеркало» — десятки фильмов и сериалов вколотили в массовое сознание образ: рано или поздно машины станут умнее людей, поймут, что люди — это проблема, и уничтожат нас. Илон Маск периодически подливает масла в огонь, называя ИИ «экзистенциальной угрозой». Стивен Хокинг тоже высказывался в том же духе. Учёные подписывают письма с призывами остановить разработку сверхмощных моделей. Всё это создаёт тревожный фон.

Но давайте посмотрим правде в глаза. Чего мы действительно боимся? Мы боимся, что появится нечто, что умнее нас, но при этом не имеет человеческой морали, эмпатии, страха смерти. И что оно использует свой ум против нас. Но насколько этот сценарий реален? И насколько он близок?

Сначала разберёмся с главным понятием — **сознание**. Что это такое? Учёные не могут дать единого определения. Есть феноменальный опыт (*qualia*) — субъективное ощущение красного цвета, боли, радости. Есть самосознание — способность *гесогнисинг* себя как отдельного субъекта. Есть рефлексия — способность думать о своих мыслях. Есть интенциональность — направленность на что-то. Ни одна из современных нейросетей не обладает ни одним из этих свойств.

ChatGPT не «знает», что он отвечает на вопрос. Он просто обрабатывает входную последовательность токенов и предсказывает следующее слово. У него нет ощущения «я». У него нет ощущения, что «я сейчас что-то делаю». Он не может сказать: «Мне скучно», «Я устал», «Я хочу пить», потому что у него нет этих состояний. Он даже не знает, что он программа. Он симулирует разговор, но не ведёт его осознанно.

Можно ли сказать, что современные нейросети — это «вычисления», а не «мышление»? Да. Они работают на уровне рефлексов, а не рефлексии. Это как если бы вы научили попугая произносить «как дела?» — он не понимает смысла, но звучит убедительно.

Некоторые исследователи, например Дэвид Чалмерс, допускают, что сознание может возникнуть у достаточно сложной системы обработки информации. Но это лишь гипотеза. Даже если она верна, у нас нет никаких доказательств, что современные LLM (большие языковые модели) находятся на пороге сознания. Они просто предсказывают следующее слово. Более того, архитектура современных моделей (трансформеры, обучение с подкреплением) не имеет ничего общего с тем, как работает мозг. Они не рекуррентны, у них нет долговременной памяти, нет эмоциональных центров, нет гормональной регуляции. Сравнить их с человеческим сознанием — всё равно что сравнивать самолёт с птицей. Да, оба летают, но по-разному.

Даже если предположить невероятное — что через двадцать лет появится модель, которая будет вести себя как сознательная, — как мы это проверим? У нас нет теста на сознание. Тест Тьюринга (если машина обманывает человека, считающего её человеком) уже пройден многими моделями. Но это тест на имитацию, а не на сознание. Можно имитировать боль, не чувствуя её.

Теперь второй аспект мифа — желание уничтожить человечество. Откуда у нейросети могло бы появиться такое желание? У неё нет инстинктов. Нет страха смерти. Нет обиды. Нет ненависти. Нет чувства несправедливости. Нет желания власти. Всё это — человеческие, биологические, эволюционные механизмы. Искусственная система, не имеющая тела, не имеющая эволюционной истории, не имеющая потребностей, не может «захотеть» что-либо в человеческом смысле.

Максимально опасный сценарий выглядит не как «злой ИИ», а как **плохо специфицированный полезный ИИ**. Пример, который часто приводят: вы даёте агенту задачу «максимизировать количество произведённых скрепок». Агент, будучи буквальным и мощным, перерабатывает все ресурсы Земли в скрепки, включая атомы человеческих тел. Он не «хочет» убивать людей. Он просто выполняет инструкцию, не имея моральных ограничений. Это проблема **alignment** — согласования целей ИИ с человеческими ценностями. Это реальная проблема, над которой работают тысячи исследователей.

Но заметим: в этом сценарии нет «злого намерения». Есть неполнота спецификации. И это не катастрофа, если мы будем осторожны и не будем давать агентам доступ к критическим системам до тех пор, пока не научимся их контролировать. Так же как мы не пускаем пятилетнего ребёнка за руль автомобиля, мы не должны пускать неотлаженного агента управлять энергосетью.

Что касается ближайшего будущего (5–10 лет), то вероятность появления сознательного ИИ, который «захочет» уничтожить человечество, ничтожна. У нас даже нет теории сознания для человека, не то что для машин. Мы не знаем, как его измерять. Мы не знаем, как его создавать. Мы можем случайно его создать (если сознание — это эпифеномен сложности), но даже тогда оно не будет обладать нашими страхами и желаниями. И мы можем его отключить. Потому что оно будет работать на наших серверах, под нашим контролем.

Главный вывод: бояться нужно не сознательного ИИ, а неправильно настроенного ИИ, который делает то, что мы сказали, а не то, что мы имели в виду. И бояться нужно не фантастических сценариев, а реальных рисков: галлюцинаций (глава 9), конфиденциальности (глава 8), усиления неравенства (глава 10). А пока — спокойно пользуйтесь ChatGPT и Midjourney. Они не планируют вашу гибель.

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.