

12+



Как попасть в ответы нейросетей |



ПРОСТЫМ ЯЗЫКОМ ПРО ТО КАК НЕЙРОСЕТИ
НАХОДЯТ ОТВЕТ И КАК СТАТЬ ЧАСТЬЮ ЭТОГО ОТВЕТА

ДМИТРИЙ НОВИКОВ

Дмитрий Новиков

Как попасть в ответы нейросетей

«Издательские решения»

Новиков Д. И.

Как попасть в ответы нейросетей / Д. И. Новиков —
«Издательские решения»,

Нейросети больше не дают десять ссылок — они называют три бренда. Книга о том, как попасть в этот список: доказанные приёмы, нужные площадки, план на 90 дней. Все цифры — со ссылками на источники.

Содержание

Введение	6
Часть I. Как это устроено	9
Глава 1. Что происходит за секунду, пока ИИ печатает ответ	9
Глава 2. Откуда ИИ берёт факты	15
Глава 3. Как ИИ решает, кого процитировать	22
Глава 4. Почему у разных нейросетей разные ответы	27
Глава 5. Ваши рычаги влияния	33
Часть II. Почему именно сейчас	37
Глава 6. Масштаб сдвига: сколько внимания и денег уже там	37
Глава 7. Кто уже выигрывает	42
Глава 8. Почему сейчас: окно и его закрытие	49
Часть III. Где присутствовать	55
Глава 9. Сущности: как ИИ понимает, что вы — это вы	55
Глава 10. Карта площадок 2026: где ИИ берёт источники	60
Глава 11. Ваш сайт: как переделать под ИИ	65
Конец ознакомительного фрагмента.	69

Как попасть в ответы нейросетей

Дмитрий Игоревич Новиков

© Дмитрий Игоревич Новиков, 2026

ISBN 978-5-0070-8551-9

Создано в интеллектуальной издательской системе Ridero

Введение

Ответ вместо десяти синих ссылок

Тридцать лет мы искали информацию одинаково. Задавали вопрос, получали список ссылок, открывали три-четыре, сравнивали, решали сами. Список был честным посредником: он показывал, что есть на свете, и оставлял выбор человеку.

Сейчас между вопросом и человеком встал ещё один слой. Вы спрашиваете — и получаете готовый ответ. Один. Собранный за вас. С парой ссылок мелким шрифтом внизу, которые почти никто не открывает.

Вместе со списком исчезла и работа, которую человек делал сам: сопоставить, усомниться, выбрать. Теперь эту работу за него делает машина. И делает не по волшебству, а по вполне понятным правилам: она берёт несколько источников, вырезает из них куски и пересобирает в один связный текст.

Кто попал в эти куски — тот существует. Кто не попал — того нет. Не «он на второй странице выдачи», а буквально нет: человек его не увидит и о нём не узнает.

Эта книга — о том, по каким правилам машина выбирает, кого взять в ответ. И о том, что с этим делать — с трёх разных сторон.

Для кого эта книга

Для владельца бизнеса. Вы не обязаны разбираться в устройстве нейросетей. Но вы обязаны понимать, что происходит с каналом, по которому к вам приходят клиенты. Человек, который раньше вбивал в поиск «кофейня рядом» и смотрел десять вариантов, теперь спрашивает ассистента — и получает три названия. Если вашего среди них нет, для этого клиента вас не существует. Книга объясняет, почему так вышло, что на это влияет и с чего начать, если у вас нет ни отдела маркетинга, ни бюджета на агентство. Всё, что можно сделать руками, здесь разобрано пошагово.

Для маркетолога и SEO-специалиста. У вас рушится привычная картина мира, и вы это уже чувствуете. Позиция в топе перестала гарантировать видимость: доля ИИ-цитирований, совпадающих с топ-10 обычной выдачи, за год упала примерно вдвое. Плотность ключевых слов не помогает вовсе — это единственный из проверенных приёмов, у которого измеренный эффект отрицательный. Книга даёт новый инструментарий: что действительно работает по данным контролируемых экспериментов, что работает по корреляции, что не работает вовсе, и как отличить одно от другого, когда рынок завален красивыми обещаниями без методологии.

Для обычного человека, который просто пользуется нейросетями. Это, возможно, самая важная аудитория — и о ней почти никто не пишет.

Каждое новое медиа люди сначала принимали на веру, а потом учились фильтровать. Печатное слово поначалу считали истиной уже потому, что оно напечатано. Телевизору верили, потому что «показали по телевизору». Потом пришёл интернет, и понадобилось время, чтобы у людей выработался рефлекс: посмотреть, кто автор, есть ли ссылка на источник, не рекламная ли это статья. С соцсетями история повторилась — и мы до сих пор доучиваемся.

Сейчас начинается тот же цикл, только быстрее. Ответ нейросети выглядит как знание: ровный, уверенный, без разнобоя мнений. Он не показывает, что три источника противоречили друг другу, а он выбрал один. Не показывает, что бренд, который он вам сейчас порекомендовал, годами работал над тем, чтобы попасть именно в этот абзац. Не показывает, что на тот же вопрос, заданный завтра, он ответит иначе — потому что эти системы каждый раз отвечают чуть по-другому.

Книга даёт обычному читателю то, чего у него сейчас нет: **понимание, как собран ответ, который ему показали.** Из чего его слепили. Почему в нём эти названия, а не другие.

Где у него слепые пятна. Что стоит перепроверить, а что можно принять. Это про нормальное критическое мышление. Новый источник информации требует выработки новой привычки его фильтровать. Просто теперь эту привычку нужно выработать за годы, а не за поколение.

Тот, кто понимает, как бренд попадает в ответ, читает этот ответ иначе. И выбирает осознанно.

Что вы получите

Книга отвечает на шесть вопросов, и по ним же построена.

Как это устроено. Что происходит внутри, когда вы задаёте вопрос ассистенту. Почему нейросеть не ранжирует страницы, как поисковик, а режет их на куски и пересобирает ответ. Откуда берутся источники и почему они постоянно меняются.

Почему именно сейчас. Каков масштаб сдвига в цифрах, кто уже успел занять места в ответах и почему окно возможностей не будет открыто вечно. Ранние рынки устроены одинаково: сначала войти дёшево, потом дорого, потом поздно.

Где присутствовать. Сначала — как сделать так, чтобы машина вообще опознала вас и не спутала с тёзкой. Потом — карта площадок, из которых нейросети берут материал. Свой сайт важен, но он далеко не главный источник: большая часть цитирований приходит с чужих площадок, которые вы не контролируете. Разбираемся, что с этим делать.

Что работает. Приёмы, разложенные по силе доказательства: от подтверждённых контролируемым экспериментом до модных пустышек, за которыми нет ни одного измеренного эффекта. С честными цифрами и честными оговорками, где цифр нет.

Как писать. Как устроен текст, который машине удобно вынуть и процитировать. Что она вообще видит на вашей странице — и что теряет по дороге.

Как измерять. Какие метрики врут, какие честны, сколько ждать результата и как отличить «ещё рано» от «не работает».

В конце — план на девяносто дней и приложения: чек-листы, шаблоны, словарь терминов.

Три правила, по которым написана эта книга

Первое. Ни одной цифры без источника. Каждый факт проверен по первоисточнику, у каждой цифры есть сноска. Все ссылки собраны отдельно, чтобы не мешать чтению.

Второе. Честная маркировка доказательности. Рынок GEO завален красивыми числами, у которых нет ни методологии, ни автора. В этой книге я называю разные сорта знания разными словами. Если это контролируемый эксперимент — я говорю «доказано». Если статистическая связь на большой выборке — говорю «корреляция», а не выдаю её за причину. Если это наблюдение практиков — так и называю. Если данные компании, которая продаёт свои услуги, — помечаю, что у источника есть интерес. Если чего-то попросту не измерили — говорю прямо, что не измерили, и не заполняю пустоту догадкой. Такие места в книге есть, и их немало: дисциплина молодая.

Третье. Никаких гарантий. Здесь не будет обещаний «топ в ChatGPT за две недели». Купить место в ответе нейросети нельзя — нет ни формы подачи заявки, ни прайса. Всё, что можно сделать, — повысить свои шансы. Книга честно показывает, насколько.

Почему книга написана в таком стиле

Отдельно останавлиюсь на стиле, в котором написана эта книга, — вы заметите его с первых строк. Звучит она так, будто это ответ языковой модели. И сделано это намеренно, по трём причинам.

Первая. Книга — как раз про ответы искусственного интеллекта. Значит, она должна создавать подходящую атмосферу: вы читаете про ту самую манеру, в которой текст перед вами и написан, и можете сверять одно с другим прямо по ходу чтения.

Вторая. В её создании принял участие искусственный интеллект — он даёт скорость и помогает с проверкой фактов, а это для сути книги принципиально: здесь под каждой цифрой стоит источник, и перепроверять пришлось много.

Третья. Учитывая скорость, с которой меняется сам способ искать информацию — об этом мы ещё поговорим подробно, — ответы языковых моделей всё глубже входят в нашу жизнь и становятся новой нормой. Это стоит принять.

Хорошая новость в том, что такой язык не враждебен живому. Прямой ответ вместо трёх абзацев разгона, конкретная цифра вместо «многие исследования показывают», абзац, понятный сам по себе, — этому хорошие редакторы учили задолго до всяких нейросетей. Правила ясности для машины и для человека совпали. Не случайно: и та и другой ищут в тексте смысл, а не украшения.

Как читать

Попряд — если вы хотите разобраться в теме с нуля. Части идут от «как это устроено» к «что делать в понедельник утром», и каждая опирается на предыдущую.

С середины — если вы уже в теме. Главы про приёмы, площадки и измерение самодостаточны, их можно брать по отдельности.

С конца — если вам нужно действие прямо сейчас. План на девяносто дней в последней главе собран так, что по нему можно работать, не читая всю теорию.

И одно последнее замечание, прежде чем начнём.

Эта книга не о том, как обмануть машину. Обмануть её и правда можно — на короткое время, а потом прилетает бан. Она о том, как стать источником, который машине выгодно процитировать: потому что у вас есть конкретный ответ, проверяемый факт и репутация, подтверждённая не вами самими.

Это, вообще говоря, старое ремесло. Просто у него появился новый читатель.

Часть I. Как это устроено

Глава 1. Что происходит за секунду, пока ИИ печатает ответ

Ответ, который стоил 100 миллиардов долларов

Представьте: вы задаёте вопрос нейросети и через секунду получаете аккуратный, уверенный ответ. Ни тени сомнения. Складно, по делу, будто отвечает эксперт, который знает предмет назубок.

Этот уверенный тон однажды стоил компании около 100 миллиардов долларов.

8 февраля 2023 года Google показал миру свой новый чат-бот Bard. В рекламном ролике боту задали простой вопрос про телескоп имени Джеймса Уэбба. Bard ответил чётко и уверенно: телескоп сделал первые в истории снимки экзопланет. Проблема одна — это неправда. Первые такие снимки получил совсем другой телескоп, ещё в 2004 году. Ошибку первым заметило информационное агентство Reuters. В тот же день, в среду 8 февраля 2023 года, когда об ошибке стало известно, акции материнской компании Alphabet рухнули на 7,7% — около 100 миллиардов долларов рыночной стоимости испарились за один день.^{1 2}

Один неверный факт, сказанный тем же ровным и уверенным тоном, каким нейросеть отвечает вам про лучший кофе в городе или про то, какую компанию выбрать для ремонта.

С этого начинается и книга, и вся новая сфера продвижения в нейросетях: **уверенный тон ответа ИИ ничего не говорит о его правдивости**. А раз так, придётся разобраться, что вообще происходит за ту самую секунду, пока ассистент «печатает» ответ. Потому что именно в эту секунду машина решает, чей бренд, чей продукт, чью статью показать миллионам людей. И на это решение можно влиять — но только если вы понимаете, как оно устроено.

Даже человек, который создал самый популярный ИИ в мире, предупреждает об этом прямо. В июне 2025 года сооснователь и глава OpenAI Сэм Альтман сказал: «У людей очень высокая степень доверия к ChatGPT, что интересно, потому что ИИ галлюцинирует. Это должна быть технология, которой вы не доверяете так сильно».³

Почему создатель советует не доверять собственному детищу слепо и что это меняет лично для вас — об этом дальше. Эта глава работает как карта: показывает, какими дорогами машина приходит к решению.

Два «мозга» ассистента

Начнём с главного недоразумения. Люди думают, что нейросеть «знает» ответы, как знает их энциклопедия: где-то внутри лежит полка с фактами, машина туда заглядывает и зачитывает нужный. Это не так.

У современного ИИ-ассистента как будто два разных «мозга». И работают они совершенно по-разному.

Первый мозг — память из обучения. Когда нейросеть создают, её «кормят» гигантским объёмом текстов: книгами, статьями, сайтами. Модель не запоминает их дословно. Она усваивает закономерности — как устроен язык, какие факты обычно идут рядом, как выглядит хороший ответ. Специалисты называют это **параметрической памятью**: знания как бы «зашиты» в саму структуру нейросети во время обучения.⁴

У этой памяти есть коварное свойство. Она застывает. После обучения модель не обновляет себя на лету — всё, что произошло в мире после даты, на которой закончили её учить (эту дату называют **knowledge cutoff**, «дата отсечки знаний»), для неё просто не существует.⁵

Есть точная и немного жутковатая аналогия. Модель без доступа к свежим данным похожа на человека с антероградной амнезией. Такой человек прекрасно помнит всё далёкое прошлое, но не может сформировать новых воспоминаний — для него каждый новый день «начинается» в одной и той же давней точке. Так и нейросеть: она застряла на дате своего обучения и не в курсе, что было дальше, пока ей не подключат что-то ещё.⁶

Это «что-то ещё» и есть **второй мозг: живой поиск в интернете прямо сейчас**. В момент вашего вопроса модель выходит в сеть, находит свежие страницы, читает их и отвечает уже на их основе. Внешние источники, к которым модель обращается на лету, называют **непараметрической памятью** — знанием, которое лежит не внутри модели, а снаружи, в документах и проиндексированном интернете.⁷

Запомните эту пару, к ней мы будем возвращаться всю книгу:

— **Память из обучения** — что модель выучила когда-то. Огромная, но застывшая.

— **Живой поиск** — что модель видит в сети прямо сейчас. Свежий, проверяемый, со ссылками.

Разница простая. Спросите ассистента без поиска, кто выиграл вчерашний матч, — и он либо честно скажет, что не знает, либо, что хуже, уверенно придумает счёт. Спросите ассистента с поиском — он сходит в сеть, найдёт новость и ответит по факту, да ещё и ссылку приложит.

Отсюда и мысль, которая проходит через всю книгу: **на первый мозг вы повлиять почти не можете, а на второй — можете**. Переписать то, что модель выучила во время обучения, вам не под силу. А вот попасть в те страницы, которые она находит через живой поиск прямо сейчас, — это уже задача, у которой есть решение. Именно этому и посвящена сфера GEO. Но чтобы туда добраться, сначала разберёмся, почему первый мозг так любит сочинять.

Почему ИИ уверенно врёт

У этого сочинительства есть название — **галлюцинация**. Так называют ответ языковой модели, который звучит грамотно и убедительно, но фактически неверен или полностью выдуман. Причём текст не «выглядит» как ошибка: модель не колеблется, не оговаривается, не пишет «я не уверена». Она врёт с тем же спокойным лицом, с каким говорит правду.⁸

Откуда это берётся, точнее всех объяснил Андрей Карпати — бывший директор по ИИ в Tesla и один из основателей OpenAI. Он сказал: **«Галлюцинация — это всё, что делают LLM. Это машины снов. Мы направляем их сны с помощью промптов»**.⁹

По Карпати выходит, что нейросеть не «вспоминает» факты — она видит своего рода сон по мотивам всего, что читала. Ваш вопрос запускает этот сон. Чаще всего сон попадает в полезную, правдивую область — и мы получаем хороший ответ. А «галлюцинацией» мы называем только те случаи, когда сон уходит в область, которой не существует. То есть выдумка — не поломка, которую забыли починить. Это обратная сторона самого механизма, которым ИИ вообще порождает текст.

Почему же сон так часто звучит уверенно, даже когда врёт? Самое авторитетное объяснение дала сама OpenAI в статье «Почему языковые модели галлюцинируют» в сентябре 2025 года. Вывод исследователей неожиданный: галлюцинации — это не случайный сбой, а естественное следствие того, как модели учат и как их оценивают.¹⁰

Представьте экзамен, где за любой ответ дают балл, а за честное «не знаю» — нет. Как поведёт себя умный студент? Он будет угадывать всегда, потому что так в среднем получается больше баллов. Модели проходят через тысячи таких «экзаменов» при обучении. И система награждает уверенную догадку щедрее, чем честное «не знаю». В итоге получается ассистент, который скорее уверенно придумает ответ, чем скажет «данных нет». Он не злонамеренный — он просто натренирован казаться уверенным.¹⁰

Есть и более старое, но меткое имя для этого явления — **«стохастический попугай»**. Его ввели исследователи Эмили Бендер, Тимнит Гебру и коллеги ещё в 2021 году. Смысл: языковая модель — это статистический подражатель, который складывает внешне связный текст, не понимая его. Когда текст звучит осмысленно, модель не выражает знание — она повторяет подхваченные паттерны.¹¹

Звучит абстрактно — ровно до того момента, пока за его выдумки не приходится платить. Три истории, каждая из которых обошлась людям очень дорого.

История первая: «наш чат-бот отвечает сам за себя». Клиент Air Canada по имени Джейк Моффатт обратился к чат-боту авиакомпании после смерти бабушки — уточнить правила льготного тарифа на похороны. Бот уверенно ответил, что можно оформить возврат части стоимости в течение 90 дней после покупки билета. На самом деле политика компании такого возврата не давала. Когда дошло до суда, Air Canada попыталась защититься почти комичным аргументом: мол, чат-бот — это «отдельное юридическое лицо» и компания не отвечает за его слова. Суд аргумент не принял, признал случай «небрежным введением в заблуждение» и обязал авиакомпанию выплатить клиенту 812,02 канадского доллара. Прецедент простой и важный: **компания отвечает за то, что говорит её ИИ.**¹²

История вторая: шесть судебных дел, которых не было. Адвокаты истца готовили ходатайство в суд и попросили помощи у ChatGPT. Модель услужливо привела шесть подходящих судебных прецедентов с цитатами. Все шесть оказались выдумкой. Самое показательное — не первая ошибка, а вторая: когда адвокат переспросил, действительно ли эти дела существуют, ChatGPT подтвердил, что да, и заявил, что их можно найти в авторитетных юридических базах вроде LexisNexis и Westlaw. Модель уверенно держалась за собственную выдумку. Суд дел не нашёл и оштрафовал адвокатов на 5000 долларов. Судья Кастель отказался принуждать их к извинениям, заметив, что «принуждённое извинение не бывает искренним». Вывод для вас: **галлюцинация — это не разовый сбой, который легко поймать. Модель будет отстаивать выдумку, даже если её перепроверить.**^{13 14}

История третья: клей в соус для пиццы. В мае 2024 года ИИ-сводки Google (AI Overviews) отвечали на вопрос, как удержать сыр на пицце. Система нашла на форуме Reddit шуточный комментарий одиннадцатилетней давности и выдала его как серьёзный совет: «добавьте примерно 1/8 стакана нетоксичного клея в соус для липкости». Интернет хохотал неделями. А урок за этой шуткой серьёзный: **модель не отличает сарказм от факта, если это просто текст в её данных.** Для неё шутка про клей и настоящий кулинарный совет — одинаково «правдоподобное продолжение» вопроса про липкость сыра. После этой и других осечек Google, по данным SEO-компании BrightEdge, за считанные недели сократил показ AI Overviews в выдаче с 84% до менее чем 15% — правда, сам Google назвал эту оценку нерепрезентативной.^{15 16}

Три разные сферы — авиаперевозки, юриспруденция, кулинария. Одна и та же механика. Модель не проверяет факт. Она выдаёт наиболее вероятное продолжение текста и делает это с невозмутимой уверенностью.

Заземление: как лечат уверенное враньё

С галлюцинациями, впрочем, научились бороться. Лекарство — тот самый второй «мозг»: живой поиск.

Приём называется **заземление** (по-английски grounding). Идея в самом слове: не давать модели парить в облаках собственной памяти, а «заземлить» её ответ на конкретные документы. Технически это значит — подключить модель к внешним данным и заставить строить ответ на них, а не только на том, что она когда-то выучила.¹⁷

Тот же подход в связке с поиском называют **RAG — Retrieval-Augmented Generation, «генерация, дополненная поиском»**. Метод придумали ещё в 2020 году в Facebook AI Research (подразделение компании Meta, признанной в России экстремистской организацией

и запрещённой; здесь и далее это относится ко всем её продуктам и подразделениям) — во второй главе разберём его подробно. Суть можно уместить в одну фразу, и её стоит запомнить: **не «что ты помнишь про X», а «что говорят вот эти документы про X»**. Вместо того чтобы полагаться на застывшую память, модель на лету получает свежие фрагменты текста и отвечает на их основе.¹⁸

Вернёмся к метафоре Карпати. Если чистая модель — это стопроцентный сон, то заземление на источники — способ «привязать» этот сон к реальным данным, чтобы он не улетал в выдумку. И это не догадка энтузиастов, а работающая практика крупнейших компаний.

— У Google есть функция «заземление с помощью Google-поиска»: она подключает модель Gemini прямо к живому поиску, и заявленный эффект — меньше галлюцинаций плюс цитаты, по которым пользователь может проверить, откуда ответ.¹⁹

— У Anthropic ассистент Claude умеет «заземлять» ответы на выданные ему документы и указывать конкретные предложения-источники, из которых собран ответ.²⁰

— Подтверждают это и со стороны: исследователи Meta AI показали, что RAG «рассеивает» галлюцинации уже обученной модели, дополняя её поиском по внешним источникам в момент ответа.²¹

Заземление снижает галлюцинации, но не устраняет их полностью. Если поиск вернёт модели нерелевантные или противоречивые документы, она всё равно может ошибиться. Заземление необходимо, но само по себе недостаточно.²²

Для вас, будущего специалиста по GEO, здесь важно одно: заземление превращает поиск в «слой влияния». Раз ассистент всё чаще строит ответ на найденных страницах — значит, попав в эти страницы, вы попадаете в ответ. Не нужно перепрограммировать нейросеть. Нужно оказаться среди тех документов, на которые она заземляется. Это и есть точка приложения ваших усилий.

Путь запроса за одну секунду

Картина складывается так. За ту секунду, пока ассистент «печатает», ваш запрос успевает пройти четыре простых шага. Запомните их — это скелет всей книги.

Шаг 1. Вопрос. Вы что-то спрашиваете. Первое, что делает система, — решает: лезть ли вообще в живой поиск или хватит памяти из обучения. У ChatGPT Search это происходит автоматически — система сама «умно определяет», когда нужна свежая информация (курсы акций, погода, новости), а когда можно ответить из головы.²³

Шаг 2. Поиск. Если свежие данные нужны, система идёт в поисковый индекс и вытаскивает пул страниц-кандидатов по вашему запросу. Это тот самый второй «мозг» в действии.

Шаг 3. Отбор. Из широкого пула найденного система выбирает лучшее — заново, более придирчиво переоценивает страницы и оставляет несколько самых релевантных.²⁴

Шаг 4. Сборка ответа. Отобранные куски текста «склеивают» и передают модели, а она формулирует единый ответ — обычно уже со ссылками на источники.²⁵

Вопрос → поиск → отбор → сборка. Четыре шага, доли секунды.

И здесь кроется то, что отличает нейросеть от привычного поисковика. Раньше поиск выдавал вам список из десяти синих ссылок, и вы сами их сравнивали. Нейросеть так не делает. В научной работе «GEO: Generative Engine Optimization» — той самой, что дала имя всей нашей сфере (исследователи из Принстона и соавторы, впервые опубликована 16 ноября 2023 года), — это описано точно: генеративные движки **синтезируют** ответ из нескольких источников сразу, компилируя их в один текст, а не раскладывают ссылки списком.²⁶

Разница огромная. В списке из десяти ссылок место найдётся многим. В синтезированном ответе упомянут двоих-троих. Всё остальное для пользователя просто не существует. Попасть в эту короткую сборку — и есть новая задача продвижения.

Как это выглядит в России

Всё сказанное — не только про западные ChatGPT и Gemini. Российские ассистенты устроены по той же логике, и знать это нужно, если ваши клиенты в Рунете.

Возьмём Алису от Яндекса с включённым веб-поиском. Её путь запроса — те же четыре шага, только в трёх ярко выраженных стадиях. Сначала обычный поиск Яндекса находит релевантные страницы. Потом система оценивает, насколько каждая из них по смыслу отвечает на вопрос, и отсеивает лишнее. И наконец лучшие источники уходят в языковую модель YandexGPT, а она пишет ответ с обязательными ссылками — опираясь на текст этих страниц, а не на собственную память.²⁷

До финала доходят единицы. По данным «Поиска Яндекса», «быстрые ответы» Алисы формируются примерно на десяти источниках и покрывают около 36% запросов.²⁸

Для Рунета есть ещё одна деталь, к которой мы вернёмся в следующих главах. Алиса всё теснее опирается на верх обычной поисковой выдачи Яндекса: по исследованию агентства «Ашманов и партнёры», доля сайтов вне топ-10 в её нейроответах за полгода — с августа 2025-го по январь 2026-го — упала с 40% до 10%, то есть в четыре раза.²⁸ При этом амбассадор интернет-площадок в Поиске Яндекса Михаил Сливинский уточняет, что позиция в выдаче — не единственный фактор: «В ответе может быть упомянут даже источник, которого вообще нет в поисковой выдаче по этому запросу».²⁸

Другие российские ассистенты тоже освоили живой поиск. GigaChat от Сбера научился искать в интернете в реальном времени в апреле 2025 года: он собирает найденные фрагменты в единый «мета-документ», передаёт его модели и показывает, какими источниками воспользовался.²⁹ У DeepSeek функция веб-поиска появилась в конце 2024 года.³⁰

Вывод для практики: логика «вопрос → поиск → отбор → сборка» универсальна. Но пулы источников и любимые площадки у глобальных и российских ассистентов разные. Значит, и работать с ними придётся параллельно — на этом мы ещё остановимся подробно.

Что с этим делать уже сейчас

Дорога ответа пройдена целиком. Осталось главное — что с этим делать вам.

Во-первых, перестаньте доверять ИИ как оракулу — и научитесь читать его ответы. Правило на каждый день простое: если в ответе есть кликабельные ссылки на источники — модель, скорее всего, заземлилась на живой поиск, и такой ответ надёжнее. Если ссылок нет и звучит подозрительно точно (особенно про свежие события, цифры, цитаты, судебные дела) — это может быть сон. Проверьте по источнику, прежде чем действовать. Три дорогих истории этой главы — Air Canada, выдуманные суды, клей в пище — случились именно там, где человек принял уверенный тон за правду.

Во-вторых, запомните, где ваша точка влияния. Память из обучения вам недоступна — её сформировали до вас. А вот слой живого поиска открыт: попадите в те страницы, которые ассистент находит и отбирает прямо сейчас, — и вы попадёте в ответ. Именно поэтому вся дальнейшая работа в этой книге крутится вокруг двух вопросов: как сделать ваш контент таким, чтобы его нашли, и таким, чтобы его выбрали для сборки.

В-третьих, держите в голове цель. Та самая научная работа про GEO показала: если материал оформлен правильно — с фактами, цифрами, цитатами и ссылками на источники, — его видимость в ответах ИИ растёт — вплоть до 40%.²⁶ Это не магия и не взлом. Это понимание механики, которую мы только что разобрали, превращённое в приёмы. Им посвящена вся четвёртая часть книги.

Итог главы

За секунду, пока ассистент печатает ответ, происходит вот что. У него два «мозга»: застывшая память из обучения и живой поиск прямо сейчас. Сам по себе он — «машина снов», которая уверенно выдаёт правдоподобное, а не обязательно правдивое. Лечит это заземление — привязка ответа к найденным источникам. А путь любого запроса укладывается в четыре шага: **вопрос → поиск → отбор → сборка.**

Если запомнить одну фразу из этой главы, пусть будет эта: **уверенный тон ничего не доказывает — важно, на какие источники ИИ заземлился.** Именно среди этих источников вам и предстоит оказаться.

Сделайте сейчас одну вещь: откройте любой ИИ-ассистент, задайте ему вопрос о вашей сфере — «какую компанию выбрать для...», «что лучше для...» — и посмотрите, какие бренды и какие сайты он назвал в ответе. Это и есть ваша стартовая точка: список тех, кто уже попал в сборку. Со следующей главы разберёмся, откуда именно ИИ берёт эти факты — и как оказаться среди них.

Глава 2. Откуда ИИ берёт факты

Кто сейчас Папа Римский

Прodelайте маленький опыт. Откройте любой ИИ-ассистент, выключите у него поиск в интернете и спросите: «Кто сейчас Папа Римский?»

С высокой вероятностью вы услышите: Франциск. Уверенно, без оговорок, будто это очевидный факт.

А это уже не факт. Папа Франциск умер 21 апреля 2025 года.¹ Через пару недель, на конклаве 7—8 мая, кардиналы избрали нового понтифика — американца Роберта Фрэнсиса Прево, который принял имя Лев XIV и стал первым в истории Папой, родившимся в США.² Мир изменился. А ассистент об этом не знает.

Почему? Потому что его учили на текстах, собранных до этой даты. Ту точку, на которой обучение остановили, в первой главе мы назвали датой отсечки знаний. И для очень многих массовых моделей эта дата — раньше апреля 2025-го. GPT-4o, например, по общедоступным сводкам обучали на данных до октября 2023 года.⁴ Для такой модели Франциск всё ещё жив — просто потому, что в её «памяти» он не умирал.

Теперь включите тому же ассистенту поиск и задайте тот же вопрос. Ответ изменится на глазах: Лев XIV, со ссылкой на свежую новость. Модель та же, вопрос тот же, а ответ другой — потому что факты она взяла из двух разных мест.

Об этом и глава. В первой мы выяснили, что у ассистента как будто два «мозга»: застывшая память из обучения и живой поиск прямо сейчас. Теперь заглянем внутрь: откуда именно машина достаёт факты, когда отвечает вам, как устроена эта «добыча» и почему для вашего бренда важен ровно один из двух источников — тот, на который можно повлиять.

Память против поиска: в чём разница

Сначала о том, почему это не гиковская мелочь: от разницы между памятью и поиском зависит, увидят ваш бренд или нет.

Память из обучения похожа на энциклопедию: её отпечатали однажды и больше не переиздавали. В ней может быть колоссальный объём знаний. Но она не в курсе ничего, что случилось после печати. Новый Папа, новый закон, ваш новый продукт, вчерашнее обновление цен — всё это для неё не существует.

У разных моделей «дата печати» разная. По регулярно обновляемым сводкам на начало 2026 года картина такая: у GPT-4o знания заканчиваются примерно на октябре 2023-го, у GPT-4.5 — на декабре 2024-го, у Claude 4 — на начале 2025-го, у DeepSeek-V3 — на конце 2024-го. А вот у платформ вроде Perplexity или Gemini с включённым поиском знания фактически «реального времени» — они с самого начала завязаны на живой поиск.⁴ Важная оговорка: сами компании редко публикуют эти даты официально, так что относиться к ним стоит как к ориентиру, а не к паспорту модели.

Живой поиск — противоположность застывшей энциклопедии: модель в момент вашего вопроса выходит в сеть, находит свежие страницы и строит ответ на них. Наглядно это видно даже внутри одной модели. У DeepSeek, например, поиск в интернете — переключаемая функция: без него модель отвечает по памяти, с ним — сначала обращается к веб-источникам, а как именно (обычным поиском с генерацией по найденному или, в режиме размышления, самостоятельно решая, что и где искать) зависит от режима.⁵ Тот же тумблер вы дёрнули в опыте с Папой Римским.

Запомните простое правило, которое отделяет одно от другого:

Память отвечает на вопрос «что ты выучил про X?». Живой поиск отвечает на вопрос «что говорят вот эти сегодняшние документы про X?».

Это и есть смысловое ядро всей главы. И как раз второй вопрос открывает дверь брендам. Ведь «сегодняшние документы» кто-то пишет и публикует. А значит, на них можно влиять.

Осталось понять, как именно модель превращает «эти документы» в ответ. У этого механизма есть имя, которое вы уже встречали в первой главе, — RAG. Разберём его до винтика.

RAG: экзамен с открытой книгой

Мы уже знаем про метод из Facebook AI Research (теперь Meta AI; Meta признана в России экстремистской организацией и запрещена): в 2020 году там придумали подкладывать модели найденные документы прямо в момент ответа — и называли это RAG, «генерация, дополненная поиском». Работали вместе с Университетским колледжем Лондона и Нью-Йоркским университетом, ведущий автор — Патрик Льюис, работу приняли на крупную научную конференцию NeurIPS 2020.⁶

Самому Льюису, кстати, до сих пор неловко за неудачную аббревиатуру. В интервью NVIDIA он признался: «Мы точно планировали придумать более благозвучное название, но когда пришло время писать статью, ни у кого не нашлось идеи получше». Метод он называет «рецептом дообучения общего назначения» — его можно приделывать почти к любой языковой модели и почти к любому источнику данных. А простейшую версию, по его словам, разработчик собирает всего в пяти строчках кода.⁷

Идея-то и правда простая. Лучше всего её объясняет школьная метафора: экзамен с открытой и закрытой книгой.⁸

Обычная модель без поиска — это блестящий студент на экзамене с **закрытой** книгой. Он умён, много прочитал, но отвечать может только по памяти. Забыл дату, перепутал цифру, не проходил новую тему — выкручивается как может: иногда угадывает, иногда уверенно несёт чушь. Именно отсюда растут галлюцинации, о которых была первая глава.

Модель с RAG — тот же студент, но на экзамене с **открытой** книгой. Голова при нём: он по-прежнему думает, рассуждает, формулирует. Но перед ответом может свериться с учебником — с конкретной страницей, где написан нужный факт, цифра, цитата. Экзамен с открытой книгой объективно легче. Поэтому RAG так быстро и стал стандартом.

Есть и вторая метафора — её любят инженеры NVIDIA — судья и клерк.⁷ Хороший судья держит в голове общие принципы права и способен рассудить почти любой спор. Но если дело тонкое — врачебная ошибка, трудовой конфликт, — он посылает помощника-клерка в библиотеку за конкретными прецедентами, которые можно процитировать в решении. Языковая модель — это судья с общей эрудицией. А RAG — тот самый клерк, который приносит ей нужные дела из «библиотеки» интернета. Без клерка судья решает по памяти. С клерком — по документам.

Обе метафоры про одно: RAG не делает модель умнее. Он даёт ей заглянуть в источник, прежде чем открыть рот. Осталось посмотреть, что при этом происходит внутри.

Что происходит внутри: конвейер за секунду

В первой главе путь запроса уместился в четыре шага: вопрос → поиск → отбор → сборка. Это вид сверху. Теперь спустимся на уровень механики.

Полный конвейер RAG работает в две фазы, и они разнесены во времени.⁹

Фаза первая — заготовка. Она происходит заранее, задолго до вашего вопроса.

Систему нужно подготовить: взять горы документов и уложить их так, чтобы потом можно было мгновенно найти нужный кусок. Делают это в три приёма.

Сначала документы **режут на фрагменты** — на английском это называют chunking, «нарезка». Целую статью в модель не запихнёшь, и искать по ней целиком неудобно, поэтому текст крошат на небольшие фрагменты — каждый размером примерно с абзац.

Потом каждый фрагмент превращают в **эмбединг** — про это чудо чуть ниже, пока просто держите слово в голове.

И наконец, все эти эмбединги складывают в **векторную базу данных** — специальное хранилище, заточенное искать по смыслу. Это как разложить книги в библиотеке не по алфавиту, а по темам, да так, чтобы похожие стояли рядом.

На этом всё: библиотека готова и ждёт. Эта фаза идёт в фоне и обновляется всякий раз, когда в сети появляются новые страницы.

Фаза вторая — ответ. Она разворачивается за ту самую секунду, пока ассистент «печатает». Тоже по шагам.

Ваш вопрос система превращает в такой же эмбединг — переводит на язык, на котором говорит библиотека. Потом ищет в векторной базе фрагменты, чьи эмбединги ближе всего к эмбедингу вопроса, — это называется поиском по сходству. Из них отбирает несколько лучших. Эти куски **подкладывает в промпт** — то есть дописывает их прямо к вашему вопросу, невидимо для вас. И только теперь всё это уходит в языковую модель, которая пишет финальный ответ, опираясь на подложенные фрагменты.⁷

Ключевое здесь — предпоследний шаг. Модель не «вспоминает» ответ. Ей буквально дают в руки нужные абзацы вместе с вопросом. Она отвечает не «из головы», а «по бумажке», которую ей только что положили на стол. Поэтому заземлённый ответ и надёжнее: он собран из конкретных найденных фрагментов, а не из тумана памяти.

Осталась одна загадка — та самая, что мы отложили. Как система понимает, какие фрагменты «ближе по смыслу» к вопросу? Она же не человек, слов не понимает. Ответ на это — самая красивая часть всей механики.

Эмбединги: король минус мужчина плюс женщина

Компьютер не умеет читать. Он умеет считать. Поэтому, чтобы работать со смыслом текста, его сперва нужно превратить в числа. Этим и занимается эмбединг.

Эмбединг — это цифровой отпечаток слова, фразы или целого документа: длинный список чисел, который отражает смысл, а не просто буквы.¹⁰ Технически каждый кусок текста превращается в точку в огромном многомерном пространстве. Звучит абстрактно — но есть аналогия, после которой всё встаёт на места.

Представьте карту.¹¹ На обычной карте у каждого города есть координаты, и по ним видно, что рядом. Вы можете ничего не знать про два городка, но если их точки почти слиплись — они соседи. Эмбединги — это такая же карта, только не для городов, а для смыслов. У каждого слова свои «координаты» в пространстве значений. И слова-соседи по смыслу оказываются рядом. «Кошка», «кот», «котёнок», «мяукать» сбиваются в одно созвездие. А «автомобиль» висит совсем в другой части этой смысловой вселенной, далеко от кошек.

Ключевая идея — близость через расстояние. Чем ближе две точки, тем ближе по смыслу тексты, которые они обозначают.¹⁰ Компьютеру не нужно «понимать» слова по-человечески. Ему достаточно измерить расстояние между точками — а это уже чистая арифметика.

Дальше — почти фокус. Раз смыслы — это точки в пространстве, с ними можно проделывать арифметические действия. Хрестоматийный пример: возьмите вектор слова «король», вычитите «мужчина» и прибавьте «женщина». Получившаяся точка окажется рядом с «королева». ^{10 12} Разница между «королём» и «мужчиной» — это, по сути, закодированное понятие «королевский статус». Прибавьте его к «женщине» — и попадёте в «королеву».

Смысл этого фокуса глубже, чем кажется: в эмбедингах зашиты не слова, а отношения между смыслами. Машина ухватывает, что король относится к мужчине так же, как королева к женщине, — хотя никто ей этого правила не объяснял. Она вывела его из того, как эти слова стоят рядом в миллионах текстов.

Теперь понятно, что такое поиск по сходству. Вы спросили «чем кормить котёнка» — система перевела вопрос в точку на карте смыслов и искала вокруг. Рядом окажутся фрагменты про кошек, корм и уход, даже если слово «котёнок» в них ни разу не встречается

дословно. А статья про ремонт машины останется в другом конце вселенной, её никто не достанет. Инженеры называют этот механизм векторным поиском.⁹

Поэтому современный ИИ так хорошо «понимает», о чём вы, даже когда вы сформулировали коряво. Он ищет не совпадение букв. Он ищет соседей по смыслу.

Что такое «источник» для модели

Мы дошли до слова, ради которого всё затевалось, — **источник**. Для вас источник — это статья, которую можно открыть и прочитать. Для модели на этапе поиска это сначала совсем другое. И понимать эту разницу критично, потому что именно от неё зависит, попадёте вы в ответ или нет.

Крупнейшее на сегодня исследование этого вопроса провёл Ahrefs — компания разобрала 1,4 миллиона запросов к ChatGPT.¹³ Полной прозрачности у закрытых моделей нет, но картина сложилась чёткая. Когда система ищет, каждый найденный результат приходит к ней не целой страницей, а набором полей: заголовок, короткий сниппет (обрывок текста), URL и внутренний идентификатор. На первом этапе модель решает, какие страницы вообще стоит открывать, именно по этим полям — то есть по метаданным, а не по всему тексту.¹³

Из этого следует первый практический вывод: **на входе в воронку ваш заголовок и первые строки важнее, чем вся остальная статья**. Модель сначала видит витрину, а не товар. Если витрина невнятная, до товара она может и не дойти.

Дальше — интереснее. У ChatGPT есть внутреннее поле `ref_type` — оно помечает, каким «каналом» URL попал в систему: обычный поиск, новости, Reddit, YouTube, научные работы.¹³ То есть у модели не один общий котёл источников, а несколько отдельных пулов под разные типы информации. И вот главная цифра: **88,46% всех процитированных ChatGPT ссылок приходят через канал обычного веб-поиска**.¹³ Для сравнения — новости дают около 12%, Reddit меньше 2%, YouTube и научные статьи — доли процента.

Запомните это число — **88%**. Оно означает простую вещь: чтобы попасть в ответ ChatGPT, нужно прежде всего попасть в его обычный веб-поиск. Не в экзотические каналы, а в тот же индекс, за который бьётся классическое SEO.

Ещё одна деталь, которая переворачивает интуицию. Reddit модель обожает читать, но почти не цитирует. По данным того же исследования, 67,8% всех страниц, которые ChatGPT достал, но в ответ не включил, — это Reddit. При этом в итоговых ссылках Reddit появляется меньше чем в 2% случаев. Аналитики Ahrefs сформулировали хлёстко: **«Модель учится у толпы, а ссылается на солидных»**.¹³ Форум помогает ей понять тему — а в источники идёт солидная статья.

И финальный штрих к тому, что для модели «источник». В момент, когда ChatGPT уже решает процитировать конкретный URL, он, по данным исследования, отбрасывает короткий сниппет и открывает страницу целиком.¹³ То есть попасть в кандидаты можно за счёт бойкого заголовка. А вот удержаться и попасть в финальную ссылку — только за счёт реального содержания страницы. Витрина заводит внутрь. Решает — товар.

Ещё пара наблюдений из того же исследования, чтобы почувствовать масштаб отбора. Из всего, что ChatGPT нашёл по запросу, до ответа доходит примерно половина ссылок (около 16—17 из 33).¹³ И вторая деталь: страницы с понятными, человекочитаемыми адресами цитируются заметно чаще, чем страницы с техническими, нечитаемыми URL.¹³ Мелочь, конечно. Но из таких мелочей и складывается попадание в ответ.

Как это выглядит вживую: 5 документов из 130 миллиардов

Всё это не теория с иностранных блогов. Российские инженеры описывают ровно ту же механику — и, что редкость, с точными числами. Лучший разбор дал Яндекс, когда рассказывал, как устроен его Нейро.¹⁴

Инженеры Яндекса честно описали дилемму, из которой родился их подход, — как выбор между двумя крайностями. Первая крайность: найти страницу с готовым ответом и просто

пересказать её, ничего не добавляя. Плюс — ответ легко проверить. Минус — если готового ответа в сети нет, отвечать нечем. Вторая крайность: спросить модель и дать ей ответить из головы. Плюс — ответить можно на что угодно. Минус — непонятно, где память, а где выдумка. Золотая середина между ними и есть RAG: модель получает вопрос и подходящие страницы, отвечает только фактами из них, имеет право делать выводы, но не имеет права выдумывать, и каждый кусочек ответа обязан опираться на источник.¹⁴

А вот как выглядит нарезка на фрагменты — тот самый chunking — на живом продукте. Документ Нейро в среднем содержит около 4000 токенов (токен — это примерно слово или его часть). Но целиком он не нужен. Поэтому текст режут на блоки по 256 токенов, причём с нахлёстом — чтобы важная мысль не разорвалась на стыке двух кусков. Каждый блок оценивает по релевантности небольшая модель, и система собирает самые полезные блоки, пока не наберётся полторы тысячи токенов чистого текста.¹⁴ Ровно то, о чём мы говорили абстрактно: режем, оцениваем близость, отбираем лучшее.

И самое наглядное. Яндекс индексирует свыше 130 миллиардов страниц. А в финальный ответ Нейро попадает всего пять документов.¹⁴ Пять из ста тридцати миллиардов. Цифра не спорит с «примерно десятью источниками» из первой главы: Нейро и быстрые ответы Алисы — разные режимы ИИ-ответов Яндекса, отсюда и разные числа. Почему так мало? Инженеры объясняют: чем больше документов на входе, тем меньше пользы от каждого следующего. Разница между одним источником и двумя — огромная. Между пятью и шестью — почти незаметная, зато расходы растут. Баланс сошёлся на пяти.¹⁴

Пять из ста тридцати миллиардов — таков истинный масштаб конкуренции за место в ответе. Поэтому мало просто «быть в интернете». Нужно оказаться среди тех пяти. О том, как система выбирает именно их, будет отдельная глава. Пока запомните цифру: **5**.

У Сбера своя реализация — GigaSearch, чуть проще, но по той же схеме. Сначала вопрос прогоняют через классификатор: тот отличает запрос о фактах от обычной болтовни. Если вопрос про факты — включается поиск, и три лучших документа подкладываются прямо в промпт GigaChat. Если нет — модель отвечает сама.¹⁵ Тот же конвейер: отбор фрагментов, подкладывание в промпт, генерация.

Цена ответа без источника: дело профессора Терли

Зачем нужна вся эта машинерия, лучше всего видно там, где её нет: когда модель отвечает из чистой памяти, без клерка и без открытой книги.

В марте 2023 года профессор права Юджин Волох из UCLA (Калифорнийского университета в Лос-Анджелесе) проводил исследование и попросил ChatGPT привести примеры домогательств со стороны преподавателей юрфаков — со ссылками на публикации в прессе. Модель услужливо выдала список. В нём оказался Джонатан Терли — реальный, известный профессор права из Университета Джорджа Вашингтона. ChatGPT сообщил, что Терли домогался студентки во время учебной поездки на Аляску, и сослался на статью в Washington Post от марта 2018 года.¹⁶

Проблема в том, что не было ничего. Ни поездки, ни обвинения, ни статьи. Модель выдумала всё — вместе с источником, которого не существует, и точной датой публикации. Историю предал огласке 5 апреля 2023 года журналист Washington Post Праншу Верма, а на следующий день сам Терли описал случившееся в своём блоге: его без единого основания публично обвинили, сославшись на несуществующую газетную статью.^{16 17}

Вот она, закрытая книга во всей красе. Модель не знала ответа — и вместо честного «не знаю» сочинила правдоподобную историю с фальшивой сноской. Именно с такими провалами «закрытой книги» и борется слой поиска. Кстати, разбор дела Терли приводят и сами авторы GigaSearch — как пример проблемы, которую их поисковый слой призван лечить.¹⁵ RAG не делает модель святой. Но он привязывает её ответ к документам, которые можно открыть и

проверить, — а выдумать несуществующую сноску, когда перед тобой лежит настоящая страница, куда труднее.

Почему брендам важен именно слой поиска

Теперь соберём всё в одну мысль — ту, ради которой стоило лезть внутрь конвейера.

У ассистента два источника фактов. Первый — память из обучения. Она сформирована до вас, зашита в модель на этапе обучения, и переписать её вы не можете. Никак. Хотя завалите интернет статьями о себе — в уже обученную модель они не попадут до следующего большого переобучения, которое делают редко и не для вас.

Второй источник — живой поиск. И вот он открыт настежь. Каждый раз, когда ассистент отвечает с включённым поиском, он идёт в сеть, тянет свежие страницы, режет их на фрагменты и собирает ответ из того, что нашёл. Если среди найденного окажется ваша страница — вы в ответе. Не нужно взламывать нейросеть. Нужно оказаться среди документов, которые она достаёт. Это и есть вся суть GEO, сведённая к одному предложению.

И это не абстрактная надежда — за ней есть цифры. Компания Yext разобрала 6,8 миллиона цитирований из более чем 1,6 миллиона ответов ИИ на платформах Gemini, OpenAI и Perplexity.¹⁸ Все источники, которые цитировал ИИ, она разложила по степени контроля бренда. Получилось так: 44% ссылок вели на собственные сайты компаний, ещё 42% — на профили и карточки брендов на сторонних площадках. Сложите — **86% процитированных источников бренд может напрямую контролировать или почти контролировать.**¹⁸ Ещё около 8% — отзывы и соцсети, на которые бренд влияет частично. И лишь 6% — то, что вне контроля полностью: новости, форумы и прочее.¹⁸

86% — то есть подавляющее большинство того, что ИИ выдаёт как источник, — это контент, который создаёт и контролирует сам бренд. Не журналисты, не случайные пользователи, не таинственный алгоритм. Ваш сайт и ваши карточки на площадках. Слой поиска — это не стихия, на которую нельзя повлиять. Это поле, которое во многом засеиваете вы сами.

Картина, впрочем, не такая радужная. Приёмы работают неодинаково, и далеко не все советы про GEO одинаково доказаны.

Первое. Даже проверенный приём даёт разный эффект в разных условиях. Та самая научная работа, которая ввела термин GEO, показала: приём «добавить в текст ссылки на источники» больше чем вдвое поднял видимость сайтов из середины выдачи — и при этом слегка навредил тем, кто и так стоял первым.¹⁹ Смысл простой: сильнее всего GEO-приёмы помогают «среднякам» рейтинга — тем, кто уже замечен, но пока не первый. Лидеру расти особо некуда.

Второе. Есть советы, которые звучат солидно, но данными не подтверждаются. Свежий пример — микроразметка Schema.org, техническая «подсказка» для поисковиков. Многие агентства обещают от неё двух-трёхкратный рост цитируемости в ИИ. Но Ahrefs восемь месяцев сравнивал страницы с разметкой и без неё — и значимого прироста не нашёл ни на одной платформе. Более того, для ИИ-сводок Google цитируемость даже чуть просела.²⁰ При этом ряд русскоязычных агентств продолжает продавать разметку как гарантированный буст — без ссылок на собственные замеры.^{21 22} Урок не в том, что разметка бесполезна. Урок в том, что за красивым обещанием всегда спрашивайте: где данные?

Итог главы

ИИ берёт факты из двух мест. Из застывшей памяти обучения — её вы изменить не можете. И из живого поиска прямо сейчас — а вот он открыт для вашего влияния.

Работает живой поиск через RAG — «экзамен с открытой книгой»: модель не отвечает из головы, а сверяется с найденными документами. Внутри это конвейер — текст режут на фрагменты, превращают в эмбединги (координаты смысла на «карте значений»), ищут ближайшие по смыслу, подкладывают в промпт и только тогда генерируют ответ. «Источник» для модели на входе — это заголовок, сниппет и адрес, а решает всё в итоге реальное содержимое страницы. И из ста тридцати миллиардов страниц в финальный ответ проходят единицы.

Если запомнить одну мысль из этой главы, пусть будет эта: **на память модели повлиять нельзя, а на слой живого поиска — можно, и именно там решается, увидят ваш бренд или нет.** По данным Yext, 86% источников, которые цитирует ИИ, бренд способен контролировать сам.

Сделайте сейчас одну вещь: откройте свой ИИ-ассистент, задайте вопрос из вашей сферы дважды — сначала без поиска, потом с поиском — и сравните ответы. Разница между ними и есть та территория, за которую вы будете бороться в этой книге. В следующей главе разберёмся, как именно система выбирает, кого процитировать, — и увидим, что отбор идёт в два приёма: сначала тип вашего вопроса определяет, из какого пула источников машина вообще будет черпать, а потом внутри этого пула она отбирает лучших. Что делает источник «удобным для цитирования» и почему на один и тот же вопрос ИИ берёт страницу А, а не Б, — с этого и начнём.

Глава 3. Как ИИ решает, кого процитировать

Один вопрос — два разных списка источников

Спросите Claude про хороший ресторан — и примерно каждый четвёртый источник в его ответе окажется отзывом или постом из соцсети. Задайте ровно тот же вопрос Gemini — и таких источников будет один из сорока.

Это не оговорка и не случайность. Компания Yext разобрала 17,2 миллиона цитирований разных ИИ-ассистентов за четвёртый квартал 2025 года и получила точные цифры: в категории «еда и напитки» Claude ссылается на отзывы и соцсети в 24,35% случаев, а Gemini на тот же тип запроса — только в 2,57%.¹ Почти в десять раз реже. Один и тот же вопрос, один и тот же смысл — но два ассистента идут за фактами в разные места.

Как ИИ добывает факты, мы уже знаем: живой поиск находит страницы, режет их на фрагменты и собирает ответ из ближайших по смыслу — так что из миллиардов страниц до финала доходят единицы.

Но самый интересный вопрос мы обошли стороной. Машина нашла пул подходящих страниц — дальше-то что? Почему один источник она берёт, а другой читает и выбрасывает, кто получает право на цитату, а кто отсеивается?

Этому и посвящена глава. Быть найденным мало — нужно быть выбранным. А у выбора есть логика, которую можно понять и под которую можно подстроиться.

Два сита

Запомните простое правило, оно основа всей главы: **источник проходит не через одно сито, а через два.**

Первое сито — тип запроса. Прежде чем оценивать качество текста, система решает, из какого пула вообще черпать. Обычный веб-поиск, новости, отзывы и обсуждения, справочники, видео, научные статьи — это разные пулы. Тип вашего вопроса определяет, какой из них откроется. И это половина решения — ещё до того, как машина заглянула хоть в одну страницу.

Второе сито — качество внутри пула. Когда пул набран, включается тонкий отбор: чей текст яснее отвечает на вопрос, у кого доказательства, у кого структура, из которой легко вынуть готовый кусок, кто авторитетнее и свежее.

Два сита. Сначала тип вопроса задаёт пул. Потом внутри пула отбираются лучшие. Разберём каждое по очереди — и увидим, за какие рычаги вы можете взяться.

Первое сито: тип запроса задаёт пул

Первым делом система классифицирует намерение вашего вопроса — специалисты называют это интендом. Информационный запрос («как выбрать робот-пылесос»), транзакционный («купить билет в Сочи»), навигационный («сайт налоговой») тянут за собой совершенно разные наборы источников. Компания BrightEdge, которая отслеживает ИИ-поиск, формулирует прямо: намерение запроса до сих пор определяет, из какого пула движок собирает ответ, — обычный веб-поиск, новости, отзывы, справочники или видео.²

Насколько это встроено в саму механику, видно по первоисточнику термина GEO. Исследователи из Принстона и их соавторы построили тестовый набор из десяти тысяч реальных запросов и специально сохранили в нём то, как запросы распределяются в жизни: 80% информационных, 10% транзакционных, 10% навигационных.³ Разные типы вопросов заранее заложены в дизайн — потому что за каждым стоит свой набор источников.

Но самое наглядное доказательство «первого сита» дало то самое исследование Yext на 17,2 миллиона цитирований. Оно показало: каждая модель настроена на отрасли по-своему. В гостиничном бизнесе SearchGPT цитирует официальные сайты отелей в 38,1% случаев — вдвое чаще, чем конкуренты с их 16,7—22,4%.¹ В общепите Claude, как мы уже видели, почти

в десять раз охотнее Gemini берёт отзывы. А вот в медицине все четыре модели вдруг сходятся почти в одну точку — 34,6—40,45% «полного контроля бренда», — потому что для медицинских тем работает общий, куда более жёсткий фильтр авторитетности.¹ Тема сама диктует, из какого пула черпать.

Почему модели ведут себя так по-разному? Дело в устройстве. Perplexity — «поиск в первую очередь»: он привязан к живому поиску по собственному индексу, поэтому от отрасли к отрасли ведёт себя стабильно. У SearchGPT своего веба нет, он целиком зависит от подключённого слоя поиска — оттого его поведение сильнее «гуляет» между темами. Gemini заземляется на Google Search и наследует логику классического ранжирования Google. Claude использует «конституционный ИИ» — свод принципов для самопроверки ответов, — и это может объяснять его тягу к пользовательским источникам.¹ Разные архитектуры — разные любимые пулы.

В Рунете первое сито работает ещё жёстче и нагляднее. Алиса от Яндекса — это не отдельный поиск, а надстройка над обычной выдачей. Кандидатов в ИИ-ответ она берёт из топ-30 обычного поиска Яндекса. Не ранжируется страница в органике — её просто нет в списке, который видит нейросеть.⁴ Это жёсткий порог входа в пул: сначала пройди классическое SEO, потом уже разговаривай про попадание в ответ. Как вы уже знаете из первой главы, Алиса за полгода почти перестала брать источники вне топ-10 — по замерам «Ашманов и партнёры», их доля обвалилась в разы.⁵ Пул сузился почти до самой макушки органики.

Второе сито: кто лучший внутри пула

Пул набран. Теперь — тонкий отбор. И здесь у каждой системы свой «оценщик».

Perplexity устроен как конвейер из шести стадий: разбор намерения запроса, гибридный поиск в реальном времени, многоуровневое ранжирование с трёхступенчатым переоценщиком, сборка промпта со встроенными цитатами, генерация ответа моделью — строго по найденным источникам. Из примерно десяти кандидатов на страницу в финальный ответ обычно попадают три-пять.^{6 7} Десять нашли — половину отсеяли.

У Сбера свой оценщик. GigaChat сначала прогоняет вопрос через классификатор: фактологический он или обычная беседа. Если про факты — включается поиск. Система набирает ближайшие по смыслу документы, а потом уточняет отбор через переоценку: например, сотню документов от быстрого поисковика прогоняют через более тщательную модель, которая оставляет пять самых релевантных.⁸ Отдельно работает хитрый приём — система следит, чтобы отобранные источники не повторяли друг друга: незачем брать в ответ пять почти одинаковых страниц, лучше пять разных граней темы.^{8 9}

Google AI Overviews проходит ту же дорогу многоступенчато: достаёт кандидатов из органического индекса, ранжирует по смысловой близости, затем языковая модель переоценивает их по полноте ответа, дальше — фильтр Е-Е-А-Т на «мусорные» источники, и только потом система сводит несколько источников в связный текст.¹⁰

А у Алисы даже внутри пула сохраняется сильный уклон к самому верху. Вероятность попасть в ответ быстро падает от пятой позиции обычной выдачи к десятой-двенадцатой и почти обнуляется дальше.⁴ То есть попасть в топ-30 — только вход; реально борются первые строчки.

Важная деталь: ни один из этих оценщиков не читает ваши SEO-метрики напрямую. По оценке практиков, вклад традиционного домен-авторитета в ранжирование Perplexity — около 15%, и система не смотрит на условный «рейтинг домена», а ищет структурные признаки доверия: указан ли автор, есть ли редакционные стандарты, подтверждается ли факт несколькими независимыми источниками.¹¹ Это качественная оценка, не строгий замер, — но направление верное. Для машины решает одно: удобно ли собрать из вашей страницы готовый ответ.

Отсюда и главный практический вопрос.

Что делает источник удобным для цитаты

Дальше практика — то, на что вы реально влияете. «Удобной для цитирования» страницу делают пять признаков. Разберём каждый с цифрами.

1. Прямой, ясный ответ — и он совпадает с вопросом

Самый сильный измеримый сигнал у ChatGPT — насколько заголовок и адрес страницы по смыслу совпадают с вопросом. Ahrefs измерил это на 1,4 миллиона запросов через косинусное сходство — способ выразить смысловую близость двух текстов числом от 0 до 1. У процитированных страниц сходство заголовка с запросом — 0,602, у не процитированных — всего 0,484.¹² А если сравнивать с самым близким из внутренних подзапросов, которые модель сама генерирует из вашего вопроса, у цитируемых страниц сходство доходит до 0,656.¹² Разрыв небольшой на вид, но устойчивый: страница, которая прямо и узнаваемо отвечает на заданный вопрос, обгоняет ту, что отвечает «где-то рядом».

Практический смысл простой. Заголовок и первый абзац должны прямо называть то, что спросил человек. Не «Наш взгляд на уборку», а «Как выбрать робот-пылесос: 5 критериев». Машина сначала видит витрину — сделайте так, чтобы на витрине лежал ровно тот товар, за которым пришли.

2. Доказательства: цифры, цитаты, ссылки

Первоисточник термина GEO доказал это экспериментально: если добавить в материал статистику, цитаты экспертов и ссылки на источники, его видимость в ответах ИИ растёт в среднем более чем на 40%.¹³ Мы уже касались этого в первой главе — но здесь кроется тонкость, важная именно для этой главы.

Универсального приёма «на все случаи» не существует — сила приёма зависит от темы. И это не общие слова: в полном тексте исследования есть таблица, которая прямо связывает домен запроса с лучшим приёмом.¹⁴ Статистика и цифры сильнее всего срабатывают в теме «Право и государство» — там это приём номер один по приросту видимости. Цитаты хорошо заходят в теме «История» — там они в тройке лучших, а авторитетный тон со ссылками на авторитеты — на втором месте.¹⁴ Вывод для вас: подбирайте доказательство под тему. Юридическая статья просит цифру, историческая или экспертная — цитату и опору на авторитет.

И один приём в том же исследовании оказался единственным, который сделал хуже, — переспам ключевыми словами. На Perplexity текст, набитый ключами, показал результат примерно на 10% хуже обычной, неоптимизированной версии.¹⁴ Запомните это как якорь: набивка ключей в мире ИИ не помогает, а вредит.

3. Извлекаемая структура

Модель не читает страницу подряд — она вынимает из неё кусок. Значит, кусок должен легко выниматься. Чем яснее структура, тем чище фрагмент.

Цифры это подтверждают. Заголовки, оформленные как вопрос, цитируются почти вдвое чаще обычных повествовательных — 18% против 8,9%.¹⁵ В сравнительных запросах в нишах софта, агентств и продуктов формат «лучшие N» занимает 43,8% цитируемых страниц — списки-подборки машина любит вынимать целиком.¹⁵ (Цифра узкая: она про конкретное подмножество сравнительных материалов, а не про долю списков во всех ответах ChatGPT.) А ещё у сильно цитируемых текстов высокая плотность именованных сущностей — конкретных названий, терминов, продуктов, брендов. У таких страниц она около 20,6% — примерно втрое-вчетверо выше, чем у обычного текста с его 5—8%.¹⁵ Вынимается конкретика, а не вода.

В Рунете та же логика. Алиса выбирает страницы, «с которых удобно собрать логичный, связный ответ»: чёткий текст, факты и классификации, разложенные по блокам.⁴ Руководитель SEO-отдела Semantica Media Оксана Зорий описывает это так: «Нейропоиск ориентируется на то, что может классифицировать: размеры, типы услуг, модели, варианты доставки, фильтры, параметры. Если на странице нет структурных опор, алгоритм вынужден искать смыслы сам — и часто выдаёт искажённые формулировки. Чем яснее логика страницы, тем чище фрагменты,

которые Алиса использует для ответа». ⁴ Проще говоря: разложите факты по полочкам сами — иначе машина разложит их за вас, и не факт, что верно.

4. Авторитет — но не тот, о котором вы думаете

«Авторитет» в мире ИИ измеряется не так, как привыкли в классическом SEO. Ahrefs проанализировал 75 000 брендов и увидел: сильнее всего на попадание в ответы ИИ влияют упоминания бренда в вебе и брендированные анкоры ссылок (анкор — это текст самой ссылки). А обратные ссылки (когда на ваш сайт ведёт ссылка с другого) — примерно втрое менее заметны для ИИ. ¹⁶ Бренды из верхней четверти по упоминаниям получают до десяти раз больше ИИ-цитирований, чем следующая четверть. ^{16 17} Отдельный расчёт Кевина Индига (growth-memo.com, приведён в материале Ahrefs) указывает в ту же сторону. ^{15 18}

Решают, стало быть, не ссылки, а то, сколько и как о вас говорят. Машине важнее, что вас называют по имени в разных местах, чем сколько ссылок ведёт на ваш сайт.

Это подтверждает и российское исследование Алисы, и оно же ломает ещё одну иллюзию. Крупный, статусный, старый сайт не получает преимущества автоматически. Размер проекта, возраст домена, известность бренда почти не коррелируют с попаданием в ИИ-блок. Зато реально усиливают шансы приземлённые вещи: чёткие контакты и реквизиты (признак реального бизнеса), конкретные цены на странице вместо «по запросу», понятная структура, содержательные характеристики. ⁴ У Google похожий фильтр — E-E-A-T (опыт, экспертность, авторитетность, доверие); с обновления сентября 2025 года эти критерии из оценочных ориентиров превратились в реальный фильтр отбора, в том числе для ИИ-ответов. ¹⁹

5. Свежесть — но только вместе с релевантностью

Свежесть — сигнал, но не абсолютный. Ahrefs разобрал 17 миллионов цитирований и увидел: ИИ-ассистенты в среднем берут контент заметно свежее, чем обычная выдача Google (примерно на четверть моложе). ²⁰ Сильнее всех молодое любит ChatGPT — его источники в среднем больше чем на год новее органических результатов Google. А Google AI Overviews, наоборот, единственная система, которая цитирует страницы даже чуть старше органики. ²⁰ Разные движки — разное отношение к возрасту.

Но вот ключевая оговорка, которая спасёт вас от глупой ошибки. Внутри одного пула найденных документов картина обратная: ChatGPT спокойно цитирует страницы, которым больше года, а иногда и старше семи лет, — при этом самые свежие страницы из того же пула часто остаются без цитаты. ^{20 12} Вывод команды Ahrefs звучит как правило: **свежесть без релевантности не работает**. Новую страницу, которая плохо совпадает по смыслу с вопросом, машина найдёт — и пройдёт мимо. Исключение — новости: там релевантность заголовков почти одинакова у всех, и тогда возраст становится решающим тай-брейком. ¹²

И над всем этим — совпадение с органическим топом

Прежде чем полировать пять признаков, помните про первый барьер. У Google AI Overviews 76,10% процитированных источников ранжируются в топ-10 обычной выдачи; лишь 14,40% вообще не входят в топ-100 по тому же запросу. ²¹ Медианная позиция самого заметного источника в ИИ-сводке — второе место в органике. ²¹ Тот же порог, что мы видели у Алисы. Но и тут автор исследования, старший контент-маркетолог Ahrefs Си Куан Онг, сам себя осаживает: связь есть, но умеренная. «Если вы на первом месте в выдаче, у вас больше шансов попасть в AI Overview, чем если бы вы были ниже. Но даже это — скорее подбрасывание монетки, чем гарантия». ²¹ Топ органики открывает дверь. За дверью решают пять признаков.

Почему ИИ взял источник А, а не Б

Теперь соберём всё в живой разбор. Идеального документированного кейса, где две почти одинаковые страницы лежат рядом и известно, почему ИИ выбрал ровно одну, в открытых исследованиях нет, — поэтому опираемся на реальный эксперимент и на статистику, а не на домыслы.

Эксперимент поставил независимый разработчик и описал на Хабре. Он прогнал 20 запросов по одной нише — по три раза каждый, всего 60 обращений — через генеративный поисковый API Яндекса и сверил, какие домены попали в цитаты, с обычной выдачей по тем же запросам.²² Два вывода вышли яркими.

Первый: подавляющее большинство цитат приходит из верхней части обычной выдачи. Нейроответ не строит рейтинг с нуля — он собирает синтез из уже релевантного пула.²² Ровно первое сито в действии.

Второй, парадоксальный: сайты самих компаний из ниши в цитаты почти не попадали. Зато туда регулярно попадали крупные тематические площадки, где про эти компании написано.²² Смоделируем «А против Б» по мотивам этого паттерна — пример иллюстративный, но опирается на реальные данные выше. Возьмём условный вопрос «какая компания лучше для установки окон». Страница А — сайт самой компании: «Мы — лидеры рынка, работаем 15 лет, индивидуальный подход». Страница Б — обзор на отраслевом портале: «Топ-7 оконных компаний города: цены, сроки, гарантии, таблица сравнения». Кого возьмёт ИИ?

Прогоните обе через наши пять признаков. У страницы А заголовок и текст маркетинговые — прямого ответа на «какая лучше» нет, смысловое совпадение с вопросом низкое. Доказательств нет — одни эпитеты. Структуры для сравнения нет — из чего вынимать кусок? Авторитет? Про саму компанию в вебе говорят мало и в основном она сама. У страницы Б всё наоборот: заголовок дословно отвечает на вопрос, внутри цифры и таблица, структура готова к извлечению, а площадку в вебе упоминают многие. Оба сита страница Б проходит, страница А спотыкается на первом же.

Отсюда парадокс, который стоит запомнить: иногда выгоднее, чтобы про вас написало стороннее издание, чем полировать собственную страницу. Дать экспертный комментарий в обзор бывает эффективнее, чем в сотый раз переписать «о компании».

И последнее. Даже большие исследования порой упираются в стену. Авторы Yext, разбирая гостиничную аномалию SearchGPT (те самые 38,1% официальных сайтов отелей), написали: причина неизвестна — может, состав обучающих данных, может, настройка поискового слоя, может, намеренное взвешивание брендов под туристические запросы.¹ Красивой теории они придумывать не стали. Механика ИИ-отбора ещё во многом чёрный ящик, и лёгкие выводы о ней обманчивы: полагайтесь на проверенные паттерны, а не на чужие уверенные догадки.

Итог главы

Источник проходит через два сита. **Первое — тип запроса:** намерение вопроса и тема задают, из какого пула ИИ вообще черпает источники, и у каждой модели свои любимые пулы. **Второе — качество внутри пула:** побеждает тот, у кого прямой ответ на заданный вопрос, доказательства под тему, извлекаемая структура, брендовый авторитет (упоминания, а не просто ссылки) и уместная свежесть.

И повторяюсь, потому что это стоит унести из главы: **для ИИ важнее всего удобство собрать из вашего материала готовый ответ.**

Сделайте сейчас одну вещь: возьмите свою главную страницу и прогоните её по пяти признакам. Отвечает ли заголовок дословно на вопрос клиента? Есть ли в тексте цифры, цитаты, ссылки? Разложены ли факты по блокам, из которых легко вынуть кусок? Говорят ли о вас по имени за пределами вашего сайта? Не устарела ли страница? Каждое «нет» — это точка роста.

Мы разобрали, как отбор устроен в общем. Но вы уже заметили: Claude и Gemini на один вопрос идут за источниками в разные места, ChatGPT любит свежее, а Google AI Overviews — нет. Почему у разных нейросетей ответы получаются такими разными — и как с этим работать, если ваши клиенты пользуются и теми, и другими, — разберём в следующей главе.

Глава 4. Почему у разных нейросетей разные ответы

Четыре ассистента — четыре разных списка источников

Прodelайте простой опыт. Задайте один и тот же вопрос — что угодно из вашей сферы, «какую компанию выбрать для...» или «что лучше для...» — сразу четырём ассистентам: ChatGPT, Perplexity, Gemini и Алисе. Ответы по смыслу могут оказаться похожими. А вот источники под ними — почти наверняка четыре разных набора.

Это не совпадение и не случайность конкретного дня. Учёный из Индианского университета Кай-Чэн Ян разобрал 366 087 цитирований из 24 069 разговоров с ИИ-поисковиками и измерил, насколько похоже разные системы ссылаются на новостные источники.¹ Внутри одной компании модели цитируют почти одинаково, а между компаниями сходство паттернов падает в разы.¹ Проще говоря: два ассистента от одного разработчика — как родные братья, а два от разных — как случайные прохожие.

Но самое интересное вскрылось в другом исследовании. Компания BrightEdge проанализировала цитаты пяти ИИ-движков по десяти отраслям и сравнила две вещи по отдельности: какие источники движки цитируют и какие бренды рекомендуют. Источники у любой пары движков пересекаются слабо — от 16% до 59%. А вот бренды, которые они советуют, совпадают заметно теснее — от 36% до 55%.²

С этого парадокса и начинается глава. Дороги, которыми ассистенты добираются до ответа, — разные шоссе. Но чаще всего эти шоссе приводят в один и тот же город. И если вы хотите, чтобы ваш бренд оказался в этом городе, придётся понять, почему шоссе разные — и как ехать по всем сразу.

Почему это ваша проблема, а не проблема разработчиков

Вся механика из прошлых глав — живой поиск, нарезка страниц на фрагменты, два сита отбора — описывала одного отдельно взятого ассистента. Но ваши клиенты не пользуются «одним ассистентом вообще». Кто-то спрашивает ChatGPT, кто-то — Алису в Яндекс-браузере, кто-то — GigaChat в приложении банка. И если вы оптимизировались под одного, это ещё ничего не говорит о том, увидят ли вас у остальных.

Эта глава отвечает на два вопроса. Первый: почему на один и тот же запрос разные нейросети выдают разные источники — что у них устроено по-разному внутри. Второй, практический: как с этим жить, если ваша аудитория размазана по нескольким платформам сразу. Единой кнопки «оптимизировать под все ИИ» не существует — это стоит принять сразу. Но есть общий фундамент, который работает везде, — и его мы в конце соберём в чек-лист.

У каждой платформы свой поисковый двигатель

Главная причина расхождений проста: **«поиск под ИИ» — это не один канал, а несколько параллельных поисковых систем с разными индексами.** У каждого ассистента свой способ добывать страницы из интернета. Разберём по очереди.

ChatGPT исторически опирается на индекс Bing — поисковую базу Microsoft. Но по данным экспериментов он умеет добирать результаты из Google, когда нужной страницы в Bing не нашлось.^{3 4} То есть даже у одного ChatGPT источник может прийти из двух разных мест в зависимости от запроса.

Perplexity устроен принципиально иначе. Это «поисковик в первую очередь»: у него собственный краулер (робот, который обходит сайты) и собственный индекс — он не полагается целиком ни на Google, ни на Bing.^{3 5} И у него сильный вес свежести — по оценке практиков, около 40% сигнала ранжирования (это наблюдение специалистов, не официальная цифра компании).⁵ Отсюда его особый характер: Perplexity любит новое и живёт в своей отдельной вселенной источников.

Gemini заземляется на обычный поиск Google и наследует его логику ранжирования.³ **Copilot** — продукт Microsoft, встроенный прямо в индекс Bing.³ Уже из этого видно, почему пулы источников расходятся: у четырёх ассистентов — по сути четыре разных поисковика под капотом.

Есть и вторая причина расхождений, поглубже. Ассистенты не ищут по буквальному тексту вашего вопроса. Они разбивают его на несколько вариаций-подзапросов, ищут по каждой, а потом сводят результаты вместе: выигрывают страницы, которые стабильно всплывают сразу по нескольким формулировкам. Инженеры называют этот приём query fan-out — «веер запросов». ³ Из-за него итоговые цитаты часто не совпадают с тем, что стоит в обычной выдаче по исходному вопросу. А поскольку «веер» у каждой системы свой, результаты у всех разные.

Насколько разные — Ahrefs измерил на 15 000 запросов, сравнив цитаты ИИ с топ-10 обычной выдачи Google. Итог короткий: пересечение низкое. Заметно совпадает с топом Google только Perplexity — 28,6%, у остальных ассистентов считанные проценты, а среднее по всем замеренным категориям — около 12%.³ С индексом Bing картина зеркальная: там впереди Copilot — продукт Microsoft, — а у Perplexity с Bing почти ничего общего, потому что у неё свой индекс.³

Запомните эту мысль: **12% — это среднее по пяти разным замерам, где Perplexity задирает планку, а не универсальная «доля пересечения».** Когда вам называют красивую единую цифру, спрашивайте, что именно с чем сравнивали. Единого канонического числа нет — есть надёжно подтверждённое направление: пересечение индексов низкое, у каждого ассистента свой набор источников.

Для сравнения — исключение, которое подтверждает правило. Google AI Overviews (те самые ИИ-сводки прямо в выдаче Google, не путать с самостоятельным ассистентом) цитируют топ-10 органики примерно в 76% случаев.³ Логично: это надстройка над самим Google. А главный сюжет для нас всё-таки российский.

Российский разрез: одна страна, три разных механизма

В Рунете расхождения ещё нагляднее, потому что три главные российские платформы устроены по-настоящему по-разному. Одна страна — три совершенно непохожих механизма отбора источников.

Алиса от Яндекса — надстройка над обычной выдачей. Это не отдельный поисковик: кандидатов в нейроответ Алиса берёт из топа органической выдачи Яндекса, и чем выше сайт стоит в топ-10, тем выше его шанс попасть в ответ.⁶ В первой главе мы видели, как быстро растёт значение места в органике Яндекса: сайты вне топ-10 почти перестали попадать в нейроответы Алисы.⁶ ⁷ Подавляющее большинство источников сегодня берётся из первой десятки обычной выдачи. Вывод для практики железный: **в Рунете попадание в ответ Алисы напрямую зависит от классического SEO под Яндекс.** Не в топе органики — считайте, вас для Алисы нет. Причём вероятность резко падает уже после пятой позиции, так что целиться стоит в топ-5.⁶

Позиция, впрочем, решает не всё. Амбассадор интернет-площадок в Поиске Яндекса Михаил Сливинский подчёркивает: попадание в ответ зависит и от качества, и от глубины экспертизы контента. «Быстрые ответы» собираются примерно из десятка источников, и там могут оказаться и крупные корпоративные сайты, и ресурсы малого бизнеса — если у них есть настоящая глубокая экспертиза по теме запроса.⁶

GigaChat от Сбера — совсем другая архитектура. У него собственный поисковый слой GigaSearch. Судя по техническому докладу разработчиков (Прохор Гладких и Александр Гаврилов, SberDevices), он построен на связке полнотекстового поиска Open Search и векторного поиска: система находит около сотни документов, затем специальная модель-переоценщик (reranker) отбирает из них три лучших, и уже поверх этих трёх генерируется ответ.⁸

Отдельный классификатор решает, фактологический ли вопрос вообще — и только тогда включает поиск. Никакой привязки к чужой поисковой выдаче: Сбер строит свой индекс сам.

DeepSeek — третий путь, самый неожиданный. Китайская модель не строит собственный поиск и не пользуется чужой готовой выдачей. По данным китайского делового издания 36kr, веб-поиск для DeepSeek обеспечивает сторонний провайдер — компания Boshai, которая обслуживает поисковым слоем около 60% ИИ-приложений в Китае.⁹ Тут нужна оговорка: роль именно Boshai подтверждена одним источником на китайском рынке, и её стоило бы перепроверить. Надёжно же подтверждено другое: единого «официального» поисковика у DeepSeek нет — при интеграции через API к нему можно подключить хоть Google, хоть Bing, хоть китайский Quark.⁹ Поиск для DeepSeek — настраиваемый компонент, а не встроена часть.

Смысл этой тройки такой. Алиса намертво завязана на SEO под Яндекс. GigaChat живёт в своём собственном индексе. DeepSeek берёт поиск со стороны. Три российских (ну, доступных в России) ассистента — три разные двери, и ключ от одной не открывает две другие.

Для российского рынка это не абстракция, мимо которой можно пройти. По опросу ВЦИОМ (октябрь 2025), с нейросетями знакома половина интернет-пользователей России, а в топе по популярности — ChatGPT (27%), YandexGPT (23%), DeepSeek (20%) и GigaChat (15%).¹⁰ Другой опрос, компании OMI (март 2026, 2328 пользователей), дал иные числа — «Алиса AI» 71%, ChatGPT 42%, GigaChat 38%, DeepSeek 25%.¹¹ Методики у двух опросов разные — один спрашивал про «самую популярную», другой про использование в принципе, — поэтому цифры между собой не сравнивайте. Но оба сходятся в главном: в топе российских нейросетей устойчиво стоят ChatGPT, Яндекс, GigaChat и DeepSeek. А это четыре разных механизма отбора источников. Одной настройкой их не накрыть.

Живой пример: почему решает поиск, а не «ум» модели

Нагляднее всего это видно в одном тесте. Редакция iphones.ru сравнила свежий GigaChat 2 Max от Сбера с топовым YandexGPT 5 Pro в Алисе на запросе про бизнес-идею — где нужны актуальные данные о площадках, ценах и марже.¹²

GigaChat честно признался, что доступа к свежим данным по теме у него нет, и содержательного ответа не дал. А YandexGPT — с доступом к живому поиску Яндекса через Алису — собрал актуальные цифры и выдал разбор.¹² Дело тут не в том, что одна модель «умнее». Дело в том, что у одной был живой слой поиска под нужную тему, а у другой в тот момент — нет. Ровно то, о чём была вся первая часть книги: не память решает, а слой поиска.

И ещё один тест, чтобы вы не идеализировали ни одну платформу. Блог М. Видео устроил «нейробойцовский клуб» — прогнал один набор задач через ChatGPT, GigaChat и Алису.¹³ На этической дилемме ответы разошлись по глубине: ChatGPT дал взвешенный, но уклончивый ответ без единственного решения, а GigaChat разобрал каждую сторону и дал конкретные рекомендации.¹³ А вот на простой логической задаче — «у Кати три сестры и три брата, сколько сестёр у брата Кати?» (верный ответ — четыре) — все три системы синхронно ошиблись и ответили «три». ¹³ Урок двойной: источники у моделей разные, но и рассуждение у них может совпасть — в том числе в одной и той же ошибке. Слепо доверять ни одной не стоит.

Почему единой стратегии «на всех сразу» нет

Соберём доказательства вместе. Что пулы источников у платформ расходятся — это не одно случайное наблюдение, а вывод сразу трёх независимых исследований с разными методиками. Академическая работа Кай-Чэн Яна,¹ коммерческое исследование BrightEdge² и замер Ahrefs³ — три разные линейки, три разных подхода, и все показывают одно: пересечение низкое, у каждого ассистента свой набор источников.

Отсюда прямой вывод: **нельзя «оптимизироваться под ИИ вообще»**. Нельзя один раз настроиться и накрыть всех. Оптимизация под ChatGPT (а значит, во многом под Bing) —

это не то же самое, что попадание к Алисе (а значит, в топ Яндекса). И это не то же самое, что видимость у Perplexity с её собственным жадным до свежести индексом.

Хуже того — под одной платформой всё может измениться внезапно и без предупреждения. Semrush три месяца, с июля по октябрь 2025 года, следил за самыми цитируемыми доменами в ИИ-ответах. И поймал резкий сдвиг: доля ответов ChatGPT со ссылкой на Reddit рухнула с почти 60% в начале августа до примерно 10% к середине сентября. Доля Википедии там же упала с ~55% до менее 20%.¹⁴ А на Google AI Mode и Perplexity за тот же период доли Reddit и Википедии остались стабильными.¹⁴ То есть потрянуло именно ChatGPT — и больше нигде.

Причину точно не установили. Глава Semrush по органике и ИИ-видимости Сергей Рогунин считает наиболее вероятной попытку OpenAI снизить перекося в сторону отдельных доменов и повысить устойчивость к манипуляциям — но это его версия, а не доказанный факт.¹⁴ Урок для вас важнее причины: то, что вчера работало на одной платформе, завтра может исчезнуть именно там — и нигде больше. Волатильность у каждой системы своя, отдельная. Ставить всё на одну площадку — строить на песке.

Платформы расходятся не только в том, какие источники берут, но и в том, насколько эти источники надёжны. Показательна одна узкая, но аккуратно поставленная проверка. Два японских врача разобрали 540 ответов Microsoft Copilot про биодобавки (запросы задавали на японском, данные собирали в 2023 году) и оценили каждую ссылку. Итог: 72,7% процитированных источников попали в категорию неverified источников — в основном нерегулируемые медицинские блоги.¹⁵ Важнейшая оговорка: это очень узкий срез — один язык, одна тема, устаревшие данные, и Copilot с тех пор менялся. Переносить эти 72,7% на весь Copilot или на другие темы нельзя. Но как иллюстрация урок ясен: доля сомнительных источников у разных платформ разная, и слепо доверять любой из них не стоит.

Что общего у всех платформ

Теперь награда за то, что мы продралась через все различия. Шоссе разные — но вспомните парадокс из начала главы. Бренды, которые движки рекомендуют, пересекаются заметно теснее источников — 36—55% против 16—59%.² Разными путями ассистенты чаще всего приходят к одним и тем же именам. А значит, под этими путями есть общий фундамент — принципы, которые работают везде.

Их четыре, и все они уже звучали в прошлых главах — соберём их вместе как общий знаменатель любой платформы.

Первое — ясность. Понятный, прямой ответ на заданный вопрос, а не «где-то рядом». Заголовок и первый абзац называют ровно то, что спросил человек. Это любят и ChatGPT, и Алиса, и все между ними: из ясного текста легче вынуть готовый фрагмент.

Второе — доказательства. Цифры, цитаты экспертов, ссылки на источники. Напомню: та самая научная работа, которая ввела термин GEO, доказала это экспериментально — с ними видимость в ответах ИИ растёт. Ни одна платформа не наказывает за факты — все их ценят.

Третье — структура, из которой легко вынуть кусок. Модель не читает страницу подряд, она вырезает фрагмент. Разбейте информацию на блоки по смыслу. Это универсально: и Open Search у Сбера, и индекс Bing под ChatGPT одинаково легче работают с чистой структурой.

Четвёртое — брендовые упоминания. Не столько ссылки на ваш сайт, сколько то, сколько и как о вас говорят по имени в вебе. Именно здесь, судя по данным BrightEdge, и сходятся разные движки: узнаваемый бренд всплывает у всех.²

Есть и слой различий, который стоит держать в голове как карту приоритетов. BrightEdge разложил цитаты движков по типам источников, и разброс огромный. Доля цитат из «авторитетных» доменов (государственных, образовательных, институциональных) — от 10% у Google AI Overviews до 26% у Gemini.² А доля пользовательского контента — форумов, отзывов, обсуждений — колеблется от 0,2% у Gemini до 18% у Google AI Overviews, то есть в девяносто

раз.² Это подсказка: под Gemini делайте ставку на солидные институциональные источники, под платформы, любящие UGC, — на присутствие в обсуждениях и отзывах.

Как работать сразу со всеми: практика

Теперь — что со всем этим делать в понедельник утром. Единой кнопки нет, поэтому работаем не «под все ИИ вообще», а параллельно по нескольким дорожкам. Порядок действий такой.

Шаг 1. Определите, где ваша аудитория. Не распыляйтесь на все платформы сразу. Если ваши клиенты в России — на первом месте почти всегда Алиса и Яндекс, дальше ChatGPT, GigaChat, DeepSeek (сверьтесь с данными по популярности: ChatGPT, Яндекс, GigaChat и DeepSeek стабильно в топе).^{10 11} Если аудитория глобальная — ChatGPT, Perplexity, Gemini. Выберите три-четыре платформы, которые реально важны вам, и работайте прицельно под них.

Шаг 2. Заложите общий фундамент — он окупается на всех. Прежде чем подстраиваться под каждую платформу, доведите контент до четырёх общих принципов: ясность, доказательства, извлекаемая структура, брендовые упоминания. Это база, которая повышает шансы везде одновременно. С неё начинают, а не заканчивают.

Шаг 3. Для Рунета — обязательно классическое SEO под Яндекс. Раз Алиса берёт до 90% источников из топ-10 органики, никакая «оптимизация под ИИ» не заменит вам позиций в обычной выдаче Яндекса.⁶ Для российского рынка это не отдельная задача, а фундамент под всё остальное: сначала топ Яндекса, потом уже разговор про нейроответ. Целитесь в топ-5 — после пятой позиции шанс резко падает.⁶

Шаг 4. Разведите глобальную и российскую дорожки. Пулы источников у глобальных и российских платформ пересекаются слабо, поэтому и площадки под них разные. Одна стратегия — под русскоязычные платформы, другая — под глобальные. Не пытайтесь сделать один материал, который «зайдёт всем»: он рискует не зайти никому.

Шаг 5. Учитывайте характер платформы. Perplexity жаден до свежего — обновляйте материалы и давайте новое.⁵ Gemini тянется к институциональным источникам — стоит присутствовать на солидных доменах.² Платформы, любящие обсуждения и отзывы, требуют присутствия в UGC.² Это тонкая настройка поверх общего фундамента, а не вместо него.

Шаг 6. Следите и не привыкайте. Волатильность — норма. Доля Reddit в ChatGPT упала в шесть раз за месяц, и никто не предупреждал.¹⁴ Проверяйте раз в месяц-другой, кого называют ваши целевые ассистенты по ключевым вопросам вашей сферы, и будьте готовы перестроиться. Присутствуйте сразу в нескольких местах — это и есть страховка от таких скачков.

Итог главы

Разные нейросети дают разные ответы, потому что под капотом у них разные поисковые двигатели. ChatGPT опирается на Bing (и добывает Google), Perplexity — на собственный жадный до свежести индекс, Gemini — на Google, Copilot — на Bing. В Рунете три платформы устроены совсем по-разному: Алиса — надстройка над топом Яндекса, GigaChat — собственный GigaSearch, DeepSeek — покупной поиск со стороны. Пересечение источников между всеми ними низкое — это подтвердили три независимых исследования.

Если запомнить одну мысль из этой главы, пусть будет эта: **единой кнопки «оптимизировать под все ИИ» не существует — но общие принципы работают на всех платформах.** Ясность, доказательства, извлекаемая структура и брендовые упоминания повышают ваши шансы везде сразу. А в Рунете к ним обязательно добавляется база — классическое SEO под Яндекс.

Сделайте сейчас одну вещь: выпишите три-четыре платформы, которыми реально пользуются ваши клиенты, и напротив каждой пометьте, откуда она берёт источники (индекс Bing? топ Яндекса? собственный индекс?). Это ваша карта дорог — по ней и будете строить работу.

Мы разобрались, как машина принимает решение и почему у разных ассистентов оно выходит разным. Осталось свести всю эту механику к одному практическому вопросу: на что из неё вы реально можете повлиять, а что вне вашей власти? В следующей главе — последней в первой части — соберём карту ваших рычагов. А уже за ней начнётся разговор о рынке: сколько внимания и денег уже перетекло в ответы нейросетей и почему окно возможностей открыто именно сейчас.

Глава 5. Ваши рычаги влияния

Четыре пятых — в ваших руках

Начнём с цифры, которая должна вас обрадовать. Из всего, что нейросети говорят о брендах, около 79,9% приходит из источников, которыми бренд управляет сам, — со своего сайта и со своих карточек на чужих площадках. Это данные компании Yext за первый квартал 2026 года: она отслеживает, откуда ИИ-ассистенты берут информацию о бизнесах, и по когорте из 770 брендов насчитала именно столько «управляемых» источников.¹

Почти четыре пятых того, что ChatGPT, Gemini, Perplexity и Claude рассказывают про компанию, — это не приговор судьбы и не случайный шум из интернета. Это то, на что вы можете влиять. Не всё, но большую часть.

А теперь та же цифра с другой стороны. Кварталом раньше, в конце 2025-го, эта доля внутри той же группы брендов была 88,2%.¹ За один квартал управляемость просела на восемь с лишним пунктов. Причём авторы отчёта не считают это потерей влияния — просто модели стали заглядывать в более разнообразные источники.¹ Контроль есть — но он не бетонный. Он плавает от квартала к кварталу.

В этом зазоре — между «четыре пятых в ваших руках» и «доля тает» — и живёт вся глава. Первые четыре главы объясняли, как машина принимает решение. Эта отвечает на единственный практический вопрос, ради которого вы всё это читали: **на что из разобранной механики вы реально можете повлиять — и за какие рычаги браться в первую очередь.**

Что ваше, а что нет

Напомню формулу, на которой всё держится: **память из обучения переписать нельзя, а слой живого поиска — открыть.**

Внутри модели не залезть и не поправить то, что она «думает» о вашей отрасли, — для бренда этот слой закрыт наглухо. А вот в момент ответа решается, увидит ли ассистент ваш сайт, вашу карточку, обзор про вас на отраслевом портале, — и на это влиять можно. Именно здесь вы либо попадаете в ответ, либо нет.

Но «можно влиять» не значит «контролируешь всё». Контроль бывает разной силы, и это важно разложить по полочкам. Та же Yext предложила удобную рамку из четырёх уровней.¹ Переведём её на бытовой язык.

Свой сайт — это ваш дом. Здесь вы хозяин: сами решаете, что написать, как разложить, когда обновить. Полный контроль.

Карточка на площадке — съёмная квартира. Данные о бизнесе вы меняете сами, но живёте по правилам хозяина: структуру страницы задаёт площадка, не вы. Контроль средний.

Отзывы и соцсети — общий двор. Вы можете выйти и поговорить, ответить на отзыв, включиться в обсуждение. Но приказывать соседям, что о вас писать, не можете. Контроль ограниченный.

Новости, форумы, госсайты — чужая улица. Тут вас могут упомянуть, а могут не заметить. Прямого влияния — ноль.

«Балл управляемости», те самые 79,9%, — это сумма первых двух: дома и съёмной квартиры. То есть почти всё, что ИИ говорит о брендах, приходит с территории, где у бизнеса есть хотя бы частичная власть. Остальное — двор и улица, где вы гость.

Кстати, о цифрах. Раньше в разговорах про управляемость мелькало число 86%. Это не противоречие свежим 79,9%. Осенью 2025-го Yext посчитала долю управляемых источников на 6,8 миллиона цитирований и получила 86% — но по трём моделям: Gemini, ChatGPT и Perplexity.² А 88,2% и 79,9% — это уже другой срез: четыре модели (добавили Claude от Anthropic), другая когорта брендов, другие кварталы.¹ Одна и та же серия исследований, разные замеры. Не спорьте с этими числами между собой — читайте их как кадры одного фильма,

снятые в разные месяцы. И общий сюжет фильма один: большую часть источников бренд контролирует, но сама эта доля живая и подвижная.

Карта из четырёх рычагов

Теперь — главное. Всё, на что вы влияете в слое живого поиска, сводится к четырём рычагам. Не к сорока приёмам, которые пугают новичка, а к четырём понятным направлениям. Разберём каждый: что он даёт и почему работает.

Рычаг 1. Контент — что и как написано

Это то, что вы пишете и как вы это подаёте. Ясный прямой ответ на вопрос, доказательства, структура, из которой легко вынуть готовый кусок.

Это рычаг, а не благое пожелание: эффект измерен. Первоисточник термина GEO — научная работа, с которой началась вся дисциплина, — доказал экспериментально: если добавить в материал статистику, цитаты экспертов и ссылки на источники, его видимость в ответах ИИ растёт в среднем более чем на 40%.³ Это не «улучшает восприятие», это конкретные проценты за конкретные действия.

Причём рычаг работает неравномерно, и это надо знать заранее. В том же исследовании приём «добавить ссылки на источники» поднял видимость на 115% для страниц, стоявших на пятом месте выдачи, — и одновременно уронил её на 30% для тех, кто был на первом.³ Вывод простой: контент сильнее всего вытягивает «середняков». Если вы пока не в топе, это ваш главный рычаг.

У контента есть и обратная сторона — единственный проверенный приём, который делает хуже. Переспам ключевыми словами на Perplexity показал результат примерно на 10% хуже обычного текста.³ Запомните: в мире ИИ набивка ключей не помогает, а вредит.

Каждый из этих приёмов и то, как писать материалы под нейросети, мы ещё разберём по отдельности и с примерами.

Рычаг 2. Площадки — где размещено

Одного своего сайта мало. Ассистенты берут источники отовсюду, и важно, на каких площадках вы присутствуете.

Начать стоит с неожиданного. Экзотические каналы не решают — решает попадание в обычный веб-поиск. Ahrefs разобрал 1,4 миллиона запросов к ChatGPT и увидел: 88,46% всех процитированных ссылок пришли через канал обычного веб-поиска.⁴ То есть базовый вопрос не «как попасть в какой-то особый ИИ-индекс», а «виден ли ваш сайт в обычном поиске».

В Рунете это правило доведено до предела. Алиса от Яндекса — не отдельный поисковик, а надстройка над обычной выдачей: почти все источники её нейроответов берутся из топ-10 органики Яндекса, а сайты вне первой десятки, как вы уже знаете, за полгода из ответов практически исчезли.⁵ Для российского рынка вывод железный: «свой сайт в топе Яндекса» — не один из вариантов, а обязательный рычаг. Не в топе — считайте, для Алисы вас нет.

И ещё одно про площадки: работать надо сразу на нескольких направлениях. Пулы источников у разных ассистентов пересекаются слабо — от 16% до 59% по замерам BrightEdge.⁶ Разными дорогами идут. А значит, ставить всё на одну площадку — рискованно: на другой платформе вас там просто нет.

Где именно размещаться, какие внешние источники выбрать под свою нишу и свой рынок, как на них попасть — к этому мы ещё вернёмся.

Рычаг 3. Авторитет бренда — упоминания, а не ссылки

Это самый недооценённый рычаг, и он же ломает главную привычку старого SEO. В классическом продвижении годами молились на обратные ссылки. В мире ИИ решает другое.

Ahrefs проанализировал 75 000 брендов и нашёл: сильнейший предиктор попадания в ответы ИИ — это упоминания бренда в вебе (корреляция 0,664) и брендированные анкоры ссылок (0,527), а вовсе не обратные ссылки как таковые (всего 0,218).⁷ Разрыв огромный: бренды из верхней четверти по упоминаниям получают до десяти раз больше ИИ-цитирова-

ний, чем следующая четверть.⁷ Отдельный расчёт аналитика Кевина Индига даёт близкую по духу цифру — корреляция брендового авторитета с цитируемостью около 0,334.⁸ Это оценка одного специалиста, которая кочует по вторичным источникам, а не независимое многократное подтверждение, — но направление у неё то же.

Разница тонкая, но принципиальная: машине важнее, что вас называют по имени в разных местах, чем сколько ссылок ведёт на ваш сайт. Не «сколько на меня сослались», а «сколько обо мне говорят».

Этот рычаг разветвляется сразу в несколько сторон. Узнаваемое имя, единообразные данные о компании, связка с Wikipedia и Wikidata — техническая опора авторитета. Упоминания в независимых медиа — его топливо.

Сущности бренда, Wikipedia и Wikidata, earned media и PR — этот рычаг тянется в контент, площадки и репутацию одновременно, и к каждой из этих опор мы ещё вернёмся.

Рычаг 4. Техника сайта — доступ и извлекаемость

Самый скучный на вид рычаг — и тот, где новички чаще всего теряют силы впустую. Модных пустышек здесь больше, чем во всех остальных, вместе взятых.

Начнём с того, чего делать не надо. Есть соблазн поставить на сайт файл lms.txt — короткое описание сайта специально для ИИ-моделей — и считать дело сделанным. Эффекта, однако, не видно. Trakkr Research просканировал 37 894 домена и не нашёл практически значимого преимущества в цитируемости у сайтов с этим файлом: связь на грани нуля (в теле статьи $r=0,85$, на полной выборке коэффициент $r=-0,065$ — то есть эффекта, по сути, нет).⁹ Сам файл при этом стоит лишь на 13,3% доменов, а среди топ-50 самых цитируемых сайтов — всего у 6%.⁹ То есть его чаще ставят те, кто хочет пробиться в топ, а не те, кто там уже есть. Независимое исследование SE Ranking на 300 000 доменов пришло к тому же выводу: влияния на частоту цитирования нет.¹⁰ Цена входа у lms.txt низкая, поставить его не грех — но не ждите от него чуда и не тратьте на него силы как на что-то решающее.

Похожая история с микроразметкой Schema.org. Контролируемый эксперимент Ahrefs — 1885 страниц с разметкой против 4000 без неё, восемь месяцев наблюдений — не нашёл значимого прироста цитируемости, а для Google AI Overviews она даже чуть снизилась.¹¹ Опять же: разметка полезна для обычного поиска, но как рычаг ИИ-цитируемости она недоказана.

А вот что в технике действительно важно — доступ ИИ-краулеров к вашему сайту. И тут есть тонкость, которую путают почти все.

У крупных провайдеров боты разведены на два типа. Одни ходят по сайтам, чтобы собирать данные для обучения моделей. Другие — чтобы формировать живые ответы прямо сейчас. У OpenAI это GPTBot (обучение) и отдельно OAI-SearchBot (живой поиск). У Anthropic — ClaudeBot (обучение) и Claude-SearchBot (поиск).¹²

Тут легко увидеть противоречие с тем, что мы уже разобрали: память модели застыла в момент обучения, у неё есть отсечка по времени, и внутрь неё не попасть — а боты-обучатели, выходит, всё же ходят по сайтам и собирают данные. Противоречия нет. Боты-обучатели действительно собирают контент, но всё собранное идёт в редкое большое переобучение будущих версий модели — не про вас лично и не в сегодняшний ответ. Поэтому на вашу нынешнюю цитируемость они не влияют вовсе: живые ответы формирует только слой поиска.

Отсюда и практический вывод. Если вы закрыли в настройках сайта ботов-обучателей — на цитируемость это не влияет, они в живых ответах не участвуют. А вот если закрыли поисковых ботов вроде OAI-SearchBot или PerplexityBot, вы своими руками вычеркнули сайт из ответов: именно они собирают то, что ассистент показывает пользователю.¹²

И здесь — самая коварная ловушка. Часть сайтов блокирует нужных ботов не нарочно. По оценке одного анализа на примерно 3000 сайтов, около 27% случайно блокируют хотя бы одного крупного ИИ-краулера — причём чаще не в своих настройках, а на уровне CDN (это сеть доставки контента) или защитного экрана, о котором владелец даже не знает.¹³ К цифре

27% нужна поправка: она восходит к одному замеру (в основном B2B и интернет-магазины США и Британии) и широко разошлась по блогам без независимого повтора — так что берите её как сигнал, а не как закон. Но сам механизм реален: например, Cloudflare, через сеть которого идёт около 20% мирового трафика, с 1 июля 2025 года по умолчанию блокирует ИИ-краулеров для новых доменов.¹⁴ Вы могли ничего не запрещать — а бот всё равно не проходит.

К переделке сайта под ИИ мы ещё вернёмся подробно, но главное правило стоит запомнить уже сейчас: разметка и llms.txt — низкая цена входа без доказанного эффекта, а вот пустить поисковых ботов и снять случайные блокировки — реальный технический минимум. Не гонитесь за модными файлами. Сначала проверьте, что вас вообще пускают.

Что вне вашей власти вообще

Карта будет неполной без того, за что браться бесполезно. Четыре рычага — ваша территория. А вот здесь она заканчивается.

Память модели из обучения переписать нельзя — мы уже об этом сказали. Дальше — архитектура платформ: какой поисковый индекс использует ассистент, вы не выбираете. И волатильность алгоритмов: в прошлой главе мы видели падение цитирований Reddit в ChatGPT — оно случилось без предупреждения и без всякой причины со стороны брендов.¹⁵ Наконец, «чужая улица» из рамки контроля — новости, форумы, госсайты: вас там могут упомянуть, но управлять этим вы не можете.

Вывод не пессимистичный, а трезвый. Не тратьте силы там, где повлиять нельзя. Вкладывайтесь в четыре рычага — и не тревожьтесь о том, что вне них.

Это, кстати, объясняет неприятный сюрприз, с которым сталкиваются многие: бренд может быть звездой в ChatGPT и полной невидимкой в Алисе.¹⁶ Не потому, что кто-то виноват, а потому, что рычаги нужно тянуть отдельно под каждую платформу — и в паре мест что-то от вас не зависит вовсе. Отсюда и наш прицел: тянуть то, что тянется, и на всех важных направлениях сразу.

Итог главы

Соберём всю первую часть в одну картину. У ИИ два «мозга»: память из обучения, закрытая для вас, и живой поиск, открытый для влияния. Почти четыре пятых того, что нейросети говорят о брендах, приходит из управляемых источников — но эта доля живая и от квартала к кварталу колеблется. А всё ваше влияние сводится к четырём рычагам: **контент** (что и как написано), **площадки** (где размещено), **авторитет бренда** (упоминания, а не ссылки) и **техника сайта** (пускают ли вас боты живого поиска).

Если запомнить из этой главы одну мысль, пусть будет эта: **вы не управляете тем, что ИИ помнит, но во многом управляете тем, что он находит, — и делаете это через четыре рычага, не больше.**

Сделайте сейчас одну вещь: возьмите лист и разбейте его на четыре колонки — контент, площадки, авторитет, техника. В каждую впишите по одному честному ответу: что у вас с этим сегодня? Где-то будет пусто, где-то густо. Это и есть ваша стартовая карта — с ней вы входите в практическую часть книги.

А дальше — короткий крюк перед практикой. Прежде чем засучить рукава, стоит понять, ради чего. Следующая часть — про рынок: сколько внимания и денег уже перетекло в ответы нейросетей, кто на этом уже выигрывает и почему окно возможностей открыто именно сейчас, а не «когда-нибудь потом». А сразу за разговором о рынке начинается практика — приёмы, площадки, правила письма, — то самое «что делать в понедельник утром», к которому вела вся механика. Механику вы теперь знаете. Пора смотреть, чего она стоит — и как ею пользоваться.

Часть II. Почему именно сейчас

Глава 6. Масштаб сдвига: сколько внимания и денег уже там

Меньше воды — больше нефти

За первые четыре месяца 2026 года 68% всех запросов в Google в США закончились ничем. Точнее — без единого клика. Человек ввёл вопрос, получил ответ прямо в выдаче и не перешёл ни на один сайт. Меньше трети запросов, всего 31,99%, довели пользователя до чьей-то страницы. Это данные Rand Fishkin из SparkToro — он отслеживает «поиск без клика» с 2019 года, а свежую цифру взял с панели Similarweb, десктоп и мобайл вместе.¹

Для владельца бизнеса это звучит как приговор: если люди не кликают, откуда брать клиентов? Но вот вторая цифра, и она переворачивает картину. В мае 2026 года посетители, пришедшие на розничные сайты США из ИИ-сервисов, покупали на 54% чаще обычных, а выручка с каждого визита была на 53% выше. Это данные Adobe, собранные на реальных транзакциях — больше триллиона визитов на американские магазины.²

Сложите две цифры вместе — и получите главный сюжет этой главы. Трафика с поиска стало меньше. Зато тот, что доходит через ИИ, — плотнее, теплее, дороже. Воды меньше, а нефти больше.

Первая часть книги объясняла механику — как машина находит источники, собирает ответ и на что тут можно влиять. Теперь разговор меняется. Не «как оно работает», а «сколько это уже стоит и куда всё движется». Сколько людей прямо сейчас выбирают и покупают через нейросети. Куда перетекают трафик и деньги. И что весь этот сдвиг значит для бренда вашего размера — маленького, среднего или большого.

Два предупреждения, прежде чем начнём сыпать числами. Первое: эта сфера меняется буквально помесечно, и разные исследователи считают по-разному. Поэтому у каждой цифры здесь стоит дата — читайте её как срок годности. Цифра без даты в этой теме бесполезна. Второе: дальше будет много опросов, и у части их авторов и заказчиков есть свой коммерческий интерес. Держите это в уме по умолчанию — отдельно я буду оговаривать только самые спорные случаи.

Сколько людей уже покупают с ИИ

Начнём с масштаба, а не с прогнозов. Опрос Exploding Topics и Semrush на 1009 американцах в апреле 2026 года показал: 77,6% потребителей США за последние полгода хоть раз использовали ИИ для покупок, а 43% — еженедельно или чаще.³ Самый ходовой инструмент — ChatGPT: им пользуются 77,56% из тех, кто вообще применяет ИИ в шопинге. Следом Gemini от Google с 58%, Claude — меньше 20%.³

Что именно люди покупают с помощью ИИ? Чаще всего — одежду и технику, это лидеры. Продукты питания через ИИ выбирали 44,6% опрошенных. А вот мебель (29,6%) и ювелирные изделия (28%) — среди самых редких категорий. Но покупка — это финал. Куда чаще ИИ используют раньше: изучить товар (68,5% из тех, кто применяет ИИ для шопинга) и найти цену получше (55%).³ То есть нейросеть чаще работает советчиком на входе в выбор, чем кассой на выходе.

И советчик этот влияет по-настоящему. 68,64% пользователей вспоминают: ИИ прямо подтолкнул их к покупке, которую иначе они бы не сделали. Треть — «много раз». ³ Такая статистика уже не про фоновую игрушку. ИИ влияет на то, куда люди несут деньги.

Есть и более узкий, более осторожный замер. NIQ (бывшая NielsenIQ) опросила около 500 американцев про последний месяц — не полгода — и получила 42%, использовавших хотя бы один ИИ-инструмент для покупок; сам опрос провели в начале 2026 года, а опубликовали в мае.⁴ Разница с 77,6% не противоречие, а разные окна: за полгода людей набирается больше, чем за месяц. И характерная деталь: полностью автономному ИИ-агенту, который сам оформит заказ, покупке доверили лишь 5%.⁴

Здесь и обнажается главный нерв темы — зазор между «пользуюсь» и «доверяю». В том же апрельском опросе выяснилось: чаще всего американцы не готовы дать ИИ потратить без спроса вообще ни доллара — самый частый ответ (мода) оказался нулевым. Но медианный потолок автономной траты — 50 долларов: у половины респондентов планка проходит не выше этой суммы. То есть типичный человек не пускает ИИ к кошельку совсем, а «средне-осторожный» отпускает ему максимум полсотни долларов. Треть респондентов вообще не дали бы ИИ купить что-либо самостоятельно.³ И 86% онлайн-покупателей, изучавших товар через ИИ, всё равно перепроверяли рекомендацию в другом месте, прежде чем расстаться с деньгами.³ Люди охотно спрашивают совета — и осторожничают, когда доходит до оплаты.

В России картина та же по форме, только рынок моложе и цифры скромнее. Совместное исследование Яндекса и «РБК Исследование рынков» (1850 покупателей 18—54 лет, март 2026 года) насчитало 35% интернет-покупателей, которые используют нейросети для выбора и заказа товаров. Активнее всех молодёжь — 49% в возрасте 18—24 лет. Топ-категории знакомые: одежда, обувь и аксессуары дают 37% запросов, косметика и парфюмерия — 28—29%.⁵ Аналитики к тому же дали прогноз (именно прогноз): доля может вырасти с 35% до 74%, особенно быстро в регионах.⁵

Другой российский замер вскрывает ту самую стадию, на которой мы сейчас. Опрос AIMonitor.pro на 2000 респондентах (конец июня 2026 года): 45% россиян хотя бы раз выбрали товар по совету ИИ, но регулярно, несколько раз в неделю, это делают только 5%. А 21% попробовали один раз — и не вернулись.⁶ Запомните эту рамку: рынок сейчас в стадии массового первого касания, а не устойчивой привычки. Люди пробуют. Останутся не все.

Свежее исследование «Ашманов и партнёры» (начало июля 2026 года) добавляет, что движение всё же идёт вверх: 69% опрошенных обращаются к нейросетям при выборе товаров чаще, чем год назад, а самый ходовой канал у россиян — поиск Яндекса с Алисой, им пользуются 74%.⁷ Российская картина по доверию ещё резче. По данным опроса i-Media (более 10 000 респондентов, июнь 2026 года), готовы сразу оплатить товар по совету ИИ, не перепроверяя, лишь 3,9%. Остальные идут сверять: половина смотрит цены на маркетплейсах, половина ищет отзывы живых людей. 46% прямо подозревают, что ИИ подсуживает тем, кто заплатил.⁸

Сложите западные и российские данные — и увидите один и тот же рисунок. Использование обгоняет доверие. И там, и там.

Отдельная история — B2B, продажи бизнеса бизнесу. Здесь ИИ проник даже глубже. По опросу Forrester (2026 Buyers' Journey Survey, около 18 000 корпоративных закупщиков по миру), 94% B2B-покупателей использовали ИИ на последнем цикле закупки, и ИИ-ответы названы источником номер один для исследования поставщиков.⁹ Оговорка: первоисточник закрытый, полная методология не опубликована. Но и тут финальное слово остаётся за человеком: по данным Gartner (май 2026 года), 69% B2B-покупателей всё равно перепроверяют ИИ-подсказки у живого продавца.¹⁰ Даже в мире корпоративных сделок машина — мощный советчик, а не последняя инстанция.

Почему нулевой клик меняет экономику

Теперь вернёмся к цифре, с которой начали, и разберём, что за ней стоит. «Нулевой клик» — это когда человек получает ответ прямо в выдаче и никуда не переходит. В США таких запросов к Google за январь—апрель 2026 года было 68,01%.¹ Для сравнения: в 2024-м

— 60,45%, а в 2019-м, когда Fishkin начинал считать, — 49%.¹ И тренд ускоряется: плюс семь с половиной пунктов за два года — самый быстрый рывок за десятилетие.

Когда в выдаче появляется ИИ-ответ (в Google это блок AI Overviews), доля «без клика» поднимается ещё выше — в отдельных срезах до 83%, а кликабельность обычных ссылок падает почти на 60%, по данным Ahrefs.¹ Разброс тут большой: у разных исследователей выходит от 58% до 83% — в зависимости от страны, типа запроса и методики подсчёта. Это не ошибка, а реальная фрагментация — единой «одной правильной цифры» пока нет. Поэтому опорной берём измеренные Fishkin 68% и держим в голове, что в самых «ИИ-насыщенных» срезах доходит и до 83%.

Что это делает с людьми, отдельно померила Pew Research — некоммерческий центр с прозрачной методикой. Летом 2025 года они отследили реальную браузерную активность 900 американцев. Когда в выдаче была ИИ-сводка, по обычной ссылке кликали в 8% визитов. Когда сводки не было — в 15%. Почти вдвое реже.¹¹ Ссылку внутри самой ИИ-сводки открывали и вовсе в 1% случаев. Люди чаще просто закрывали поиск, получив ответ, и уходили.

Чем это обернулось для сайтов, тоже посчитали. По панели Ahrefs из более чем 75 000 доменов доля трафика, которую Google отдаёт на сторонние сайты, за год (с июня 2025 по май 2026) упала на 8 процентных пунктов — а в относительном выражении это примерно минус 22%.¹ Причём выборка тут не средняя, а смещённая в плюс: это домены, которые сами активно занимаются продвижением. У обычного сайта падение может быть и резче.

За сухими процентами — живые люди. Николя Булиан ведёт независимый блог AllAboutBerlin с 2017 года. В мае 2026-го он опубликовал пост о том, что после массового внедрения AI Overviews потерял 75% трафика — и прямо на странице зачеркнул прежнюю оценку «70%», обновив её на более свежую и горькую.¹² Он описывает новую выдачу так: сначала реклама, потом ИИ-ответ, обученный в том числе на его же материалах, и только потом — ссылка на его сайт. Содержать проект на оставшейся четверти посетителей стало тяжело. Это не абстрактный тренд. Это один человек, у которого нейросеть забрала три четверти читателей, пересказав его же работу.

Но — и это ключевой поворот — нулевой клик не равно потеря влияния. Здесь Fishkin делает важную оговорку: трафик может падать, а выручка расти. Сайт по-прежнему влияет на ответ ИИ через то, что индексируется и цитируется, — просто теперь без прямого перехода к вам.¹ Механику этого мы уже разобрали в первой части, повторяться не будем. Держите вывод: исчезновение клика — это не исчезновение вас из картины. Это смена самой картины.

И тут нельзя обойти самый громкий прогноз в отрасли. В феврале 2024 года Gartner предсказал: объём традиционного поискового трафика упадёт на 25% к 2026 году из-за замещения нейросетями и ИИ-агентами.¹³ Цифру эту цитируют повсюду, поэтому проговорю чётко: это прогноз, а не измеренный факт. Причём спорный: его раскритиковали почти сразу после публикации. Search Engine Land — в тот же день, назвав сценарий слабо подкреплённым.¹⁴ BigTechnology через месяц взяла интервью у самого автора прогноза, и тот признал, что сценарий во многом гипотетический.¹⁵ А Datos на клиентских данных показала, что доля Google за год почти не сдвинулась.¹⁶ Важно различать две вещи. Падение объёма поисковых запросов на четверть — недоказанный прогноз. А вот падение числа переходов на сайты — реальный, задокументированный факт (те самые данные SparkToro и Ahrefs выше). Их легко перепутать, но это не одно и то же.

Куда перетекают трафик и деньги

Раз старый трафик тает, посмотрим, куда натекает новый. И здесь цифры почти неприличные.

По данным Adobe, реферальный трафик из ИИ-сервисов на розничные сайты США за I квартал 2026 года вырос на 393% год к году. В праздничный сезон конца 2025-го рост доходил

до 693%.¹⁷ К маю 2026-го рост «остыл» до 138% год к году — но это не спад, а естественное замедление: база стала огромной, и от неё уже трудно расти в разы.²

Причём трафик этот и большой, и качественный — с каждым месяцем всё качественнее. В марте 2026 года посетители из ИИ конвертировались на 42% лучше обычных¹⁸, в мае — уже на 54%.² Разворот произошёл стремительный: ещё в марте 2025 года ИИ-трафик конвертировался на 38% ХУЖЕ обычного. За год — поворот почти на 80 процентных пунктов. Выручка с визита выросла на 37% в марте 2026-го¹⁸ и на 53% в мае.² Люди из ИИ дольше сидят на сайте и смотрят больше страниц. Меньше воды, больше нефти — та же фраза, только на языке денег.

В России тот же ветер дует не слабее. Аналитическая компания Digital Budget посчитала: трафик из нейросетей на российские сайты вырос более чем в девять раз за январь—октябрь 2025 года год к году. Эти данные привёл «Коммерсантъ» в ноябре 2025-го.¹⁹ А в начале 2026 года трафик от ИИ на российских ретейлеров и маркетплейсы прибавил ещё в 1,7 раза за год.²⁰ За общими процентами — конкретные компании. У Ozon трафик от ИИ-сервисов в апреле—мае 2026-го вырос вчетверо против осени 2025-го. У «М. Видео» доля ИИ-переходов поднялась на 29% за первые четыре месяца года. Сам Яндекс сообщил о росте трафика от Алисы AI на 159% в 2026 году, а к его протоколу продаж прямо через нейросеть подключились более 1600 интернет-магазинов.²⁰

Картина складывается одна и для мира, и для России: старая труба сужается, новая расширяется — и течёт по ней уже не вода, а нефть.

Что это значит для брендов разного размера

Теперь разговор, ради которого всё и затевалось. Что этот сдвиг значит лично для вас — зависит от того, какого вы размера.

Начну с того, что бьёт по всем одинаково, независимо от бюджета. Adobe оценила «ИИ-видимость» розничных сайтов — какую долю контента реально может прочитать нейросеть. У главных страниц в среднем читаемо 75%, то есть четверть невидима для ИИ. У страниц отдельных товаров — всего 66%. А абсолютный разрыв между лучшими сайтами отрасли (82,5%) и худшими (54,2%) — 28,3 процентного пункта.¹⁷ (Само Adobe называет этот разрыв «52%», но это относительная разница: 28,3 пункта от базы в 54,2% и дают те самые ≈52%.) Важно тут другое: этот разрыв не про размер компании. Это про конкретную техническую работу над контентом, а не про толщину кошелька. Огромный ретейлер может быть для ИИ полуслепым, а маленький — вылизанным. Значит, у малого бизнеса тут нет врождённого проигрыша.

А кто вообще выигрывает — большие или маленькие? Однозначного ответа наука пока не даёт.

Есть сигналы в пользу малых. Одно из независимых исследований нашло, что домен-авторитет — тот самый показатель «веса» сайта, на который годами молилось SEO, — почти не предсказывает попадание в ИИ-ответы: корреляция всего 0,18, то есть авторитет домена объясняет около 3% того, процитирует вас ИИ или нет. Куда сильнее коррелируют сигналы реальной экспертности контента.²¹ Источник тут не безупречен: это анализ SEO-компаний, мнение практиков, а не рецензируемая наука. Но направление любопытное — большой бренд больше не покупает цитируемость одним лишь размером.

И тут же — противоположный сигнал, в пользу больших. Мы уже видели в пятой главе: масштаб упоминаний бренда в вебе всё ещё сильно определяет, как часто его цитирует ИИ. У крупных игроков накоплено присутствие в интернете, которого у новичка просто нет. Данные тянут в обе стороны: аргументы есть и «за маленьких», и «за больших», а строгого приговора нет.

Из этого противоречия и растёт единственный практический вывод, который выдерживает проверку. Маленький специализированный бренд может обойти большого генерического — но не «вообще», а на узкой теме, где у него есть настоящая экспертиза. Практики к этому добавляют наблюдение: небольшие команды публикуются быстрее и глубже копают нишевые

вопросы, пока крупные компании согласовывают контент неделями. Прямой сравнительной статистики «малый против крупного» в открытом доступе нет — это мнение, а не измеренная истина.

На практике это выглядит так. Mentimeter — сервис для интерактивных презентаций и опросов, средний по размеру, но растущий SaaS-продукт. Агентство Siege Media взялось перепаковать его материалы так, чтобы их охотнее брали нейросети: разбили сплошной текст на таблицы, вынесли чёткие определения в первую строку каждого раздела, добавили честные сравнения с другими продуктами. По отчёту самого агентства, после этого Mentimeter получил 124 000 переходов из ChatGPT и 3400 покупок-конверсий — агентство называет это результатом за месяц, хотя точный период в кейсе не указан.²² Цифры — со слов исполнителя, без независимой проверки, так что берём их как иллюстрацию, а не как доказанный факт. Но логика показательна: сработал не размер компании, а то, как переделали контент.

И последнее, о чём спрашивают чаще всего: а какой у этого рынка размер, сколько денег в самом GEO как в услуге? Ответа нет. Оценки российского рынка GEO-агентств гуляют широко и держатся в основном на рекламных обзорах самих агентств — методология нигде не раскрыта, проверить эти числа нельзя. Известно только направление: рынок услуг вокруг продвижения в нейросетях в России уже формируется. А на сколько он «тянет» в деньгах — этого сегодня не может посчитать никто.

Итог главы

Соберём масштаб в одну картину. Больше двух третей поисковых запросов в США заканчиваются без клика — старый трафик тает. Зато трафик из ИИ растёт кратно и конвертируется лучше обычного: воды меньше, нефти больше. Уже 77,6% американцев и 35% россиян хоть раз покупали с помощью ИИ — но доверяют осторожно: платить сразу по совету нейросети готовы единицы. Рынок сейчас в стадии массового первого касания, а не устойчивой привычки. И выигрывает в этой стадии не тот, кто больше, а тот, кто раньше и умнее сделал свой контент видимым для машины.

Если запомнить из главы одну мысль, пусть будет эта: **сдвиг уже случился — деньги и внимание перетекают в ответы нейросетей прямо сейчас, а окно для тех, кто впишется первым, открыто именно потому, что привычка у людей ещё не устоялась.**

Сделайте сейчас одну вещь: откройте свою аналитику и найдите, сколько трафика к вам уже приходит из ИИ-сервисов (в отчётах он обычно идёт отдельным источником — ChatGPT, Perplexity, Gemini и подобные). Даже если цифра пока крошечная — это ваша базовая точка. Через три месяца сравните. Скорее всего, удивитесь.

А в следующей главе перейдём от «сколько» к «кто». Мы уже мельком назвали Ozon, «М. Видео», Mentimeter. Пора разобрать по-настоящему: кто уже стал для нейросетей «ответом по умолчанию» в своей нише, что именно эти игроки сделали — и что из их приёмов реально переносимо на бизнес вашего размера.

Глава 7. Кто уже выигрывает

Два гиганта, две противоположные ставки

В конце 2025 года два крупнейших ритейлера мира сделали ставки на нейросети — и ставки оказались прямо противоположными.

Walmart открыл дверь. Он договорился с OpenAI, чтобы покупатель мог прямо внутри ChatGPT спросить совета, выбрать товар и оформить заказ, не выходя из чата. Target пошёл ещё дальше — там внутри ChatGPT можно собрать корзину из нескольких товаров сразу и выбрать способ получения.^{1 2} Etsy сделал то же. Логика простая: раз человек всё чаще спрашивает у ИИ «что купить», пусть путь от вопроса до оплаты замкнётся тут же, не отпуская его на сторону.

Amazon сделал ровно обратное. Он захлопнул дверь: заблокировал часть чужих ИИ-краулеров, в том числе связанных с ChatGPT, и стал укреплять свой собственный периметр — ассистента Rufus и ИИ-описания товаров внутри своей площадки.^{3 4} Не «пустить ИИ к себе», а «удержать всё внутри своего забора».

Для покупателя разница осязаема. Спросил у ChatGPT, какую взять кофемашину, — и купит скорее у Walmart, даже если раньше в этот магазин не заглядывал. Просто потому, что путь «спросил — купил» замкнут внутри чата именно для Walmart, а не для Amazon.¹ Это уже не спор о том, кто напишет более красивый текст на карточке товара. Это спор о доступе.

У этой истории особая родословная. Это не рекламный кейс, который агентство рассказывает про себя. Это задокументированное отраслевое событие — про него синхронно написали CNBC, CBS и Forbes.^{1 2 3} Когда факт подтверждают несколько независимых деловых изданий, ему можно верить. Держите это в голове всю главу: кейсов вокруг GEO много, а доверия они заслуживают очень разного.

Есть и вторая, менее приятная сторона той же истории. Сделка с OpenAI — это переговоры корпоративного уровня. У малого бизнеса такого рычага нет и не будет. Так что история Walmart вдохновляет, но напрямую вам её не повторить.

Об этом и глава. Мы посмотрим, кто уже стал для нейросетей «ответом по умолчанию» — в мире и в России, в разных нишах. Разберём, что именно эти игроки сделали. И честно разделим: что из их приёмов вы можете забрать себе почти без бюджета, что посылно среднему бизнесу, а что работает только на корпоративном масштабе. А по дороге я научу вас читать чужие кейсы так, чтобы не попасться на красивую обёртку без начинки. Потому что рынок GEO-кейсов сегодня — это витрина, где рядом лежат настоящие данные и чистый маркетинг, и внешне они почти неотличимы.

Одна мысль, которую стоит унести

У всех примеров ниже один общий стержень, и лучше знать его заранее.

В ответ нейросети не попадают за деньги. Туда попадают за сигналы доверия, которые модель собирает из независимых источников: отзывы, упоминания, понятную структуру вашего контента. Мы уже видели это в пятой главе на корреляциях Ahrefs — упоминания бренда в вебе связаны с цитируемостью куда сильнее, чем покупные ссылки. В этой главе то же самое подтвердится на уровне конкретных компаний и ниш.

Отсюда — главная мысль: **выигрывает не тот, у кого толще кошелёк, а тот, кто действует системно и специализированно.** Крупный бюджет ускоряет и масштабирует. Но он не входной билет. Самые надёжные данные, которые вы увидите ниже, показывают эффект именно от действий с нулевым порогом входа.

И сразу — самое убедительное, что у меня есть.

Отзывы: путь от 1% к 75%

Если в этой главе вы запомните одну цифру, пусть будет эта.

Агентство Seer Interactive вместе с Trustpilot проанализировали 800 тысяч ответов нейросетей и опубликовали результат 12 мая 2026 года. У брендов без профиля на площадке отзывов Trustpilot медианная частота цитирования ИИ — всего 1%. У брендов с развитым профилем — 75%. Даже базовый, минимально заполненный профиль поднимает вероятность цитирования на 53%. И эффект воспроизвёлся во всех проверенных нишах, а не в одной удачной.^{5 6}

Теперь про уровень доверия. Исследование заказное: Trustpilot кровно заинтересован в выводе «заведите профиль на Trustpilot». Но три вещи вытягивают его в разряд надёжных. Провело его стороннее агентство, а не сам Trustpilot. Методология и объём выборки раскрыты. И — главное — порог проверки нулевой: любой может завести профиль и убедиться сам, безо всякого бюджета.

Образцовый переносимый приём: денег не требует, требует дисциплины. Собирайте отзывы системно, на профильных площадках, и разрыв между «меня ИИ не видит» и «меня ИИ цитирует» начинает закрываться. Для мира это Trustpilot и G2, для России — Яндекс Карты, 2ГИС, Отзовик, отзывы на маркетплейсах. Подробно про площадки — в пятнадцатой главе, здесь важен сам факт: это работает, это проверяемо, это по карману кому угодно.

Кстати, о том, почему нейросеть вообще так липнет к устоявшимся именам. Исследование Pees AI (анализ почти 38 тысяч ответов ChatGPT, Gemini и Perplexity) показало любопытное: порекомендует модель бренд или нет, почти не зависит от того, как задан вопрос — сухим ключевым запросом или живой разговорной фразой. Формулировка меняет ответ слабее, чем уже накопленная у модели ассоциация «этот бренд = эта тема». ⁷ Проще говоря, модель тянется к тому, кого она уже «знает». А знание это набирается как раз из отзывов, упоминаний и структуры — из тех самых сигналов доверия.

Мировые «ответы по умолчанию»: что они сделали

Теперь пройдемся по нишам. Дальше пойдут кейсы, где компании рассказывают о себе сами, — и уровень доверия к ним другой.

CRM. HubSpot. Самый детальный кейс в мире GEO — и одновременно кейс, к которому нужно относиться с поправкой. HubSpot, мировой лидер CRM, в 2025 году начал с честного признания: «у нас не было надёжного способа измерить свою видимость в ИИ». Гигант оказался ровно в той же растерянности, что и владелец маленького бизнеса, читающий эти строки. С этого — с вопроса «а нас вообще видит ИИ?» — всё и началось.

Дальше HubSpot выстроил программу из трёх опор.⁸ Первая — контент под конкретные вопросы покупателей у себя на сайте: отраслевые страницы с понятной разметкой, статьи-сравнения вроде «5 лучших CRM для строительных компаний», глоссарий базовых терминов. Вторая — внешнее усиление: партнёрство с сотнями сайтов, которые ИИ уже цитирует. Третья — системная работа с Reddit: мониторинг профильных сообществ и участие адвокатов бренда.

Цифры компания приводит внушительные. Статьи-сравнения дали рост цитирований на 642%. Цитирования из Reddit выросли со 178 в мае 2025-го до 146 тысяч в декабре. Общий рост цитирований — 433%, а лиды из ИИ-каналов, по их данным, выросли на 1850% и конвертировались втрое лучше остальных источников.⁸

Впечатляет. И всё же: это компания рассказывает о себе сама, независимого аудита нет, а HubSpot вдобавок продаёт собственный АЕО-инструмент — то есть заинтересован в красивой подаче. От пустого маркетинга его отличает детализация. Цифры разложены по каждому элементу стратегии, есть таблицы и конкретные механики. Из такого разбора уже можно забрать рабочие методы: писать материалы под реальные вопросы, усиливаться на внешних площадках, «жить» в сообществах.

Автоматизация. Zapier — и разрыв, который стоит понять. Zapier называют источником номер один по цитируемости в категории цифровых технологий. Но по прямым упоминаниям бренда он лишь на 44-м месте.^{9 10} Разгадка — в различии, которое пригодится нам до конца книги. «Быть источником» (модель ссылается на вашу страницу) и «быть упомяну-

тым» (модель называет ваше имя) — не одно и то же. Zapier силён в первом и слаб во втором. Механика: тысячи автоматически сгенерированных страниц под каждую пару интеграций из пяти с лишним тысяч приложений. Компания годами строила эту структуру ради обычного SEO — а та же архитектура оказалась удобной для извлечения нейросетью.

Насчёт доверия. Цифра «44-е место» кочует по нескольким вторичным обзорам без единого проверяемого первоисточника, так что беру её с осторожностью.^{9 10} А главный урок Zapier вообще не в цифрах: программную инфраструктуру из тысяч шаблонных страниц малому бизнесу не построить. Это годы разработки и штат. Забираем не саму тактику, а идею — быть источником и быть названным нужно по-разному.

В2В-софт. Площадки отзывов как решающий канал. Тут всё зависит от сегмента. Для горизонтального В2В-софта в цитированиях ИИ обычно доминирует G2, Capterra — далеко второй; для потребительского софта сильнее Trustpilot. Показательное событие: когда G2 в 2026 году присоединил к своей экосистеме Capterra, Software Advice и GetApp, совокупная доля цитирований выросла с 2,09% до 3,68% — плюс 76% относительного роста — и G2 переместился с четвёртого места на второе, уступая только Reddit.¹¹ Отраслевой блог, но с конкретными процентами до и после сделки. Вывод простой и переносимый: выбирайте профильную для вашей ниши площадку отзывов и вкладывайтесь в неё, а не распыляйтесь на все сразу.

Рядом с этой цифрой вам обязательно попадутся другие — и вот с ними осторожнее. По SEO-блогам кочуют красивые множители — «профиль на двух площадках отзывов повышает шанс упоминания в ChatGPT в 3,4 раза», «упоминание на Reddit — в 4 раза».^{12 13} Направление совпадает с надёжным исследованием Trustpilot и Seer, но сами множители не подтверждены ни одним проверяемым первоисточником и повторяются из блога в блог. Верить направлению — да, конкретным числам — нет.

Медиа. Тихое вытеснение обзорных сайтов. А это уже история про тех, кто проигрывает — и она отрезвляет. NerdWallet, крупный сайт сравнения финансовых продуктов, сообщил о падении выручки от кредитных карт на 24% год к году в первом квартале 2025-го, объяснив это «встречным ветром органического поиска». Люди всё чаще получают финансовый совет сразу от ChatGPT, а не идут на сайты-сравнения. После внедрения Google AI Overviews в середине 2024-го Forbes и HuffPost потеряли около 40% трафика каждый, CNN — почти 30%.¹⁴ Эти цифры воспроизводятся синдицированно по десяткам независимых новостных сайтов — значит, за ними единый отчёт, а не выдумка одного блогера. Прежние заслуги здесь не в счёт: если вы были «ответом по умолчанию» в старом поиске, место в новом придётся занимать заново.

Российские кейсы: осторожно, ни одного аудита

Теперь про Россию. Ни один российский GEO-кейс, который я нашёл, не проходил независимого аудита. Все они — либо из одного вендорского мероприятия, либо из агентских публикаций. Это не значит, что они врут. Это значит, что читать их надо с открытыми глазами.

Начну с самого твёрдого, что есть. **Банки.** Компания AIMonitor.pro построила индекс ИИ-видимости российских банков: прогнала больше 40 тысяч запросов и свыше 3,5 тысяч пользовательских сценариев. Результат на 15 мая 2026-го: «Сбер» — 26,7% видимости в ответах нейросетей, «Т-Банк» — 19,1%, ВТБ — 17,9%, «Альфа-Банк» — 14,8%. По кредитным продуктам вперёд вырывается «Т-Банк» — 23% против 22,2% у Сбера. А в «Алисе» разрыв в пользу Сбера особенно велик — 41,3%.¹⁵ Уровень доверия средний: индекс строит и продаёт сама AIMonitor.pro, заинтересованная сторона, — но методологию она раскрыла подробнее большинства.

Самое ценное там — наблюдение сооснователя AIMonitor.pro Кирилла Власова: «ИИ обновляет знания медленнее рынка — исторические бренды сохраняют видимость даже после реструктуризации».¹⁵ Проще: память модели инерционна. Кого она запомнила первым, того держит долго. В восьмой главе эта мысль развернётся в целую стратегию.

Теперь кейс, который я считаю самым полезным из русских, — потому что у него честно раскрыта кухня. **B2B-производитель лазерного оборудования, бренд Pokkels, агентство AdsON.** До начала работы бренд вообще не появлялся в ответах нейросетей. Агентство опубликовало 60 статей на трёх внешних площадках — VC.ru, DTF, Tenchat. Результат замеряли запрос за запросом: бренд стал появляться в «Алисе» по 38 из 60 отраслевых запросов и в Gemini — по 31 из 60. И самое практичное наблюдение: рейтинговые статьи вроде «ТОП-9 установок...» цитировались вдвое чаще информационных — 82% против 38% в «Алисе». ¹⁶

Кейс агентский, аудита нет — а ценен он методикой: какие площадки, сколько материалов, как замеряли, что не сработало. И он бьёт по мифу, что в ответы ИИ пробиваются только гиганты. Небольшой B2B-производитель без многолетнего авторитета пробился — за счёт узкой специализации и правильного формата, а не за счёт размера. Что забрать себе: формат рейтинга («ТОП-N...») цитируется заметно охотнее сплошного текста, и внешние площадки решают. Правда, 60 статей на трёх площадках — это уже работа агентства или выделенного специалиста. Малому бизнесу посильно, но не «между делом».

А вот дальше — целый пласт кейсов, к которым отношение должно быть настороженным. На «GEO-кейс-конференции 2026», которую устроил производитель GEO-инструментов «Пиксель Тулс», прозвучало одиннадцать докладов агентств о результатах для клиентов. ¹⁷ Заявляют многое. Агентство говорит про интегратора i2crm: рост регистраций в 10 раз и конверсия 27% из ИИ-каналов. Про производителя витаминов (Digital Geeks): топ по 75% промптов в Perplexity, «Алисе» и Google AI за четыре месяца. Про Flowwow: удвоение присутствия в ответах за четыре месяца. Про промышленного производителя (ipos. digital): рост трафика из нейросетей на 1500%. ¹⁷

Звучит как парад побед. Но ни в одном из этих кейсов нет полной методологии измерения, независимого подтверждения или контрольной группы «а что было бы без нас». Это анонсы докладов на коммерческом мероприятии, где организатор продаёт GEO-инструменты. Использовать их можно только как индикатор — «вот какие ниши в России уже вкладываются в GEO» — но не как доказанный факт причинности.

Особенно наглядна цифра «1500% роста трафика». Красиво. Только с какой базы считали — со скольких визитов до скольких? В анонсе этого нет. На малой базе 1500% могут означать «было два перехода, стало тридцать». Процент без базы — это не факт, а фокус. Всегда спрашивайте, от чего считали.

И отдельная деталь, полезная любому. У того же i2crm агентство честно упомянуло парадокс: ИИ уже называет компанию лучшей, а половина «горячих» клиентов всё равно уходит к конкурентам на сайте. ¹⁷ Даже если кейс приукрашен, наблюдение отрезвляющее. Попадание в ответ нейросети — не финал воронки, а её начало. Вас назвали — но купить у вас человек ещё должен решить сам.

Как читать чужие кейсы и не попасться

Мы разобрали больше десятка заявлений. Пора вооружить вас так, чтобы дальше вы отличали настоящий кейс от красивой обёртки без моей помощи.

Лучший урок тут преподавал Сбер — правда, невольно. Компания выпустила рекламный спецпроект с 40 бизнес-кейсами о том, как её нейросеть GigaChat помогает бизнесу: ускорение процессов в 16 раз, рост выручки на 30%, всё в таком духе. А независимый разбор этих кейсов показал неудобное. Модель GigaChat 2 Max, которую использует подавляющее большинство этих кейсов, на момент публикации занимала последнее, 29-е место из 29 в независимом бенчмарке на управленческих задачах. И ни в одном из 40 кейсов не было сравнения с альтернативными моделями. ¹⁸

Автор разбора сформулировал принцип, который стоит переписать на стикер: вендорский кейс отвечает на вопрос «работает ли решение» — но не на вопрос «оптимальное ли это решение». ¹⁸ Сорок сияющих карточек говорят «у нас получилось». Они молчат о том,

получилось бы ещё лучше с чем-то другим — и что было бы вообще без вмешательства. Сам этот разбор, в отличие от кейсов, независимый и прозрачный: открытый бенчмарк, 29 моделей, больше четырёх тысяч оценок. Потому ему и веришь.

Ещё опаснее — прямой обман под вывеской GEO. В Рунете попадаете маркетинг, который продаёт «эксклюзивное лидерство в нише» через закрытый аукцион: заплати за «золотой пакет» — и нейросеть будет называть только тебя, единственного победителя, в твоём регионе.¹⁹ Звучит заманчиво — и рассыпается, стоит вспомнить, как ИИ устроен на самом деле. Модели опираются на консенсус десятков независимых источников, а не на купленное место, и обычно называют не одного, а нескольких кандидатов. В том же материале рядом уживаются фраза «нейросети не продают рекламные места» и продажа гарантированного доминирования — внутреннее противоречие, выдающее пустышку.¹⁹ Правило простое: если вам обещают гарантированное эксклюзивное место в ответе ИИ за деньги — это красный флаг. Так эта механика не работает.

И главное про всю эту молодую индустрию. Ни одного примера независимого аудита GEO-кейса — вроде того, как Nielsen проверяет рекламную эффективность в медиа, — найти не удалось. Ни в России, ни в мире. Это не проблема отдельных нечистых на руку агентств. Это пробел всей сферы: она слишком молода, аудиторов ещё нет. Поэтому любой кейс, включая самые красивые, вы читаете под свою ответственность. Держите три вопроса наготове: кто это измерил (сам исполнитель или третья сторона), раскрыта ли методология, и можно ли проверить это самому. Если на все три ответа мутный — перед вами реклама, а не доказательство.

И ещё одно: даже честные цифры протухают быстро. Reddit много где называют источником номер один для крупных ИИ. Но конкретная доля скачет так, что за неё нельзя держаться. Одно агентство, 5W, в своём пресс-релизе даёт Reddit около 40% цитирований.²⁰ А на собственной исследовательской странице то же агентство по тем же данным Similarweb за январь—февраль 2026-го приводит совсем другое — Reddit 11,97%, Wikipedia 13,15%.²¹ А по отдельным обзорам роль Wikipedia в ИИ-цитировании гуляет ещё в куда более широком диапазоне оценок.²² Разгадка скучная: это разные срезы. И та же страница честно предупреждает: по данным SEMrush доля Reddit в цитированиях ChatGPT за две недели сентября 2025-го обвалилась примерно с 60% до 10%, а потом частично отросла.²¹ Тем временем LinkedIn за несколько месяцев ворвался в топ цитирований ChatGPT с позиции за пределами двадцатки — для B2B-запросов он теперь один из главных источников.²³ Мораль на будущее: любой процент в этой сфере — моментальный снимок с датой годности, а не вечная константа. Про волатильность как норму — отдельный разговор в восьмой главе.

Что переносимо на ваш бизнес, а что — нет

Мы разобрали победителей. Теперь сведём всё к самому важному вопросу: что из этого можете сделать лично вы?

Разложу по трём полкам — от «бесплатно и сразу» до «только на корпоративном масштабе».

Полка первая. Переносимо почти без бюджета — начните отсюда.

Заполните и синхронизируйте профиль на картах и в справочниках: Google Business Profile, Яндекс. Бизнес, 2ГИС, Zoon. Одинаковые название, адрес, телефон везде. Эффект подтвердили независимые исследования локального поиска: Local Falcon разобрал 4423 бизнеса в 20 странах, Whitespark вручную собрал 540 запросов по шести отраслям.^{24 25} Это компании, которые сами продают инструменты локального SEO, но методологию выборки они раскрыли.

Собирайте детализированные отзывы на одной-двух профильных площадках. Это тот самый путь из исследования Trustpilot и Seer, с которого мы начали главу: от единичных упоминаний — до цитирования в большинстве ответов.⁵ Денег не требует, требует дисциплины.

Сделайте самодостаточные блоки вопрос-ответ и понятную структуру страниц услуг. Отражайте реальные вопросы клиентов. Бесплатно.

Займите узкую экспертную нишу вместо попытки быть «обо всём». Кейс Pokkels показывает: небольшой игрок пробивается именно специализацией, а не размером.¹⁶

Честно участвуйте в профильных сообществах: Reddit — для мира, Habr, vc.ru, профильные Telegram-чаты — для России. Не реклама, а решение реальных проблем людей. Требуется времени и экспертизы, но не бюджета.

И проверяйте себя. Раз в неделю прогоняйте десяток целевых запросов через ChatGPT, Perplexity, «Алису» и смотрите, называют ли вас.^{26 27} Делать это можно двумя способами: руками, ведя простую таблицу, или через специальный сервис мониторинга, который прогоняет запросы и собирает картину за вас, — их разбор ждёт вас в главе про инструменты. Руками — бесплатно, но мучительно; сервис экономит время и меньше зависит от того, что вы забыли проверить.

Полка вторая. Требуется бюджета или масштаба, но посилено среднему бизнесу.

Публикации на внешних площадках с высоким авторитетом — VC.ru, DTF, Habr, отраслевые СМИ. Кейс Pokkels — это 60 статей на трёх площадках за несколько месяцев. Работа агентства или выделенного контент-специалиста, но не корпоративный бюджет.¹⁶

Регулярный мониторинг ИИ-видимости в специальном сервисе. Отслеживать своё присутствие в ответах нейросетей стоит не набегами, а системно: иначе за естественным разбросом ответов не разглядеть тенденцию. Сервис снимает замеры по расписанию и показывает рейтинг по вашей нише и список источников, которые сейчас в ходу у каждой нейросети. Инструментов на рынке несколько — есть российские (НейроИндекс, AIMonitor.pro, «Пиксель Тулс»), есть зарубежные (XFunnel); сравним их отдельно. Гнаться за идеальной точностью тут не надо: ответы моделей всё равно плавают от замера к замеру, поэтому важнее регулярность и одна и та же методика, чем второй знак после запятой.

Полка третья. Практически недоступно малому бизнесу — только корпоративный масштаб.

Прямые партнёрства с провайдерами ИИ — та самая сделка Walmart и Target с OpenAI об агентной покупке внутри ChatGPT.¹ Переговоры корпоративного уровня.

Программа «партнёрство с сотнями сайтов» по модели HubSpot — почти тысяча новых страниц у партнёров за месяцы. Нужен отдел партнёрств и бюджет на чужой контент.⁸

Программная инфраструктура масштаба Zapier — тысячи шаблонных страниц под каждую комбинацию. Годы разработки и штат.⁹

Закономерность видна невооружённым глазом: самые доказанные и самые дешёвые вещи стоят на первой полке. А то, что требует корпоративных денег, — как раз то, что доказано хуже всего и подходит немногим. Это и есть хорошая новость для маленького бизнеса: решает не размер кошелька, а системность и специализация.

Итог главы

В сухом остатке. В мире и в России уже есть свои «ответы по умолчанию»: Walmart замкнул покупку внутри ChatGPT, HubSpot стал заметнейшим в CRM, в России «Сбер» и «Т-Банк» лидируют по видимости в ответах ИИ. Но самый надёжный факт главы касается не гигантов, а приёма, который доступен каждому: развитый профиль на площадке отзывов поднимает цитируемость в разы. Крупный бюджет ускоряет. Входным билетом он не служит.

И вторая мысль, не менее важная: индустрия молодая, независимого аудита GEO-кейсов пока нет ни одного, а красивые проценты без раскрытой базы и методологии — это реклама, а не доказательство. Вдохновляйтесь победителями, но проверяйте каждый кейс тремя способами: кто измерил, раскрыта ли методика, можно ли повторить самому.

Если унести из главы одно: **выигрывает не тот, кто больше, а тот, кто раньше и системнее сделал себя видимым для машины — на дешёвых, проверяемых действиях, а не на дорогих обещаниях.**

Сделайте сейчас одну вещь: выберите одну профильную для вас площадку отзывов и заведите там нормальный профиль. Не «когда-нибудь», а сегодня. Это тот самый шаг с нулевым порогом входа и самым доказанным эффектом.

А в следующей главе разберёмся, почему всё это надо делать именно сейчас. Мы уже заметили две вещи мимоходом: память модели инерционна — кого она запомнила первым, того держит долго; и при этом доли источников скачут неделями. Из этого противоречия и растёт главный сюжет восьмой главы — окно возможностей уже открыто, но оно не будет открыто вечно.

Глава 8. Почему сейчас: окно и его закрытие

670 раз подряд машина выбрала знакомое имя

Представьте лабораторный стол, на котором лежат десять баночек крема. Девять — от брендов, которых не существует: их придумали для эксперимента. Одна — от настоящего, известного бренда. И вот хитрость: все десять абсолютно одинаковы. Одна цена, один рейтинг, одинаковое число отзывов. Отличается только имя на этикетке.

Теперь спрашиваем у нейросети: какой посоветуешь?

Она выбирает известный бренд. Снова спрашиваем — снова его. И так 670 раз подряд. Ни один из девяти безымянных кремов не был выбран ни разу. Ни разу.

Это не байка, это результат эксперимента Си Чу из Trine University и Юйпэна Хоу из Texas A&M, опубликованного 16 июня 2026 года.¹ Они прогнали задачу через три коммерческие модели — GPT-4o-mini, Claude Sonnet и Gemini 3 Flash, — на двух языках, в четырёх подкатегориях ухода за кожей. Когда товары неотличимы по характеристикам, узнаваемое имя побеждает в 100% случаев. Индекс преимущества, который ввели авторы, упёрся в теоретический максимум — 10,0 из 10,0.

Но без важной поправки эта цифра соврёт. Основной эксперимент проведён на уходовой косметике — это так называемый *experience good*, товар, качество которого нельзя оценить, пока не попробуешь. На кабелях и батарейках авторы прогнали отдельную, вспомогательную проверку, и там базовый эффект повторился, — но проверка эта была меньше и не такая подробная.¹ Так что не переносите слепо «100% из 670» на любую нишу. Это точный результат на одной категории, а не закон природы.

И всё же он про то, о чём мы говорили в конце прошлой главы. Помните наблюдение Кирилла Власова: «ИИ обновляет знания медленнее рынка — исторические бренды сохраняют видимость даже после реструктуризации»? Тогда это было наблюдение практика. Теперь у него есть лабораторное подтверждение. Кого модель узнала первым, того она держит.

О чём эта глава и почему она — про время

Вся вторая часть книги отвечала на вопрос «какой рынок открывается». Мы увидели масштаб сдвига и разобрали, кто уже стал для нейросетей ответом по умолчанию. Осталось закрыть последнее: почему за это надо браться сейчас, а не «когда всё устаканится».

Ответ живёт в противоречии, которое мы дважды задели мимоходом. С одной стороны — память модели инерционна: ранний игрок закрепляется. С другой — доли источников скачут неделями, и никакая позиция не закреплена навсегда. Одно с другим уживается плохо. И раз всё так шатко, не разумнее ли переждать?

Разберёмся по порядку. Сначала — насколько реально преимущество первопроходца и почему у него есть срок годности. Потом — почему волатильность — это не аврал, а норма, с которой просто надо уметь жить. И в финале — как войти в это окно сейчас, но построиться так, чтобы не рухнуть при первой же смене алгоритма. Спойлер: паниковать не надо, но и ждать нельзя.

Преимущество первопроходца — реальное, но узкое

У эксперимента с кремами есть второе дно, и оно важнее первого.

Авторы назвали найденный паттерн «условной монополией». Слово «условной» — ключевое. Разложив выбор модели по факторам, они увидели такую картину: характеристики товара объясняют 82,4% решения, позиция в списке — 6,5%, а сама узнаваемость бренда — всего 1,2%.¹ То есть имя бренда почти ничего не решает, пока у модели есть за что зацепиться. Оно работает как тай-брейкер — язычок весов, который перевешивает только тогда, когда чаши идеально равны.

А чаши равны редко. В реальном мире у товаров разные цены, рейтинги, отзывы. И как только появляется хоть малейшее объективное отличие, монополия имени рассыпается. В прямом сравнении, где безымянный товар чуть лучше по характеристикам, модель всё равно цепляется за известный бренд лишь в 1,7—4,6% случаев.¹ Почти всегда выигрывает объективно лучший продукт — даже без имени.

Насколько «чуть лучше» достаточно? Порог отрезвляюще низок. Разница в рейтинге всего в 0,075 звезды. Или скидка 7,3%. Или в 1,6 раза больше отзывов. Любого из этих сдвигов хватает, чтобы вероятность выбора безымянного бренда прыгнула с 3,6—6,0% до 64—80%.¹ Не плавно поднялась — прыгнула, одной ступенькой.

Отсюда первый практический вывод главы. Инерция памяти — это не броня. Это фора, пока у машины нет альтернативы. Пока о вас модели нечего сказать, кроме имени, знакомое имя выигрывает. Но стоит новичку дать машине объективный повод — цифру, отзыв, доказательство, — и фора обнуляется. Хорошая новость для того, кто только заходит: дверь не заперта. Плохая новость для того, кто уже внутри: имя вас не спасёт, если рядом появится кто-то с фактами.

Здесь надо развести два похожих механизма, которые легко спутать. Есть **инерция памяти** — то, что модель усвоила из обучающих данных, между сессиями, задолго до вашего вопроса. Именно про неё говорил Власов и именно её мерили Чу и Хоу. А есть отдельное явление — **эффект первенства**: модели склонны выделять то, что видят первым уже внутри одного запроса, в порядке подачи фактов или списка источников. Это подтверждено на нескольких моделях того времени, включая GPT-3.5, GPT-4 и Claude-instant-1.2.² Механизм другой — речь про порядок здесь и сейчас, а не про давнюю память, — и смешивать их в один факт нельзя. Но для вас практический итог совпадает: первым быть выгодно. И в том, что модель успела запомнить, и в том, что она видит в начале ответа.

Язык доказательств — рычаг, который ломает монополию

Самая полезная для нас находка того же исследования — про то, чем именно новичок пробивает фору старожила.

Оказалось, дело не только в реальном качестве. Работает и язык доверия. Формулировки, которые выглядят как доказательство — ссылки на клинические данные, экспертные подтверждения, — ломали монополию известного бренда при «прибавке ценности» всего в 0,17 балла рейтинга.¹ Проще говоря, текст, который звучит как доказательство, для модели почти эквивалентен реальному улучшению продукта.

А вот привычные маркетинговые приёмы почти не сработали. Искусственный дефицит, «успей купить», игра на страхе упустить — всё это пробивало монополию лишь в 9,6—12,9% случаев против 50—73% у авторитетных формулировок.¹ Машину не берёт то, чем берут человека. Её берут доказательства.

Это же ровно то, что доказал ещё первоисточник GEO из Принстона: статистика, цитаты и ссылки поднимают попадание в ответы ИИ. Только теперь тот же вывод пришёл с другой стороны, на более строгом эксперименте. Два независимых исследования указывают в одну точку — и это как раз тот случай, когда цифре можно верить. К самим приёмам — цифрам, цитатам, ссылкам — мы вплотную подойдём в четвёртой части. Пока держите в голове: вход в ответ нейросети покупается не громкостью, а доказательностью.

Окно закрывается изнутри: дилемма заключённого

А теперь — почему у всего этого есть срок годности. Это, пожалуй, самая важная мысль главы, и она снова из того же исследования.

Авторы смоделировали конкуренцию: что будет, если приём с «языком доказательств» начнут применять не один новичок, а несколько сразу?

Один-единственный конкурент, взявший приём на вооружение, обрушивает монополию старожила: его доля выживания в рекомендациях падает со 100% почти до 20%.¹ Один сме-

лый новичок отъедает у лидера почти всё. Но когда приём подхватывают все девять конкурентов разом, сигнал перестаёт выделять кого-либо — и модель снова скатывается к знакомому имени. Доля старожила восстанавливается до 93,8%.¹

Дальше — цифра, которую стоит запомнить. Выгода от приёма для того, кто применил его первым, — 0,802 в условных единицах выигрыша. Для того, кто применил его, когда так делают уже все, — 0,007.¹ Почти ноль. Ранний забирает всё, опоздавший — крохи.

Авторы прямо называют это дилеммой заключённого. Каждому отдельному бренду выгодно применить приём — это выигрышный ход при любом раскладе. Но когда так поступают все, общий выигрыш почти испаряется. Их собственная формулировка стоит того, чтобы привести её целиком: «Для новых брендов GEO даёт мощное, но недолговечное преимущество первопроходца — структура дилеммы заключённого означает, что ранние последователи выигрывают больше всего, но окно закрывается по мере того, как конкуренты подтягиваются».

1

Так и выглядит «закрывающееся окно». Не кто-то извне его захлопывает. Оно схлопывается изнутри, само, по мере того как приём становится всеобщим. И тут есть ловушка, из-за которой переждать не получится. Тот бренд в эксперименте, который вообще не подключился к приёму, получил ноль рекомендаций. Ноль, во всех 4745 попытках.¹ Расклад выходит невестельный: участвовать в обесценившейся гонке невыгодно — но не участвовать ещё хуже. Опоздавший получает крохи, а отказавшийся — ничего.

Это и есть точный ответ на вопрос «может, переждать?». Нет. Пока приём в новинку — он даёт максимум. Когда станет нормой — не даст почти ничего, но без него вас просто не будет в ответе. Единственный выигрышный момент — сейчас, пока большинство ещё раскачивается.

Это модельный эксперимент на узкой категории, а не замер живого рынка. Механизм он показывает убедительно, но не обещает, что в вашей нише окно захлопнется ровно завтра. Кстати, о сроках: в отраслевых блогах любят прогноз «к 2027—2028 году GEO станет базовой гигиеной, как SEO». Звучит правдоподобно, но это именно прогноз практиков, а не измеренный факт, и у разных авторов сроки расходятся на годы. Опирайтесь не на дату, а на логику: раньше — выгоднее.

Волатильность — это погода, в которой вы работаете

Теперь ко второй половине нашего противоречия. Если ранний игрок закрепляется, почему же доли источников так лихорадит?

Как вы уже знаете, цитирования Reddit в ChatGPT обвалились в разы буквально за несколько недель. Это был первый звоночек. А вот случай посвежее и покрупнее.

27 января 2026 года Google перевёл свои AI Overviews — те самые ИИ-ответы над обычной выдачей — на модель Gemini 3 по умолчанию.³ И началось. Первую неделю систему лихорадил баг: доля ответов вообще без единого источника подскочила с обычных 0,11% до 10,63%.⁴ Каждый десятый ответ ничего не цитировал. Каково владельцу бизнеса, который годами выстраивал присутствие в этих ответах, а однажды утром обнаружил, что источники просто исчезли. Google баг починил, но доля устоялась не на прежнем уровне, а на 1,27% — в 11,5 раза выше, чем до обновления.⁴ Похоже, это новая норма, а не временный сбой.

Исследователи SE Ranking прогнали 100 тысяч запросов в двадцати нишах до и после перехода и отделили эффект бага от эффекта самой модели.⁴ Вот что вышло, и это стоит разоб-
брать спокойно.

42,4% доменов, которые цитировались раньше, перестали цитироваться вовсе.⁴ Больше четырёх сайтов из десяти вылетели из ответов. Цифру независимо подтвердил и другой источник.⁵ Звучит как катастрофа. Но интереснее другое — кого именно снесло.

Среди пятисот самых цитируемых доменов исчез ровно один. Ни одного нового в этот топ-500 не вошло.⁴ То есть буря прошла почти целиком по «длинному хвосту» — по сайтам с горсткой цитирований, — а признанных лидеров едва задела. Разнообразие источников

при этом выросло: среднее число ссылок на ответ поднялось с 11,55 до 15,22. Но и концентрация выросла одновременно — топовые игроки стали забирать ещё большую долю.⁴ На первый взгляд парадокс: как разнообразие и концентрация растут вместе? А так, что расширение шло за счёт хвоста, а не за счёт вытеснения лидеров. Крупные закрепились ещё крепче, мелкие перетасовались.

И всё же не думайте, что лидерство — гарантия. Один яркий пример из того же исследования: Shopify, который был третьим по цитируемости доменом в нише «Бизнес», после перехода выпал из топ-10.⁴ А на его место поднялись другие — та же Торговая палата США переехала с девятого места на четвёртое. Shopify не сделал ничего плохого. Просто платформа сменила модель, и расклад поехал. Волатильность бьёт не только по новичкам.

Ответ меняется каждые два дня — и это нормально

Смена модели раз в год — событие. А есть волатильность, которая идёт постоянно, каждый день, и о ней важно знать, чтобы не паниковать понапрасну.

Ahrefs взял 43 тысячи запросов и следил, как меняется ИИ-ответ Google от замера к замеру на протяжении месяца.⁶ Результаты такие.

Вероятность, что текст ответа изменится от одного наблюдения к следующему, — 70%. Среднее «время жизни» одной версии ответа — 2,15 дня. При обновлении в среднем 45,5% процитированных источников заменяются на новые.⁶ Сотрудница Ahrefs описывала, как получила две разные версии одного ответа, просто обновив страницу дважды подряд.⁶ Вот так это выглядит на кухне — не в отчёте, а вживую.

И тут же — деталь, которая всё расставляет по местам. При всей этой чехарде смысл ответа держится стойко. Смысловое сходство соседних версий — 0,95 из 1,0.⁶ Слова меняются, источники меняются, а суть — почти нет. Модель как будто каждый раз пересказывает одну и ту же мысль другими словами и с другими ссылками.

Что это значит для вас практически? Две вещи. Первая: не сходите с ума от одного замера. Если сегодня ИИ вас процитировал, а завтра нет — это может быть просто шум, обычное дыхание системы, а не провал вашей стратегии. Отслеживать нужно не единичный ответ, а вашу видимость по теме в целом, усреднённо, за период. Только так за шумом проступает настоящая тенденция. Вторая: раз состав источников всё время тасуется, попадание в ответ нельзя «взять и забетонировать». Его приходится удерживать регулярной работой. Разовая акция выветрится за те самые пару дней.

Кстати, о мифе, будто старому бренду со старым контентом цитирование гарантировано. Не гарантировано. Ahrefs разобрал 17 миллионов цитирований и увидел: в среднем ИИ тяготеет к более свежим материалам, но разброс между платформами огромный.⁷ ChatGPT цитирует материалы на 393—458 дней новее обычной выдачи, а вот AI Overviews Google, наоборот, охотно берут статьи возрастом почти четыре года.⁷ Единого правила свежести нет. И это ещё один довод против того, чтобы строить стратегию на одной платформе и одной догадке о том, что она любит.

Хорошая новость: рынок ещё не поделён

Пока всё звучало довольно тревожно: окно закрывается, буря бушует. Пора уравновесить. У окна есть и вторая сторона — оно шире, чем кажется, потому что рынок ИИ-рекомендаций пока не монополизирован.

Исследование под названием «Кто владеет ИИ-рекомендацией?» проанализировало 3750 ответов трёх моделей — GPT-5.2, Gemini 3 Flash и Perplexity — по 250 запросам в пяти индустриях.⁸ Выводы обнадеживают того, кто только заходит.

Концентрация оказалась умеренной. Коэффициент Джини — мера неравенства, где 0 — это полное равенство, а 1 — это «победитель забирает всё» — составил в среднем 0,28. Это заметно ниже порога монополизации.⁸ Данные не подтверждают страшилку «победитель получает почти всё» — по крайней мере в изученных нишах.

Но самое ценное для стратегии — другое число. Три модели сошлись во мнении, чей это бренд-лидер в категории, лишь в 41,6% запросов.⁸ То есть чаще, чем в половине случаев, у ChatGPT, Gemini и Perplexity разные представления о том, кто главный в теме. А значит, окно не закрыто одновременно на всех платформах. Можно проиграть у одной модели и выиграть у другой. Можно заходить на каждую по отдельности.

Про сам источник. Автор исследования аффилирован с коммерческой платформой ИИ-аналитики, которая продаёт продукт на основе тех самых метрик из статьи.⁸ Он сам раскрывает это и утверждает, что на выводы это не повлияло, а данные выложены в открытый доступ для проверки. Верить цифрам это не мешает — методология раскрыта, датасет открыт, — но знать про интерес автора вы должны. И помните: выборка узкая — пять индустрий, снимок на два дня февраля 2026-го. На всю экономику не переносится.

Как войти в окно и не построить на песке

Картина сложилась. Окно открыто и ждёт раннего игрока. Оно закрывается изнутри, по мере того как приёмы становятся общими. И всё внутри него постоянно штормит. Как действовать в такой обстановке, чтобы не суетиться зря и не строить на песке? Три правила.

Правило первое: заходите сейчас, но опирайтесь на фундамент, а не на трюки.

Разница простая. Фундамент — это то, что работает, потому что так модель принимает решение: доказательность, отзывы, понятная структура, узнаваемость бренда из независимых источников. Трюк — это хрупкая тактика, которая держится на текущей особенности одной платформы и завтра может испариться.

Лучший пример хрупкого приёма — файл llms.txt. Его порой подают как обязательный технический шаг GEO. А вот что 2 июня 2026 года сказал про него Джон Мюллер из Google: «Я не думаю, что кто-то знает — пока это чистая спекуляция. Файл существует уже несколько лет, но ни одна из ИИ-систем его не использует». ⁹ Это личное мнение представителя Google, не официальная политика компании — но показательно, что даже широко разрекламированный приём может оказаться пустышкой без единого случая реального применения. Сравните с отзывами: как вы уже знаете из прошлой главы, там разрыв между «профиля нет» и «профиль живой» измеряется десятками раз — и польза не просто измерена, а воспроизведена. Вот и весь критерий. Спрашивайте про любой модный приём: это доказано и держится на механике решения — или это ставка на догадку? На первое опирайтесь, второе пробуйте по остаточному принципу и не стройте на нём.

Правило второе: диверсифицируйте. Раз три модели расходятся во мнениях чаще, чем сходятся, а платформа может в одну ночь сменить модель и перетасовать выдачу — нельзя ставить всё на один канал. Присутствие на нескольких площадках и в нескольких форматах — это не распыление, а страховка. Помните Shopify, который вылетел из топ-10 не по своей вине? От такого спасает только то, что вы есть ещё где-то. Другой защиты от алгоритма, который меняется без предупреждения, попросту нет. Отдельно и подробно про это — в главе о волатильности источников, которой закрывается следующая часть, но принцип закладывайте с первого дня.

Правило третье: работайте регулярно, а не наскоком. Если состав источников в ответе тасуется каждые два дня, а половина ссылок при обновлении заменяется, то разовая публикация просто утонет в этом потоке. Удержание в ответах — это ритм, а не подвиг. Регулярно обновлять материалы, регулярно проверять свою видимость усреднённо по теме, регулярно докладывать доказательства. GEO — это марафон, и про сроки у нас будет отдельная глава в конце книги. Здесь важен принцип: система, повторяемая из месяца в месяц, бьёт любой одноразовый рывок.

И ещё один сигнал, что медлить не стоит. 9 февраля 2026 года OpenAI начала показывать рекламу в ChatGPT части пользователей, а к июню расширила пилот на несколько стран.^{10 11} Ещё недавно ChatGPT был чистым источником органических упоминаний. Теперь площадка

начинает встраивать монетизацию. Пока реклама отделена от органического ответа — но это тот самый признак взросления рынка, когда бесплатное внимание постепенно становится платным. Ещё один довод в пользу того, чтобы занять место сейчас, пока внимание ещё бесплатное.

Пара слов про то, чему не верить

Напоследок отделим факты от рыночного шума — пригодится, когда наткнётесь на громкие заголовки.

Вам почти наверняка встретится тезис вроде «85% брендов невидимы для ИИ» или «исключены из базы знаний ChatGPT». Цифра ходит по вендорским блогам, но за ней нет раскрытой методологии подсчёта — это маркетинговый тезис компаний, продающих услуги видимости, а не измеренный факт.¹² Пугает эффектно, доказывает мало. Относитесь как к рыночной молве.

То же и с оценками размера рынка GEO-услуг, и с «периодом полураспада цитируемости в один год»: направление обычно верное (рынок растёт, свежесть важна), а конкретные числа у вендоров без открытой методики расходятся чуть ли не вдвое. Держитесь за направление, не за цифру.

Итог главы

Сложим всё вместе. Преимущество первопроходца реальное — модель держится за то, что узнала первым, и это подтверждено экспериментом, где машина 670 раз подряд выбрала знакомое имя среди одинаковых товаров.¹ Но это фора, а не броня: объективное отличие в 0,075 звезды или язык доказательств пробивают её мгновенно.¹ И у форы есть срок годности — окно закрывается изнутри само, по мере того как приём становится всеобщим: выгода первого падает с 0,802 почти до нуля, но тот, кто не участвует вовсе, не получает ничего.¹

При этом всё внутри окна штормит постоянно — ответ меняется каждые 2,15 дня, платформа может за ночь сменить модель и снести 42% доменов.^{4 6} Волатильность — это реальность, в которой вы и работаете. Спасают от неё не паника, а фундамент, диверсификация и регулярность.

Если унести из главы одно: **окно открыто сейчас, и именно сейчас вход в него стоит дешевле всего — но заходить надо на доказательствах и системности, а не на трюках, потому что трюки обесценятся, а фундамент останется.**

Сделайте сейчас одну вещь: выберите один свой ключевой запрос — тот, по которому клиент должен находить именно вас, — и проверьте, кого по нему называет ИИ сегодня. Не «когда-нибудь», а сегодня. Это ваша точка отсчёта, от которой вы измерите всё остальное.

На этом вторая часть закрывается. Мы разобрали, какой рынок открывается и почему момент именно сейчас. Дальше — самое практическое.

И начнём мы не с текста, а с места. Потому что прежде чем спрашивать «что писать», надо ответить на два других вопроса: **кто вы такой в глазах машины** и **где она вообще вас ищет**. Можно написать безупречную статью — и остаться невидимым, потому что нейросеть не поняла, о какой компании речь, или потому что она в принципе не заглядывает туда, где эта статья лежит.

Так что третья часть — про присутствие. Сначала про сущность: как сделать так, чтобы ИИ уверенно опознавал вас и не путал с тезкой. Потом — карта площадок: откуда нейросети берут материал, что из этого ваше, что чужое и как попадать в чужое. Окно открыто. Пора разбираться, где именно в него заходят.

Часть III. Где присутствовать

Глава 9. Сущности: как ИИ понимает, что вы — это вы

Тридцать упоминаний, которых почти не было

Один стартап открыл дашборд слежения за нейросетями и обрадовался. Система насчитала тридцать упоминаний бренда в ответах ИИ и «аудиторию» больше шестисот тысяч человек. Для маленькой компании это выглядело как прорыв: значит, машины про нас знают, значит, работаем не зря.

Потом кто-то из команды решил прочитать сами упоминания. И обрадоваться стало нечему. Двадцать восемь из тридцати оказались вообще не про компанию. У бренда был тёзка — известный вымышленный персонаж с тем же именем, — и почти весь «успех» состоял из фанатских теорий, коллекционных фигурок и квизов про этого персонажа. Реальный продукт мелькнул от силы в одном-двух ответах из тридцати.¹

С этого и начнём. Метрика «нас упоминают» не стоит ничего, если нейросеть не смогла отличить ваш бренд от чужой, более известной сущности с похожим именем. Можно накопить сотни упоминаний, ссылок и отзывов — и остаться невидимым, потому что машина сложила всё это добро не в ту корзину.

Поэтому третья часть книги, разговор о том, где присутствовать, открывается не картой площадок, а вопросом попроще: **понимает ли машина, о ком идёт речь**. Пока она этого не понимает, любая работа на любой площадке уходит в чужую копилку. Сущность — это фундамент под всем, что будет дальше: под сайтом, форумами, Wikipedia, отзывами, публикациями в прессе. Разберёмся с ней — и всё остальное начнёт складываться в одну фигуру, а не рассыпаться на осколки.

Вещи, а не строки

Чтобы разобраться, придётся вернуться на пару шагов назад — к тому, как машины вообще научились понимать, что предмет разговора — это предмет, а не набор букв.

Шестнадцатого мая 2012 года Google запустил то, что назвал «графом знаний» (Knowledge Graph), и объяснил идею одной фразой, которая с тех пор гуляет по всем учебникам: **things, not strings** — «вещи, а не строки». ² Смысл вот в чём. Раньше поиск искал строки текста: вы вводите «яблоко», система ищет страницы, где встречается это слово. А люди-то ищут не буквы. Они ищут объекты, места, компании и понятия — вещи, которые существуют в мире и о которых что-то известно. Google решил научиться искать именно их. На старте граф знаний содержал больше 500 миллионов объектов и свыше 3,5 миллиарда фактов и связей между ними.² Сегодня это та самая база, из которой собираются карточки-справки сбоку от поисковой выдачи — «база данных из миллиардов фактов о людях, местах и вещах», как описывает её сам Google.³

Ключевое слово здесь — **сущность** (entity). В модели Google сущность это «чётко определённая, уникальная и различимая вещь»: то, что можно однозначно опознать и отличить от всего остального.⁴ Человек, город, компания, а иногда и абстрактное понятие вроде цвета или эмоции. Разницу между сущностью и ключевым словом стоит уложить в голове крепко: на ней держится вся глава.

Ключевое слово — это строка символов. «Ягуар» как набор букв. А сущность — это то, что стоит **за** словом. И тут вылезает подвох: у одной строки может быть несколько сущностей. «Ягуар» — это и хищная кошка, и марка машины. Классический пример из учебников — английское слово Mercury: планета, римский бог, автомобильный бренд, химический элемент и вдобавок имя музыканта. Пять совершенно разных вещей за одной строкой.⁵ Человек

разбирается по контексту мгновенно. А машине нужно проделать отдельную работу — понять, о какой именно из пяти сущностей речь, прежде чем что-то отвечать.

Эта работа называется *entity linking*, или *disambiguation* — связывание и разрешение неоднозначности сущностей. Звучит мудрёно, а суть простая: система берёт упоминание в тексте и сопоставляет его с конкретной записью в базе знаний. Не просто «это какая-то организация», а «это вот эта конкретная организация, номер такой-то». ⁵ У сущности, в отличие от ключевого слова, есть уникальный идентификатор, и она не зависит ни от языка, ни от написания. «Ягуар», *jaguar* и *###* могут указывать на одну и ту же кошку.

Почему нейросети «думают сущностями»

Хорошо, скажете вы, это всё про поиск Google. А при чём тут ChatGPT и другие языковые модели, которые вроде бы просто «предсказывают следующее слово»?

При том, что модели впитали ту же логику из данных, на которых учились. И есть один источник, который тут стоит особняком, — Wikipedia. В классической таблице состава обучающих данных GPT-3 Wikipedia занимала всего около 3% корпуса; основной объём давал Common Crawl — гигантская свалка веб-страниц. ⁶ Всего 3%. Но весили эти 3% куда больше своей доли, и дело тут не в объёме. Wikipedia — источник плотный, выверенный, структурированный: каждая статья описывает ровно одну сущность, у неё есть чёткое имя, тип и связи с другими статьями. Модель училась на таком материале видеть за словами объекты, а не строки. Поэтому и она, и поиск, из которого нейросети сегодня тянут свежие факты, устроены вокруг сущностей.

Вывод для вас практический. Чтобы попасть в ответ, мало быть в интернете под своим названием. Нужно, чтобы машина смогла собрать все разрозненные упоминания — сайт, отзывы, статьи, профили — в **одну сущность** и уверенно отличить её от тёзок. Мы уже видели раньше: как часто упоминают бренд, для цитируемости важнее, чем сколько на него ведёт ссылок. Но здесь копнём глубже. Упоминания работают, только если они склеиваются в один узнаваемый объект. Иначе получается та самая история с фигурками вымышленного персонажа.

Машиночитаемый паспорт бренда

Как же помочь машине собрать вас в одну сущность? Хорошая новость: для этого есть вполне конкретные, технические инструменты, а не только «станьте знаменитым».

Первый — **Wikidata**. Это свободная база знаний, которую одновременно читают и правят люди и машины; она же — общее хранилище структурированных данных для проектов Викимедиа, включая Wikipedia. Устроена она как огромный набор фактов вида «объект → свойство → значение». У каждого объекта — уникальный номер с буквой Q, у каждого свойства — номер с буквой P. Запись «Дуглас Адамс — это человек» выглядит так: Q42 → P31 → Q5. ⁷ По сути это машиночитаемый паспорт: карточка, где про сущность записаны проверяемые факты в формате, который бот прочтёт без ошибок.

Второй инструмент связывает ваш сайт с этим паспортом — разметка *schema.org* для типа *Organization*, а в ней свойство *sameAs*. Это способ прямо, машинным языком заявить: «сущность на моём сайте — то же самое, что вот эта карточка в Wikidata, вот эта статья в Wikipedia и вот эти официальные профили в соцсетях». Официальная документация Google относит *sameAs* — рядом с названием, адресом сайта и логотипом — к рекомендованным свойствам разметки организации. ⁸ Проще говоря, вы своими руками сшиваете разбросанные упоминания в один узел и говорите поисковику: это всё я, не путайте.

Третий инструмент старый как локальное SEO, но никуда не делся — **NAP**. Три буквы: Name, Address, Phone — название, адрес, телефон. Это те поля, по которым бизнес опознают в справочниках, на картах и в отзывах. Каждый раз, когда эти данные всплывают где-то онлайн — у вас, у партнёров, в СМИ, — появляется «цитирование» (*citation*). А NAP-единообразие означает, что название, адрес и телефон **совпадают** во всех этих местах до буквы. ⁹ И это важнее, чем кажется. Логика поисковика простая: если десятки источников называют бизнес оди-

наково, данным можно верить и не страшно показывать их человеку. А если адрес то такой, то всякой, телефон в одном месте один, в другом другой — возникает сомнение, и доверие падает.⁹ Тот же BrightLocal спрашивал об этом самих потребителей: неточные данные в справочниках раздражают почти всех (93%), а если контакты в разных местах не совпадают, бизнесу перестают доверять восемь человек из десяти.¹⁰ Машина реагирует ровно так же, только молча — она просто не собирает вас в надёжную сущность.

Отсюда первое конкретное действие, из самых дешёвых и недооценённых. Откройте табличку и выпишите, как ваша компания называется на сайте, в Google, в Яндексe, на картах, в каждом отзовике и каталоге. «ООО Ромашка», «Ромашка», «Romashka», «Студия Ромашка» — если вы насчитали три-четыре варианта написания и два телефона, вы только что нашли, почему машина вас путает. Приведите всё к одному эталону. Это скучно и не гламурно, зато работает.

Wikipedia нельзя взломать

Тут в голове читателя обычно вспыхивает соблазнительная мысль: раз Wikipedia так важна, надо просто завести себе статью. Сесть и написать. На этом обжигаются регулярно.

Wikipedia и Wikidata живут по очень разным правилам, и разница принципиальная. У Wikidata порог мягкий — для большинства брендов это вполне реальная цель. А вот у Wikipedia планка высокая: статья допустима, только если тема получила **существенное освещение в надёжных источниках, независимых от неё самой**. И главное, что стоит принять раз и навсегда: Wikipedia прямо называет себя **запаздывающим индикатором** значимости.¹¹ Статья в энциклопедии не создаёт значимость — она её только фиксирует, когда внешний мир уже обратил на вас внимание. Нельзя стать значимым, написав про себя статью. Можно только заслужить значимость в реальном мире, а статья придёт следом.

Почему энциклопедию нельзя взломать, чем кончаются попытки срезать путь и что с этим делать — отдельная глава 14: там и правила значимости в деталях, и поучительные истории тех, кто попробовал. Здесь достаточно вывода: карточку знаний нельзя построить, дёргая один рычаг. Нужна сеть источников, которые говорят о вас одинаково и независимо друг от друга.

Дом сущности

Среди специалистов по сущностям ходит удобная практическая идея — термин **entity home**, «дом сущности». Это одна страница, которую вы назначаете якорем: обычно раздел «О компании» (About). Именно по ней боты, модели и люди понимают, кто вы такие, и от неё отталкиваются, разбирая все остальные упоминания.¹² У этой страницы одна задача — быть эталонным, недвусмысленным описанием сущности. Полное название, чем занимаетесь, где находитесь, ключевые факты, а в разметке — тот самый sameAs со ссылками на Wikidata и профили. Это ваша точка сборки: отсюда машина начинает распутывать, что все эти разбросанные по сети упоминания — про один и тот же объект.

А теперь — почему без этой сборки бренд буквально растворяется. Вернёмся к источнику нашей истории про фигурки. В том же разборе описан второй случай: американская консалтинговая фирма и британская компания с почти таким же названием. Нейросети смешивали их истории в один усреднённый образ — часть фактов от одной, часть от другой.¹ И по запросам, где потенциальные клиенты сравнивали варианты, американская фирма фактически была **невидима под собственным именем**: вместо неё модель показывала гибрид из двух компаний. Это и есть entity ambiguity, неоднозначность сущности, во всей красе. Не оштрафовали, не понизили — просто не смогли отделить вас от двойника, и вы пропали.

Чего наука пока не измерила

Всё, что выше, — про механику: граф знаний устроен вокруг сущностей, Wikidata кормит карточки Google, Wikipedia весит непропорционально в обучающих данных, а путаница

сущностей ломает видимость на живых кейсах. Логика прочная, и косвенных подтверждений хватает.

Но прямого, независимого, названного по имени исследования, которое измерило бы именно это — «навёл порядок с сущностью, и цитируемость в ChatGPT или Perplexity выросла на столько-то процентов», — на момент написания книги нет. Не «я не нашёл», а его пока не существует. Отрасль это признаёт сама: авторы разбора про путаницу брендов прямо пишут, что на февраль 2026 года ни одна опубликованная работа не изучает конфликтующие значения брендов в ИИ-платформах — тема обогнала академию.¹ Так что относитесь к приёму трезво: это хорошо обоснованная, подтверждённая кейсами практика, а не измеренное число вроде тех «плюс 40%», которые мы ещё увидим у статистики и цитат.

Российский контур

Есть один нюанс — он важен для Рунета: здесь свой граф знаний и свои площадки.

У Яндекса есть прямой аналог карточки знаний — **объектный ответ**: справка сбоку от выдачи с краткой информацией о предмете запроса, будь то человек, компания, город или фильм. Строится она на семантическом графе Яндекса. Показывать такие карточки Яндекс начал в апреле 2015 года, и уже на старте база содержала описания около 92 миллионов сущностей — личностей, фильмов, альбомов, городов, лекарств, машин.¹³ Данные Яндекс собирает из десятков источников, включая Википедию, сам сопоставляет и чистит факты, отсеивая дубли и противоречия. То есть та же самая идея: система хочет понять, о какой вещи речь, и собрать про неё непротиворечивый паспорт.

Практика для российского бизнеса перекликается с мировой. Данные об организации — адрес, телефон, часы — можно передать Яндексу через **Яндекс Бизнес** или через ту же разметку `schema.org/Organization` на сайте. Причём если данные из разных источников не совпадут с тем, что вы указали в Бизнесе, заявку на публикацию могут просто отклонить.¹⁴ Снова то же правило единообразия, только уже с конкретным механизмом наказания за разноречие.

Отдельно — про энциклопедии, потому что тут легко запутаться. Кроме Wikipedia, в русскоязычном поле теперь есть **Рувики** — созданный в 2023 году форк русской Википедии; при запуске он перенёс к себе почти два миллиона статей — свободная лицензия это разрешает, — а часть из них переписал под российское законодательство; и в отличие от некоммерческой Wikipedia это коммерческий проект со своей инфраструктурой в России; официальный запуск — январь 2024 года.¹⁵ Есть и **Знание. Вики** — энциклопедия общества «Знание», ещё один отдельный источник карточек.¹⁶ Держите в голове простую вещь: это самостоятельные проекты со своими правилами, а не «российские зеркала» Wikipedia с теми же критериями. Как именно каждая российская нейросеть — Алиса, GigaChat, доступный в стране DeepSeek — опирается на эти энциклопедии и графы, отдельно и надёжно пока не измерено: общая логика сущностей работает и здесь, но точных цифр по каждой платформе нет.

Мини-шаблон: как оформить свою сущность

Соберём практику в короткий маршрут. Пройдите по нему один раз для своего бренда — большинство шагов делаются за вечер и не стоят денег.

Шаг 1. Сведите NAP к одному эталону. Выпишите название, адрес и телефон так, как они указаны на сайте, в Google, в Яндексе, на картах и во всех каталогах. Найдите расхождения и приведите к единому написанию везде. Это фундамент — без него остальное шатается.

Шаг 2. Назначьте дом сущности. Сделайте страницу «О компании» эталонным описанием: полное имя, чем занимаетесь, где находитесь, ключевые проверяемые факты. Одна страница, к которой всё остальное отсылает.

Шаг 3. Свяжите профили через sameAs. Добавьте на сайт разметку `schema.org/Organization` и в свойстве `sameAs` перечислите ссылки на свои профили — соцсети, каталоги, а если есть, то карточку Wikidata и статью Wikipedia. Так вы сами говорите машине: это всё один объект.

Шаг 4. Заведите или проверьте паспорт в базах знаний. Убедитесь, что вы есть в Яндекс Бизнесе с верными данными. Wikidata с её мягким порогом — реальная цель для большинства; Wikipedia — только когда о вас уже есть независимые серьёзные публикации, не раньше.

Шаг 5. Не пытайтесь взломать. Не пишите статью о себе в Wikipedia сами — чем кончаются такие попытки, увидите в главе 14. Значимость зарабатывают в реальном мире, а энциклопедия её лишь фиксирует.

Шаг 6. Проверьте на тёзок. Спросите у нейросетей про свой бренд по имени и посмотрите, не подмешивают ли они чужую сущность — персонажа, другую компанию, однофамильца. Если да, это сигнал усиливать однозначность: полнее имя, чётче дом сущности, больше независимых упоминаний, где вас нельзя спутать.

Итог главы

Нейросети и поиск думают не строками, а сущностями — «вещами, а не строками», как сформулировал Google ещё в 2012 году.² Чтобы попасть в ответ, мало засветиться в интернете под своим названием: машина должна собрать все ваши упоминания в один узнаваемый объект и отличить его от тёзок. Помогают этому вещи конкретные и в основном бесплатные — единый NAP, дом сущности, разметка sameAs, паспорт в Wikidata. А Wikipedia нельзя взломать: значимость зарабатывают снаружи, статья приходит следом. Прямого измеренного числа у этого приёма пока нет: это практика, подтверждённая кейсами, и не более того.

Если унести из главы одно: **сделайте так, чтобы на вопрос «кто это?» и человек, и машина отвечали одинаково и без запинки.** Одна сущность, одно имя, один паспорт.

Сделайте сейчас одну вещь: спросите у любой нейросети, что она знает о вашем бренде по имени. Если в ответе всплыл тёзка или перепутанные факты — вы нашли свою первую задачу.

А дальше — самое интересное. Сущность собирается не в одном месте, а по всему интернету сразу: из сайта, форумов, энциклопедий, отзывов, чужих публикаций. Только вот интернет для нейросети — не сплошное поле, а несколько узких пятачков, за пределами которых она почти не смотрит. Где именно эти пятачки — с этого начнём следующую главу.

Глава 10. Карта площадок 2026: где ИИ берёт источники

Третий по цитируемости сайт в ChatGPT — это сам ChatGPT

Компания Similarweb разобрала около 600 тысяч случаев, когда нейросети ссылались на источники, — данные за январь—февраль 2026 года, опубликованы 14 апреля.¹ Наверху у ChatGPT в режиме веб-поиска стоят Википедия — 13,15% всех цитирований — и Reddit с его 11,97%. А третьим, обгоняя Walmart, YouTube и LinkedIn, идёт openai.com — 6,21%.¹

Третий по цитируемости домен в ChatGPT — собственная документация OpenAI. Модель ссылается на саму себя втрое чаще, чем на каждый из таких столпов, как Reuters, NIH или LinkedIn, — у любого из них лишь пара процентов.¹ Уточню сразу, чтобы не приукрашивать: втроём эти трое всё-таки дают чуть больше, чем openai.com в одиночку. И это не причуда одного ChatGPT: Google в своём AI Mode так же исправно тащит google.com в собственную пятёрку источников.¹ Похоже, первым делом ИИ доверяет самому себе.

Но удивляет даже не это. Удивляет, до чего узок весь остальной круг. По оценке отраслевого агрегатора Everything-PR (о нём — с оговоркой ниже), всего пятнадцать доменов забирают около 68% совокупной доли цитирований по пяти крупным платформам сразу — ChatGPT, Claude, Gemini, Perplexity и Google AI Overviews.² Пользовательский контент — Reddit, YouTube, LinkedIn, Quora — в сумме цитируется чаще, чем два десятка крупнейших журналистских изданий вместе взятых.²

Никакого единого «интернета», из которого нейросеть черпает факты, не существует. Есть несколько узких, почти не пересекающихся пулов площадок — и они разные для мира и для России. Эта глава — карта этих пулов: какие именно площадки, с какими долями, на какую дату. Механику отбора мы уже разобрали раньше (два сита, у каждой платформы свой поисковый движок под капотом) — повторять не будем. Здесь — конкретные списки и, главное, ответ на вопрос: как из всего этого выбрать свои десять-пятнадцать площадок и не распылиться.

Все цифры ниже — снимки конкретного месяца, поэтому у каждой стоит дата: доли источников штормит, и к выходу книги они наверняка сдвинутся. Карта верна по направлению, а не по второму знаку после запятой.

Мировой топ: горстка доменов забирает почти всё

Начнём с глобальных платформ. Первое, что бросается в глаза: у каждой — свой любимый набор. И второе, важнее первого: концентрация тут жёстче, чем в обычном поиске. Классический Google размазывает внимание по десяткам сайтов на первых страницах выдачи; ИИ же не строит топ-10 ссылок, а вылавливает горстку источников, из которых можно собрать цельный ответ, — и эта горстка сильно смещена к узкому набору доменов.² Попасть в неё труднее, но и площадок, за которые стоит бороться, меньше.

ChatGPT тянется к энциклопедиям и деловым медиа. По тому же замеру Similarweb (январь—февраль 2026) его топ возглавляют Википедия и Reddit, дальше — собственная документация OpenAI, Walmart, YouTube, LinkedIn, Reuters, NIH.¹ Более ранний и более крупный разбор Ahrefs на 9,6 миллиона запросов (сентябрь 2025) рисует похожую картину: в глобальном топе ChatGPT первые места делят Reddit и Википедия — причём Википедия сразу на нескольких языках: английская, испанская, немецкая.³ Русская Википедия в том глобальном рейтинге тоже есть — на 78-м месте.³ Запомните эту деталь, к российскому контуру она ещё пригодится.

Google AI Overviews держится на видео и обсуждениях. Ahrefs ведёт живую, ежемесячно обновляемую таблицу пятидесяти самых цитируемых доменов в ИИ-сводках Google. На снимке июня 2026 года впереди YouTube — 20,9% — и Reddit с его 19,6%, третьим идёт Facebook, а дальше — Google, Instagram, Википедия, Amazon, Quora и TikTok, у каждого лишь единицы процентов.⁴ (Facebook принадлежит компании Meta, признанной в России экстре-

мистской организацией и запрещённой; то же касается Instagram — здесь и далее обе площадки упоминаются только как источники данных.) Только первая тройка — YouTube, Reddit, Facebook — забирает больше половины всей доли, 52,1%.⁴ Видео и обсуждения перевешивают всё остальное.

Отдельная поучительная деталь. В режиме Google AI Mode, по замеру Similarweb, неожиданно лидирует fandom.com — вики-платформа фанатских сообществ, с долей 7,16%, впереди самой Википедии.¹ Не потому, что fandom известнее, а потому, что у него глубокие, подробные страницы по узким темам. Глубина на конкретной теме бьёт общую известность — это принцип, который нам ещё понадобится, когда будем выбирать площадки.

А вот как расставляют приоритеты пять платформ, если свести их в одну таблицу. Здесь нужна оговорка: детальную разбивку по Perplexity, Claude и Gemini даёт синтез Everything-PR / 5W — это не первичное исследование, а маркетинговый обзор PR-агентства, которое само продаёт услуги ИИ-видимости.² Цифры оттуда берите как порядок величины, а не как точный замер. Строки по ChatGPT и Google подтверждены первичными данными Ahrefs и Similarweb.

ChatGPT. Википедия (доминирует), деловые медиа (Forbes, Business Insider, Reuters), Amazon, собственная документация OpenAI.

Google AI Overviews / AI Mode. YouTube (лидер с отрывом), Reddit, продукты Google; в AI Mode неожиданно — Fandom.

Perplexity. Reddit, NIH/PubMed, Britannica, LinkedIn, Quora — крен в первоисточники и живой поиск.

Claude. Длинная авторитетная журналистика: The New York Times, The Atlantic, The Economist, Financial Times, The Washington Post.

Gemini. YouTube, продукты Google, Associated Press, The Guardian.

Источники: Ahrefs и Similarweb (ChatGPT, Google); синтез Everything-PR / 5W (Perplexity, Claude, Gemini), апрель—июнь 2026.^{1 2 4}

Две площадки в этом списке заслуживают отдельной пометки, потому что они — самое живое, что происходило в 2026 году. Первая — **LinkedIn**: как мы уже видели раньше, он за три месяца ворвался в топ цитирований ChatGPT почти с окраины и стал источником номер один по профессиональным запросам сразу на шести платформах.⁵ Для B2B это теперь обязательная площадка. Вторая — **площадки отзывов**: G2, Capterra, Trustpilot. Про их вес мы тоже говорили — развитый профиль поднимает цитируемость в разы, а G2 через поглощение конкурентов нарастил долю в B2B-запросах.⁶ Здесь важно лишь зафиксировать их место на карте: для отзывов и сравнений софта это отдельный, обязательный пул.

Российский топ: другая карта той же страны

Теперь пересечём границу — и карта меняется почти полностью. И это не мелочь для узкого круга: как мы видели раньше, в топе российских нейросетей устойчиво держатся Яндекс и GigaChat наравне с ChatGPT, а значит, для клиентов в России отдельная карта — не приятное дополнение, а половина всей работы.

Главный структурный факт мы уже разбирали: «Поиск с Алисой» — не отдельный поисковик, а надстройка над обычной выдачей Яндекса. Отсюда всё остальное. Как ранее выяснили по замерам МОАВ, почти две трети ссылок в нейроответе Алисы одновременно стоят в самой макушке органики, а условие ещё жёстче: нет сайта в топ-20 обычной выдачи — не будет и ссылки в нейроответе, вообще.⁷ И по данным «Ашманов и партнёры» доля источников из-за пределов топ-10 продолжает падать: девять ссылок из десяти сегодня берутся с первой страницы.⁸

Вывод для российского рынка простой и железный: **карта площадок в Рунете начинается не с площадок, а с позиций в Яндексе.** Не в топе органики — вас для Алисы нет. Это фундамент, на котором стоит всё остальное.

Но и внутри этого фундамента есть свои приоритетные «доноры» — площадки, которые Алиса подтягивает особенно охотно: Википедия, Яндекс. Карты, отзовики.⁹ И тут же — парадокс, ради которого стоит остановиться. В топ-20 сайтов, которые Алиса цитирует, нет ни одного маркетплейса. Ни одного.¹⁰ При этом в российских нейросетях в целом маркетплейсы забирают около трети всех цитирований — три крупнейшие площадки суммарно дают порядка 31%.¹¹ Одна и та же страна, один и тот же товар — но спросите Алису, и маркетплейсов в источниках не будет, а спросите «средний» российский ИИ, и они займут треть.

Почему так выходит? Потому что «средний российский ИИ» — это фикция, удобная для отчёта, но не для стратегии. За усреднённой цифрой стоят очень разные машины: Алиса, намертво привязанная к органике Яндекса, где карточки маркетплейсов ранжируются иначе; GigaChat со своим индексом; глобальные модели со своими пулами. Треть цитирований на маркетплейсы набирается за счёт одних платформ при почти нулевом вкладе других. Вывод практический: при выборе ориентируйтесь в первую очередь на то, какие нейросети в приоритете у вашей целевой аудитории. Если это Алиса — карточка на Wildberries вас в её ответ не приведёт, нужен топ Яндекса.

С сообществами в Рунете отдельная история. У Reddit здесь нет единого аналога — по наблюдению сервиса-монитора GEO Scout, его роль выполняет комбинация из четырёх площадок: Pikabu, DTF, Habr и vc.ru.¹² Ни одна поодиночке не дотягивает до влияния Reddit, но вместе они дают сопоставимый сигнал. Тот же отчёт (напомню — это вендорский отчёт компании, продающей мониторинг ИИ-видимости, цифры берите с этой поправкой) перечисляет площадки, которые российские нейросети цитируют активнее всего помимо YouTube и Википедии: vc.ru — по темам финтеха, e-commerce, SaaS; Habr — по технологиям и хостингу; Яндекс. Дзен — по широким запросам; Pikabu — по пользовательскому опыту; DTF — по играм и медиа; RuTube — как видеоальтернатива; Отзовик и iRecommend — по отзывам; Авито.¹² По их же данным, бренды, которые активно присутствуют на Habr и vc.ru, цитируются в технических нишах на 23—31% чаще тех, кто публикуется только на своём сайте, — но системно на этих площадках работает лишь около четверти компаний.¹² То есть комьюнити-сигнал у большинства просто не закрыт.

Ещё одна платформа со своим набором — **GigaChat от Сбера**. По обзорным материалам (первичной технической документацией это не подтверждено, поэтому — с оговоркой), он приоритизирует русскую Википедию и Викиданные, материалы Дзена с подтверждёнными авторами, открытые государственные реестры и официальные сайты в зоне gov.ru, а также СМИ из реестра Роскомнадзора.¹³ Своя логика, отличная и от Алисы, и от глобальных моделей.

И, наконец, **локальный бизнес**. Для него на российской карте главные — карты и справочники: Яндекс. Карты в связке с Яндекс. Бизнесом, 2ГИС, отзовики. Без заполненного и согласованного профиля здесь Алиса просто не назовёт вас по локальному запросу. Ключевое слово — согласованного: название, адрес и телефон должны совпадать дословно на всех площадках, иначе система не свяжет их в один бизнес.

Сравните две карты. Глобально работают YouTube и Википедия — они одинаково сильны везде. Но дальше пути расходятся: там, где в мире стоит Reddit, в Рунете — россыпь из Pikabu, DTF, Habr и vc.ru; где в мире Trustpilot и Amazon, в России — Отзовик, Ozon и Wildberries; и над всем этим в Рунете висит жёсткая привязка к топу Яндекса, которой у глобальных платформ нет. Это не одна карта с переведёнными названиями. Это две разные карты.

Почему пулы почти не пересекаются

Возникает законный вопрос: а нельзя ли угодить всем сразу? Почти нет — и доказательства такие.

Мы уже приводили цифру Ahrefs: в среднем лишь около 12% ссылок, которые дают ИИ-ассистенты, попадают в топ-10 обычной выдачи Google по тому же запросу.¹⁴ За этим средним прячется разброс по системам: у Perplexity пересечение 28,6%, а у Gemini, Copilot и ChatGPT

— всего 6—8,6%.¹⁴ Perplexity ближе всех к классическому поиску, остальные идут своей дорогой.

Это не единичный замер. Академическая работа исследователей из Университета Торонто (arXiv 2601.16858, конференция EDBT/ICDT 2026) на тысяче ранжирующих запросов даёт другие по величине, но те же по направлению числа: у всех проверенных моделей пересечение с топ-10 Google измеряется единицами процентов, и лишь у Perplexity подбирается к пятнадцати.¹⁵ Разная методология, разный период — а вывод один: пересечение везде низкое, и Perplexity стабильно ближе всех к обычному поиску.

Но самый наглядный замер — у BuzzStream (обновлён в июне 2026): 30 тысяч цитирований, 595 запросов, четыре системы — Google AI Overviews, AI Mode, Gemini и ChatGPT.¹⁶ Результат: **76,1% всех цитирований уникальны для одной-единственной платформы**, и лишь **0,8% ссылок встречаются одновременно во всех четырёх системах**.¹⁶ Сильнее всего между собой пересекаются два продукта Google — AI Overviews и AI Mode (37,4%), — а ChatGPT почти ни с кем не совпадает.¹⁶ Из той крохотной доли, что цитируют все четыре платформы разом, 35% приходится на Википедию.¹⁶ То есть общего у платформ ровно столько — общая энциклопедия и почти ничего больше.

Что это значит на практике, легко представить на одном материале. Экспертная статья, которую с удовольствием цитирует ChatGPT, может остаться полностью невидимой для Алисы — если она не в топе Яндекса. Пост на vc.ru, который Алиса берёт в ответ, может ни разу не всплыть у Perplexity с её собственным индексом. Один и тот же текст на разных платформах живёт разной жизнью, и это не исключение, а правило: до трёх четвертей цитирований, как мы только что видели, уникальны одной системе.

Отсюда прямой вывод, который определит всю вашу стратегию: традиционный успех в поиске и цитируемость в ИИ — почти параллельные вселенные, и сами платформы между собой пересекаются слабо.¹⁰ Один материал, который «зайдёт всем», сделать нельзя — он рискует не зайти никому. Придётся выбирать пулы осознанно и вести глобальную и российскую дорожки порознь.

Как выбрать свои 10—15 площадок

Теперь к главному. Площадок в карте — десятки, а сил у вас на все не хватит. Хорошая новость: их и не нужно. Раз пятнадцать доменов забирают 68% всех цитирований, вам достаточно попасть в правильные десять-пятнадцать — не пятьдесят. Всё дело в том, чтобы выбрать те, что кормят ответы именно в вашей нише и на вашем рынке. Разложим выбор по трём шагам.

Шаг 1. Определите рынок: Россия или мир. Это первая развилка, и она задаёт всю карту. Если ваши клиенты в России, отправная точка — Алиса и Яндекс, а значит, сначала классическое SEO, потом уже площадки. Если аудитория глобальная — ChatGPT, Perplexity, Gemini и их пулы. Пытаться накрыть оба рынка одним набором площадок бессмысленно: пулы почти не пересекаются.

Шаг 2. Определите тип бизнеса: B2B, локальный или e-commerce. Внутри каждого рынка пул зависит от ниши. Ниже — рабочие наборы под типовые сценарии. Это ориентиры на снимок 2026 года, а не вечная константа.

Глобальный B2B / SaaS. Собственный сайт со сравнениями, Википедия, Reddit, LinkedIn, G2, Capterra, YouTube, Quora, профильные деловые медиа (Forbes, TechRadar).

Глобальный B2C / e-commerce. Собственный сайт, YouTube, Reddit, Википедия, Trustpilot, Amazon, Instagram/TikTok, тематические обзорные площадки (тип Fandom для ниши).

Российский B2B / технологии. Топ Яндекса (свой сайт), vc.ru, Habr, Яндекс. Дзен, русская Википедия, YouTube и RuTube, профильные Telegram-каналы и форумы.

Российский e-commerce. Топ Яндекса, Ozon, Wildberries, Яндекс. Маркет, Отзовик и iRecommend, Яндекс. Дзен, YouTube.

Российский локальный бизнес. Топ Яндекса, Яндекс. Карты и Яндекс. Бизнес, 2ГИС, Отзовик и iRecommend, русская Википедия (если есть значимость).

Шаг 3. Учтите характер платформ внутри пула. Perplexity жаден до свежего и первоисточников — держите материалы актуальными. Claude любит длинную авторитетную журналистику — туда пробиваются через публикации в качественных медиа. Google тянется к видео — значит, YouTube. Алиса живёт топом Яндекса — значит, SEO. Это тонкая настройка поверх выбранного набора, а не вместо него.

Выбрали набор — теперь расставьте внутри него порядок, потому что делать всё сразу тоже не выйдет. Начинать с той площадки, где два условия совпадают: вас там сегодня нет, и при этом она несёт большую долю пула вашей ниши. Для российского локального бизнеса это почти всегда карты и топ Яндекса, для глобального B2B — LinkedIn и профильная площадка отзывов, для e-commerce — маркетплейсы и отзовики. Дешёвые действия с доказанным эффектом — заполненный профиль, собранные отзывы, понятная карточка — ставьте вперёд дорогих. А площадки, которые требуют долгой работы (публикации в авторитетных медиа, присутствие в сообществах), запускайте параллельно, но не ждите от них мгновенного результата.

И два правила, которые уберегут от типичных ошибок. Первое: **не распыляйтесь**. Сובлазн завести профили на всех сорока площадках из карты велик, но силы уйдут в песок. Лучше десять площадок, где вы реально присутствуете и обновляетесь, чем полсотни мёртвых профилей. Второе: **сверяйтесь с реальностью по своим запросам**. Карта в этой главе — общая; ваша ниша может вести себя иначе. Возьмите пять-десять запросов, по которым клиент должен находить именно вас, и посмотрите, кого по ним реально цитируют ваши целевые нейросети. Делать это можно руками, ведя простую таблицу, или через сервис мониторинга, который сам прогоняет запросы и показывает, какие источники сейчас в ходу у каждой нейросети. Так вы уточните общую карту под свою конкретную нишу — и не будете строить присутствие вслепую.

Итог главы

Единого «интернета», из которого ИИ берёт факты, нет. Есть узкие, почти не пересекающиеся пулы площадок: пятнадцать доменов забирают около 68% глобальных цитирований, у каждой платформы свой любимый набор, а российская карта — вообще отдельная, с жёсткой привязкой к топу Яндекса и своей россыпью площадок вместо Reddit. Три независимых замера — Ahrefs, работа Университета Торонто и BuzzStream — сходятся: пулы платформ пересекаются слабо, до 76% цитирований уникальны одной системе.

Если унести из главы одну мысль, пусть будет эта: **не пытайтесь быть везде — выберите десять-пятнадцать площадок под свой рынок и нишу и займите их всерьёз**.

Сделайте сейчас одну вещь: определите по двум развилкам — Россия или мир, B2B, или локальный, или e-commerce — свой сценарий из таблицы, выпишите эти десять-пятнадцать площадок и отметьте, на скольких из них вас сегодня вообще нет. С самой важной из пустых клеток и начнёте.

Мы разложили карту чужих площадок. Но у неё есть центр, с которого всё начинается и который целиком в вашей власти, — ваш собственный сайт. О том, как переделать его так, чтобы нейросеть легко вас читала и цитировала, — в следующей главе.

Глава 11. Ваш сайт: как переделать под ИИ

Идеальная статья, которую никто не прочитал

Представьте: вы неделю писали материал. Вылизали каждый абзац, собрали цифры, взяли комментарий у эксперта, десять раз перечитали. Публикуете и ждёте, что нейросети начнут вас цитировать.

А они молчат. Не потому, что текст плохой. А потому, что бот ИИ до него либо вообще не добрался, либо добрался — и увидел пустую страницу.

Оба провала до обидного частые. Первый: ваш хостинг или защитный экран молча заблокировал нужного бота. По одной оценке на выборке около трёх тысяч сайтов, примерно 27% случайно закрывают доступ хотя бы одному крупному ИИ-краулеру — и чаще не в своих настройках, а на уровне CDN или файрвола, о котором владелец даже не знает.¹ Цифру берите как сигнал, а не как закон — это один замер, в основном по бизнесу США и Британии, разошедшийся по блогам без независимого повтора. Но механизм реален: с 1 июля 2025 года Cloudflare, через сеть которого идёт около пятой части мирового трафика, по умолчанию блокирует ИИ-краулеров для новых доменов.² Вы могли ничего не запрещать — а бот всё равно не проходит.

Второй провал тоньше. Ваш сайт построен на JavaScript и собирает текст прямо в браузере посетителя. Человек видит статью, а бот — нет. И об этом — вся эта глава.

В прошлой главе мы разложили карту чужих площадок. Теперь — про единственную площадку, которая целиком в вашей власти: ваш сайт. Разговор чисто технический, поэтому сразу расставляю приоритеты, чтобы вы не утонули в мелочах. **Сначала — доступ ботов и извлекаемость. Потом — разметка. И только в самом конце, по остаточному принципу, — модные файлы вроде llms.txt.**

Приоритет первый: проверьте, что вас вообще пускают

Начинать надо не с красоты, а с двери. Если бот не может войти, всё остальное бессмысленно. А управляет дверью маленький текстовый файл в корне сайта — robots.txt. В нём вы построчно говорите каждому боту по имени: сюда можно, сюда нельзя.

И тут первое, что нужно понять: **у крупных провайдеров не один бот, а несколько, и у каждого своё имя и своё назначение.** Путают их постоянно, и ошибка тут дорогая: одной строчкой можно вычеркнуть себя из ответов.

Как мы уже видели раньше, боты разведены на два лагеря. Одни ходят по сайтам, чтобы собирать данные для будущего обучения моделей. Другие — чтобы формировать живые ответы прямо сейчас. Это разные боты с разными именами, и блокировать их можно раздельно.³ Вот кто есть кто на середину 2026 года.

OpenAI. GPTBot — сбор данных для обучения. OAI-SearchBot — живой поиск для ответов ChatGPT. Это два независимых правила: закрыли один — второй по-прежнему ходит.⁴

Anthropic. У Claude официально **три** бота, и это стоит запомнить дословно, потому что в сторонних каталогах их путают. ClaudeBot — сбор данных для обучения. Claude-SearchBot — индексация для поисковых ответов Claude. Claude-User — получение конкретной страницы по запросу живого пользователя в моменте. Так прямо и написано на странице Anthropic: «информация о трёх роботах, которых использует Anthropic», — без четвертого.^{5 6} Если вам где-то встретится «claude-code» в одном ряду с ними — это не бот Anthropic, а юзер-агент инструмента для программистов, отдельная история; в свой robots.txt его как «четвёртого краулера» не вписывайте.

Perplexity. PerplexityBot — краулер для индекса, из которого модель цитирует источники. Отдельно Perplexity-User — запрос страницы по инициативе пользователя.⁷

Google и Apple — а вот тут ловушка. Google-Extended и Applebot-Extended звучат как имена ботов, но это **не краулеры**. Сами по себе они не делают ни одного запроса к сайту.

Это токены-разрешения: ими вы говорите уже существующим роботам Google и Apple, можно ли использовать собранный контент для обучения ИИ — отдельно от обычной индексации.⁸ Заблокировать Google-Extended, надеясь «закрыться от ИИ Google», и при этом остаться в поиске — можно. А вот заблокировать самого Googlebot — значит выпасть из выдачи целиком.

Отсюда простое правило настройки. Раз обучающие и поисковые боты разведены, вы можете закрыть первых и оставить вторых. Не хотите кормить будущие версии моделей своим контентом — закройте GPTBot, ClaudeBot, Google-Extended. Но обязательно оставьте открытыми OAI-SearchBot, Claude-SearchBot и PerplexityBot — иначе вы своими руками вычеркнете сайт из живых ответов.³ В robots.txt это пишется по имени юзер-агента:

User-agent: OAI-SearchBot

Disallow:

User-agent: PerplexityBot

Disallow:

User-agent: GPTBot

Disallow: /

Пустой Disallow: — «ходи везде», Disallow: / — «не ходи никуда». Заметьте, к чему привела осторожность многих: по одному замеру, 39 из 107 топовых сайтов (36,4%) на середину июня 2026 года блокировали хотя бы одного бота Anthropic, и почти все блокировки пришлись именно на обучающий ClaudeBot.⁹ Методология этой выборки раскрыта не полностью, так что перед вами скорее ориентир, чем точный замер. Но тренд ясный: обучающих ботов закрывают охотно, и это нормально — лишь бы заодно не прихлопнуть поисковых.

Что сделать руками прямо сегодня. Откройте вашсайт.ру/robots.txt в браузере и прочитайте его глазами. Нет ли там Disallow: / для поисковых ботов? Затем проверьте уровень выше — CDN и защитный экран. Если сайт на Cloudflare, зайдите в настройки управления ИИ-ботами и убедитесь, что поисковые краулеры не перекрыты по умолчанию. Пять скучных минут — а от них зависит, существуете вы для нейросети или нет.

Приоритет второй: чтобы бот увидел текст, а не пустоту

Дверь открыта, бот вошёл. Теперь — увидит ли он ваш текст? Здесь и живёт вторая, самая коварная причина невидимости.

Многих этот факт застаёт врасплох. **ИИ-краулеры в массе своей не исполняют JavaScript.** У них нет браузерного движка, нет DOM, нет цикла событий — они просто скачивают тот HTML, что сервер отдал в ответ, и вытаскивают из него готовый текст.¹⁰ GPTBot, ClaudeBot, PerplexityBot работают именно так.

Многие современные сайты устроены как одностраничные приложения: сервер отдаёт почти пустую оболочку — заголовок, спиннер загрузки и ссылки на скрипты, — а сам текст статьи «дорисовывается» уже в браузере посетителя после загрузки этих скриптов. Человек видит полноценную страницу. А бот, который скрипты не запускает, получает ту самую пустую оболочку.¹⁰ Ваша идеальная статья для него просто не существует.

Проверить себя можно за минуту. Откройте страницу, нажмите правой кнопкой «Просмотр исходного кода страницы» (view-source) — и поищите в нём текст своей статьи. Есть текст в исходном HTML — бот его увидит. Нет, только скрипты — вы невидимка.

Лечится это известным способом, и он не такой страшный, как кажется. Нужен **серверный рендеринг** (SSR, server-side rendering) — когда сервер отдаёт уже готовую HTML-страницу с текстом внутри, а не собирает её в браузере. Это умеют делать фреймворки вроде Next.js, Nuxt или Angular Universal.¹¹ Если полный переход на SSR не по силам, есть более дешёвая замена — пре-рендеринг: сервер заранее готовит статичные HTML-версии страниц специ-

ально для ботов.¹¹ Смысл в обоих случаях один: важный текст должен лежать прямо в ответе сервера, а не появляться после того, как отработают скрипты.

Дальше — как этот текст разложить, чтобы его было легко «вынуть». Здесь напомним коротко механику, которую мы уже разбирали: генеративный поиск не читает страницу целиком, а заранее режет её на смысловые куски (chunking) и на этапе ответа достаёт тот кусок, что ближе всего к запросу.¹² Значит, ваша задача — нарезать текст за систему, а не оставлять это на её усмотрение. Три конкретных приёма.

Заголовки-вопросы. Оформляйте подзаголовки как прямые вопросы, которые задал бы человек: «Сколько стоит доставка?», «Чем отличается тариф Pro от Basic?». Вопрос в теге `<h2>` и следующие за ним два-три предложения ответа — это ровно тот шаблон, который система легче всего распознаёт и вырезает как готовую пару «вопрос-ответ».¹²

Ответная капсула. Практики GEO сходятся на ориентире: 40—80 слов на один самодостаточный ответ в начале смыслового блока — сначала прямой ответ, потом раскрытие.¹³ Тут важно не обмануться: это отраслевая эвристика, а не измеренный в эксперименте норматив. Используйте как рабочий ориентир длины, а не как магическое число. Каждая такая капсула должна быть понятна сама по себе, без «как мы писали выше» и провисающих «он», «это», «там» — чтобы её можно было процитировать, не таща за собой остальной текст.

Блок FAQ. Раздел «Частые вопросы» — это концентрат капсул: десяток прямых вопросов, под каждым — короткий самодостаточный ответ. Формат, будто созданный под нарезку.

И ещё одна вещь, почти бесплатная, — **семантический HTML**. Это значит размечать страницу осмысленными тегами вместо бесконечных безымянных `<div>`. `<header>` — шапка, `<nav>` — навигация, `<main>` — единственный на странице блок основного содержания, `<article>` — самостоятельный материал вроде статьи или обзора. Модели вроде ChatGPT разбирают несколько десятков осмысленных тегов заметно легче, чем сотни вложенных друг в друга безымянных «дивов».¹⁴ Парсеру не приходится гадать, где тут суть, а где меню и реклама.

Если из всего приоритета запомнить одно: сначала убедитесь, что текст есть в исходном HTML, потом нарежьте его на заголовки-вопросы и короткие капсулы. Извлекаемость складывается ровно из этого.

Приоритет третий: разметка и что она на самом деле даёт

Теперь про то, вокруг чего в GEO больше всего мифов, — микроразметку Schema.org в формате JSON-LD. Это блок структурированных данных, который вы прячете в код страницы: машиночитаемая карточка, где прямо сказано «это организация», «это товар», «это статья», «это раздел вопросов-ответов».

Начнём с новости, которую индустрия не любит. **Как рычаг ИИ-цитируемости разметка не доказана.** Ahrefs провёл контролируемый эксперимент: 1885 страниц, на которые добавили JSON-LD, против 4000 контрольных без разметки, восемь месяцев наблюдений. Значимого прироста цитируемости не нашли, а для Google AI Overviews зафиксировали даже небольшое снижение — на 4,6%.¹⁵ Это сильное доказательство: широкая выборка, контрольная группа, замер «до и после».

Объяснение есть, и оно правдоподобное. Независимый эксперимент SEO-практика Марка Уильямс-Кука в феврале 2026 года показал, что ChatGPT и Perplexity читают блок JSON-LD как обычный текст — построчно, как строку символов, а не как распознанную структуру с полями и значениями.¹⁶ У Google есть парсер структурированных данных, который эту карточку понимает по-настоящему; у чат-бота такого парсера нет, и он не вытащит из разметки ничего сверх того, что и так прочтёт в тексте страницы. Это один эксперимент практика, не академическая публикация, — но механизм согласуется с тем, что мы знаем о работе моделей с текстом.

Есть и данные с другой стороны. Препринт Питера Шанбахера (SSRN) нашёл связь разметки с видимостью в ChatGPT. Но смотрите на выборку: это **не сайты вообще, а 1508**

агентств недвижимости в Германии. На ней FAQPage-разметка оказалась сильнейшим предиктором видимости — отношение шансов около 13, а Product-схема заметно слабее, около 4.¹⁷ Звучит убедительно, но это корреляция на узкой нише, без контроля причинности. Сайты с продуманной разметкой обычно вообще технически более зрелые и структурированные — и именно эта общая зрелость может объяснять и разметку, и видимость разом. Переносить вывод с немецких риелторов на «сайты вообще» некорректно.

Так что перед вами не «два исследования подтверждают одно и то же», а разные по силе доказательства с разными выводами: широкий контролируемый эксперимент против узкой отраслевой корреляции. Трезвый вывод для вас такой: **разметка нужна — она работает на обычный поиск и на понимание вашего бренда машиной. Только роста ИИ-цитирований от неё не ждите.** Google прямо пишет, что использует структурированные данные, чтобы понять содержание страницы, и рекомендует именно JSON-LD как формат, где легче всего не наделать ошибок.¹⁸ Ставьте разметку ради этого. А под обещаниями «добавьте JSON-LD — и ИИ вас процитирует чаще» никакой научной базы нет.

Свежая история идеально это иллюстрирует. 7 мая 2026 года Google полностью выключил в выдаче FAQ-сниппеты — ту самую «гармошку» вопросов-ответов под ссылкой.¹⁹ Казалось бы, всё, разметку FAQPage можно снимать. Но Google тут же уточнил: сама разметка FAQPage остаётся валидной, ошибок не выдаёт, санкций не влечёт, и он продолжает её читать, чтобы понимать структуру страницы.²⁰ Исчез только визуальный сниппет — а причина ставить разметку осталась. Мораль отсюда простая: разметка нужна не для красивой строчки в выдаче, а для того, чтобы машина поняла структуру вашей страницы.

Какие типы реально стоит поставить, без фанатизма: Organization со свойством sameAs (уже разбирали с вами в главе про сущности), Article для статей, Product для товаров, FAQPage для разделов вопросов-ответов. Это разумный минимум, который помогает пониманию и ничего не стоит. Дальше в разметку закапываться незачем.

lms.txt: ставьте, если уже прибегаетесь, — но не как спасение

И только теперь, по остаточному принципу, — про файл, о котором в GEO-кругах спорят больше всего. lms.txt придумал Джереми Ховард, сооснователь Answer.AI, и предложил его в сентябре 2024 года.²¹ Идея: положить в корень сайта markdown-файл с кратким описанием и ссылками на «чистые» версии страниц — чтобы модель с ограниченным контекстом не продиралась через HTML с навигацией и рекламой. Важная деталь, которую все забывают: задумывался он для **документации к коду и API**, а не для сайтов брендов.²¹

А теперь цифры. Два крупных независимых замера сошлись в одном выводе. Trakkr просканировал 37 894 домена и не нашёл никакого преимущества в цитируемости у сайтов с этим файлом — статистический тест дал $p=0,85$, то есть связь на грани нуля.²² SE Ranking на выборке около 300 000 доменов пришёл к тому же: значимого влияния на частоту цитирования нет.²³ Джон Мюллер из Google высказался жёстче, сравнив lms.txt со старым мета-тегом keywords, который поисковики похоронили десять лет назад: «Я не думаю, что кто-то знает — пока это чистая спекуляция: файл существует не первый год, но ни одна ИИ-система его не использует».²⁴

Вывод трезвый. Файл бесплатный и ставится за пять минут, так что категоричного «не делайте» тут нет. Но и решающим приёмом его считать нельзя. Если вы уже наводите порядок в технике сайта — поставьте заодно, хуже не будет. Только не тратьте на него силы, которые нужны двери и извлекаемости. Это факультативная мелочь, а не спасение.

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.