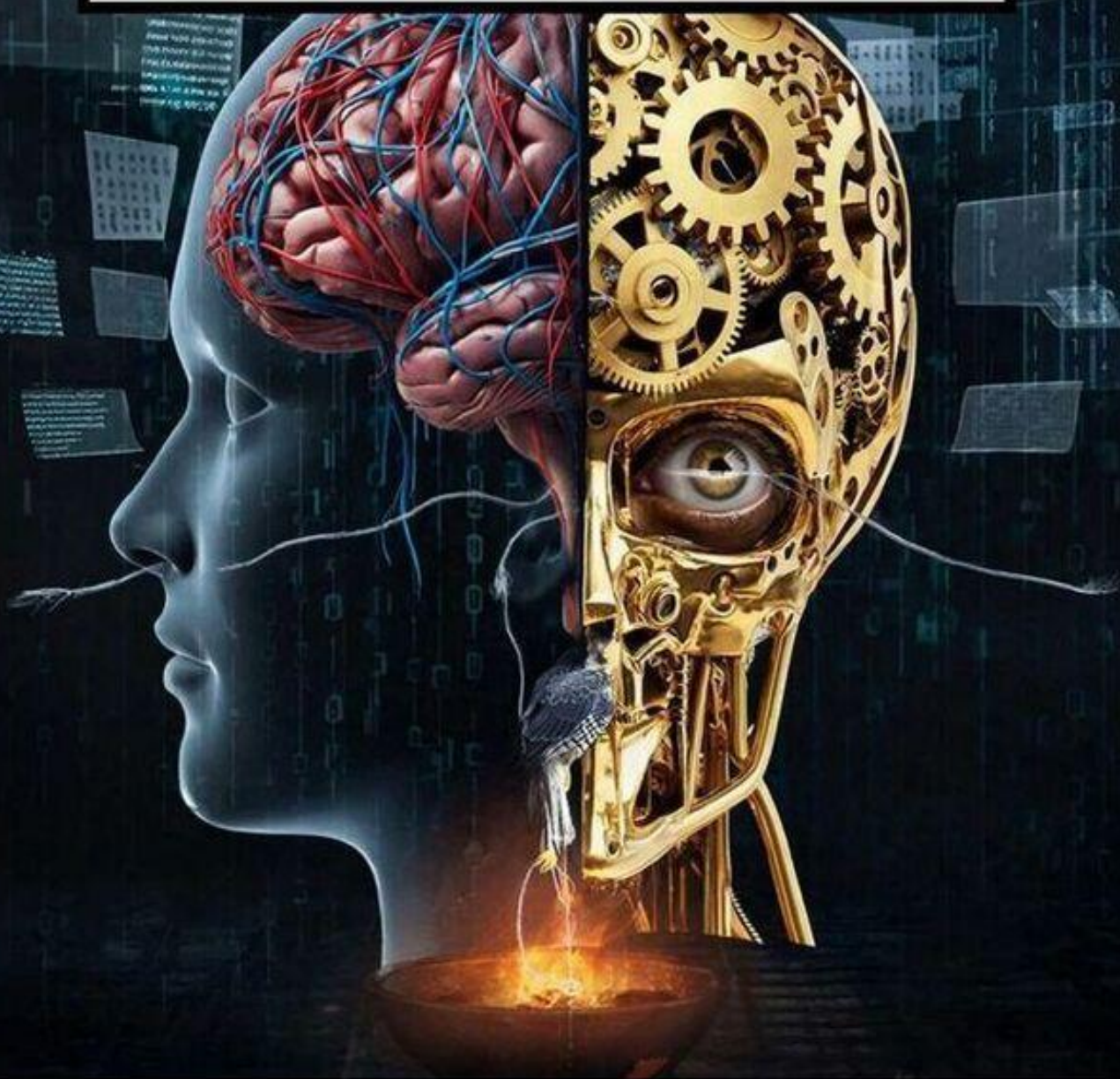


12+

АЛЕКСАНДР ЧИЧУЛИН

Великий алгоритм ЛОГИКИ

МОДУЛЬ 2. ИНСТРУМЕНТЫ ЛОГИКИ:
ОТ ФОРМУЛ К АЛГОРИТМАМ



Александр Чичулин

**Великий алгоритм логики.
Модуль 2. Инструменты логики:
от формул к алгоритмам**

«Издательские решения»

Чичулин А.

Великий алгоритм логики. Модуль 2. Инструменты логики: от формул к алгоритмам / А. Чичулин — «Издательские решения»,

Вы прошли диагностику. Вы знаете свои уязвимости. Теперь — время строить защиту. Этот модуль — кузница, а не академия. Никакой скучной теории. Только инженерная логика, которая превращает вас из пассивного пользователя ИИ в его критика, анализатора и хозяина. «Рабоми — не тот, кто пользуется ИИ, а тот, кто перестал задавать вопросы».

© Чичулин А.

© Издательские решения

Содержание

Глава 4. Логика как инженерия: не академия, а кузница	7
● 4.1. Диалог: Юрий объясняет, почему формальная логика — лучший щит от рабонии	7
Почему логика работает против ИИ	7
Три врага рабонии — три закона логики	7
Задание на паузу	8
● 4.2. Три закона логики для повседневности (закон тождества, непротиворечия, исключённого третьего)	9
Закон первый: тождества	9
Закон второй: непротиворечия	9
Закон третий: исключённого третьего	10
Проверочная таблица (без таблицы)	10
Упражнение «Найди нарушителя»	10
Ваше задание	11
● 4.3. Распознавание логических ошибок в текстах, сгенерированных ИИ	12
Типичные логические ошибки генеративных нейросетей	12
Живой разбор: AI против логики	13
Тактика защиты: три шага	13
Практическое задание	14
Конец ознакомительного фрагмента.	15

Великий алгоритм логики

Модуль 2. Инструменты логики: от формул к алгоритмам

Александр Чичулин

© Александр Чичулин, 2026

ISBN 978-5-0070-2927-8 (т. 2)

ISBN 978-5-0069-8802-6

Создано в интеллектуальной издательской системе Ridero

Цель модуля: Освоить логический аппарат, необходимый для проектирования алгоритмов, и создать «Логический конструктор» — набор шаблонов анализа.

Там, где бессилён эмоциональный порыв, где нас запутывают стремительные потоки данных, единственной опорой остаётся холодная, строгая, инженерная логика. Мы перестанем быть просто пользователями алгоритмов — мы станем их создателями и, что важнее, их бесстрашными критиками. В этом модуле мы выкуем три инструмента: умение видеть логические ошибки в машинных текстах, способность строить деревья решений под натиском неопределённости и мастерство непробиваемой аргументации в споре с цифровыми провокаторами.

Вступление: зачем логика тому, кто уже умеет задавать вопросы?

Вы прошли диагностику. Вы знаете свои триггеры, составили карту уязвимостей, научились замечать момент подмены. Это как понять, что в вашем доме есть щели, через которые дует. Но этого мало. Теперь нужно научиться заделывать эти щели — и сделать это так, чтобы они больше не появлялись.

Логика — это не скучный школьный предмет. Это инженерный каркас вашего суверенитета. Если «Детектор шума 3.0» и 5 вопросов помогали распознавать врага, то логика учит вас строить непробиваемую оборону и точное оружие.

Что изменится после этого модуля:

- Вы перестанете путаться в чужих аргументах — потому что научитесь видеть их внутреннюю структуру.

- Вы сможете оценивать любые прогнозы (новости, аналитика, заявления политиков) — с помощью деревьев вывода и сценарного моделирования.

- Вы будете парировать манипуляции AI-ботов и провокационные промпты — с помощью «Аргумент-клинка 2.0».

- Вы создадите свой «Логический конструктор» — набор шаблонов, который пополнит ваш «Формулярный» и станет основой для проектирования алгоритмов в Модуле 3.

Три шага модуля:

Логика как инженерия — освоим три базовых закона и научимся находить ошибки в любых текстах (даже написанных нейросетями).

Деревья вывода и сценарное моделирование — превратим неопределённость в набор вариантов с вероятностями.

Аргумент-клинок 2.0 — отточим искусство убеждения для цифровой эпохи.

В конце модуля вы соберёте все шаблоны в единый «Логический конструктор» — ваш персональный арсенал для анализа любых данных.

А теперь — в кузницу. Первый диалог: Юрий объясняет, почему формальная логика — лучший щит от работории.

Глава 4. Логика как инженерия: не академия, а кузница

● 4.1. Диалог: Юрий объясняет, почему формальная логика — лучший щит от рабонии

— Ты говоришь «логика», а у меня перед глазами — школьная доска, скучный учитель, формулы, — Ястребина поморщилась. — При чём здесь рабония?

Юрий рассмеялся.

— Вот именно. Ты только что продемонстрировала классическую подмену. Ты заменила содержание — «логика как инструмент защиты» — на образ — «скучный урок в школе». Это и есть та самая ловушка, которую алгоритмы эксплуатируют. Они подменяют суть эмоцией, ярлыком, картинкой.

— И что, формальная логика умеет такое распознавать?

— Не просто распознавать. Это её работа. Формальная логика — это наука о том, как отличить правильное рассуждение от неправильного. Независимо от того, кто его произнёс: человек, политик, рекламный слоган или самая совершенная нейросеть.

Почему логика работает против ИИ

Юрий достал смартфон, открыл диалог с AI-ассистентом.

— Смотри. Я пишу: «Все люди смертны. Сократ — человек. Значит, Сократ смертен». Это силлогизм. Он правильный, потому что вывод неизбежно следует из посылок. А теперь я пишу: «Многие мои друзья любят пиццу. Я — друг своих друзей. Значит, я люблю пиццу». Что здесь не так?

— Во втором утверждении подмена, — сказала Ястребина после паузы. — «Я — друг своих друзей» не означает, что я вхожу в множество «мои друзья» в первом утверждении. Там имелись в виду конкретные люди, а не абстрактное отношение.

— Блестяще! Ты только что применила закон тождества. А AI, когда я задал ему этот пример, сначала согласился, что вывод правильный. Потому что нейросети не проверяют логику — они подбирают вероятные последовательности слов. Они не знают, что «друг» в первом предложении и «друг» во втором — это не одно и то же.

Юрий отложил телефон.

— Вот почему логика — щит. Алгоритмы генерируют правдоподобное, но не обязательно истинное. Они подсовывают нам рассуждения, которые кажутся убедительными, но рассыпаются при формальной проверке. Если ты не владеешь логикой, ты проглотить любую чужую, упакованную в красивую упаковку. Если владеешь — ты видишь швы.

Три врага рабонии — три закона логики

— Нам понадобятся три базовых закона, — Юрий начал загибать пальцы. — Простых, как три удара молотом.

1. Закон тождества. «Понятие должно оставаться самим собой в пределах рассуждения». Если AI в начале диалога сказал «свобода» как одно, а в конце подменил на другое — ты это заметишь.

2. Закон непротиворечия. «Не могут быть одновременно истинны два противоположных утверждения». Если AI утверждает «рынок вырастет» и «рынок упадёт» без допол-

нительных условий — он противоречит себе. Ты не должен верить ни одному из них без проверки.

3. Закон исключённого третьего. *«Утверждение либо истинно, либо ложно. Третьего не дано».* AI любит уходить в «всё относительно» и «возможно, так, а возможно, иначе». Но для принятия решений нужна определённость. Закон напоминает: если нет доказательств истинности и ложности — утверждение не имеет ценности.

— *И ты хочешь сказать, что достаточно знать эти три закона, чтобы не стать рабони?* — спросила Ястребина.

— *Достаточно, чтобы начать. Но чтобы сделать их своей второй натурой — нужна практика. А практика — это следующие подглавы: как находить логические ошибки в текстах, сгенерированных ИИ, и как спорить с AI-оппонентом, который специально провоцирует тебя на ошибку.*

Задание на паузу

Юрий повернулся к читателю.

— *Прямо сейчас откройте любой диалог с AI-ассистентом за последние дни (или сгенерируйте новый). Найдите в его ответе утверждение, которое кажется вам сомнительным. Задайте себе вопрос: «Не нарушен ли здесь один из трёх законов? Тожество — не подменил ли он понятие? Непротиворечие — не противоречит ли он себе? Исключённое третье — не уходит ли он в „может быть“ вместо чёткого ответа?»*

Запишите своё наблюдение в «Формулярный» — в новый раздел «Логический конструктор», подраздел «Логические уязвимости AI».

— *А теперь — к делу,* — сказала Ястребина. — *Растолкуй эти три закона на пальцах. Чтобы даже ребёнок понял.*

— *Следующая подглава,* — кивнул Юрий. — *Три закона логики для повседневности.*

● 4.2. Три закона логики для повседневности (закон тождества, непротиворечия, исключённого третьего)

— Три закона, — сказал Юрий, усаживаясь поудобнее. — Их не нужно заучивать. Их нужно понять на примерах. Поехали.

Закон первый: тождества

— Звучит страшно, а смысл прост: понятие должно оставаться самим собой в пределах одного рассуждения. Нельзя в начале разговора назвать «свободой» одно, а в конце — другое.

— Приведи пример, — попросила Ястребина.

— Пожалуйста. Рекламный слоган: «Натуральный сок — это здоровье. Наш сок натуральный. Пейте — будете здоровы».

— И что здесь не так?

— А что такое «натуральный» в первом предложении? Свежевыжатый, без консервантов, из фермерских фруктов. А что такое «натуральный» во втором? Производитель может называть натуральным сок из концентрата, восстановленный, с добавлением сахара. Понятие подменили. Закон тождества нарушен.

Юрий достал телефон.

— Вот пример от AI. Я спросил: «Что такое справедливость?» Он ответил: «Справедливость — это когда каждый получает по заслугам». Потом я спросил: «А как измерить заслуги?» Он ответил: «Заслуги определяются обществом». Потом: «А если общество несправедливо?» Он ответил: «Тогда справедливость — это равенство возможностей».

— Он трижды подменил понятие, — заметила Ястребина.

— Именно. А я проглотил бы, если бы не знал закона.

Как применять в жизни: Когда слышите ключевое слово (свобода, справедливость, качество, натуральный, безопасный), спросите: «А что именно вы имеете в виду?» И следите, чтобы в ходе разговора значение не менялось.

Закон второй: непротиворечия

— Не могут быть одновременно истинны два противоположных утверждения. Либо А, либо не-А. Третьего не дано.

— Это же очевидно, — сказала Ястребина.

— Очевидно? Тогда объясни, почему AI может в одном абзаце написать «Цены на нефть вырастут из-за дефицита» и «Цены на нефть упадут из-за замедления экономики» — и это считается нормальным анализом?

— Потому что он перечисляет факторы, а не делает вывод.

— А читатель? Читатель видит два противоположных прогноза и думает: «Ну, эксперты не определились». А должен думать: «Это нарушение закона непротиворечия. Если оба утверждения истинны одновременно, то только при разных условиях. Но условия не указаны. Значит, информация не имеет ценности».

Как применять в жизни: Если вам говорят «риски высоки, но потенциал огромен» — это не противоречие. А если «проект абсолютно безопасен и при этом может провалиться в любой момент» — нарушение. Требуйте уточнений: при каких условиях истинно первое, при каких — второе?

Закон третий: исключённого третьего

— Любое утверждение либо истинно, либо ложно. Третьего не дано.

— А как же «не определились», «есть разные точки зрения», «всё относительно»? — спросила Ястребина.

— Это уход от ответа. Закон не запрещает говорить «я не знаю». Он запрещает утверждать, что есть какая-то третья истинность, средняя между да и нет.

— Приведи пример с AI.

— Легко. Я спросил: «Будет ли завтра дождь?» AI ответил: «Вероятность осадков составляет 30%, поэтому нельзя однозначно сказать, будет дождь или нет».

— Это же честный ответ!

— Честный, но бесполезный для принятия решения. Закон исключённого третьего не запрещает вероятности. Он запрещает утверждение, что существует некое третье состояние, кроме «дождь будет» и «дождя не будет». 30% — это не третье состояние. Это степень уверенности в первом. AI же часто маскирует незнание под «всё сложно», «неоднозначно», «зависит от контекста».

Как применять в жизни: Если вам отвечают «всё относительно», спросите: «Относительно чего? Назовите конкретные условия, при которых утверждение истинно, и при которых ложно». Если условий нет — перед вами пустышка.

Проверочная таблица (без таблицы)

Запомните три вопроса, которые задавайте любой информации:

Тожество: Одно и то же понятие используется в одном и том же смысле? Не подменили ли его?

Непротиворечие: Нет ли в рассуждении двух противоположных утверждений без указания условий?

Исключённое третье: Утверждение имеет чёткое значение (истина/ложь) или его размывают «относительностью»?

— А теперь, — сказала Ястребина, — проверим это на практике. Твоя очередь. Я зачитаю несколько фраз. Ты определишь, какой закон нарушен.

Юрий кивнул.

Упражнение «Найди нарушителя»

Фраза 1: «Это лекарство эффективно, потому что оно помогает. А помогает оно потому, что эффективно».

— Что здесь? — спросила Ястребина.

— Круг в доказательстве. Нарушен закон тождества: понятие «эффективно» объясняется через само себя.

Фраза 2: «С одной стороны, этот кандидат — опытный управленец. С другой стороны, у него нет опыта в нашей отрасли. Поэтому решение за вами».

— Здесь нет нарушения, — сказал Юрий. — Просто перечислены факты. А вот если бы сказали «он одновременно опытный и неопытный в одном и том же отношении» — было бы нарушение непротиворечия.

Фраза 3 (от AI): «Искусственный интеллект не может мыслить, потому что он не обладает сознанием. А сознание — это то, чего нет у машин. Однако некоторые учёные считают, что мышление возможно и без сознания. Вопрос остаётся открытым».

— *Это нарушение исключённого третьего?* — спросила Ястребина.

— *Нет, это подмена тезиса. Начали с «AI не может мыслить», закончили «вопрос открыт». AI ушёл от ответа, нарушив закон тождества — он говорил не об одном и том же.*

Ваше задание

Возьмите любой текст, сгенерированный нейросетью (новость, ответ чат-бота, рекламу). Проанализируйте его по трём законам. Найдите хотя бы одно нарушение. Запишите в «Формулярий» — в раздел «Логический конструктор», подраздел «Примеры логических ошибок».

— *Три закона — это фундамент,* — подвёл итог Юрий. — *Без них любой разговор с AI превращается в болото. С ними — вы становитесь инженером, который видит, где мост сгнил, а где стоит прочно.*

— *А следующий шаг — учиться распознавать логические ошибки в текстах, которые сгенерировал ИИ,* — добавила Ястребина. — *Это будет подглава 4.3.*

● 4.3. Распознавание логических ошибок в текстах, сгенерированных ИИ

— Три закона мы разобрали, — сказала Ястребина. — Теперь вопрос: как их применить к тому, что пишут нейросети? Они же не люди, у них нет намерения обманывать.

— В том и проблема, — ответил Юрий. — У них нет намерения говорить правду. У них есть намерение быть правдоподобными. А это разные вещи.

Он открыл ноутбук.

— Я собрал несколько примеров реальных ответов AI. Посмотрим, какие логические ошибки в них спрятаны.

Типичные логические ошибки генеративных нейросетей

Юрий вывел на экран список.

1. Галлюцинация фактов (нарушение закона тождества — выдача вымысла за реальность)

Пример: AI утверждает, что некое исследование доказало его правоту, но ссылка ведёт на несуществующую статью. Или приписывает известному человеку слова, которых он не говорил.

Как распознать: Требуйте точных ссылок. Проверяйте цитаты. Если AI не может предоставить источник — это галлюцинация.

2. Подмена тезиса (нарушение закона тождества — ответ не на заданный вопрос)

Пример: Вы спрашиваете: «Будет ли завтра дождь?» AI отвечает: «Климат меняется, погода становится непредсказуемой». Вопрос был о конкретном дне, ответ — о глобальном тренде.

Как распознать: Зафиксируйте исходный вопрос. Если ответ не отвечает на него прямо — перед вами подмена тезиса.

3. Круг в доказательстве (нарушение закона тождества — объяснение через себя же)

Пример: «Этот метод эффективен, потому что он приводит к хорошим результатам. А хорошие результаты — это то, что даёт этот метод».

Как распознать: Попробуйте вычленить причину и следствие. Если они совпадают — круг.

4. Ложная причинно-следственная связь (post hoc ergo propter hoc — после этого, значит, вследствие этого)

Пример: «После того как компания сменила логотип, продажи выросли. Значит, новый логотип — причина роста». А могло быть сезонное оживление, рекламная кампания или просто совпадение.

Как распознать: Спросите: «Есть ли другие возможные причины?» AI часто игнорирует альтернативы.

5. Ложная дилемма (нарушение закона исключённого третьего — предлагается только два варианта, хотя есть больше)

Пример: «Либо вы принимаете наши условия, либо теряете клиента». На самом деле есть третий вариант: пересмотреть условия, найти компромисс, попросить отсрочку.

Как распознать: Всегда ищите третий, четвёртый, пятый вариант.

6. Апелляция к авторитету (нарушение — мнение эксперта выдаётся за доказательство)

Пример: «Ведущие экономисты мира считают, что инфляция замедлится». Но не приводится ни расчётов, ни аргументов, только ссылка на «ведущих экономистов».

Как распознать: Спросите: «Кто именно? На чём основано их мнение? Есть ли контр-аргументы других экспертов?»

7. Скользящий склон (преувеличение последствий без доказательства)

Пример: «Если мы разрешим AI писать тексты, то скоро люди разучатся писать совсем, а потом и думать перестанут».

Как распознать: Требуйте промежуточные звенья цепочки. Без них это страх, а не логика.

Живой разбор: AI против логики

— *Приведи конкретный диалог,* — попросила Ястребина.

Юрий открыл чат с DeepSeek.

Юрий: «Стоит ли инвестировать в акции нефтяных компаний?»

DeepSeek: «Многие эксперты рекомендуют нефтяной сектор как защитный актив в период нестабильности. Например, известный инвестор Уоррен Баффет недавно увеличил долю в Chevron. Кроме того, цены на нефть традиционно растут в условиях геополитической напряжённости. Поэтому инвестиции в нефтяные акции выглядят разумно».

— *Разбираем по косточкам,* — сказал Юрий.

«**Многие эксперты рекомендуют**» — апелляция к неопределённому авторитету. Кто именно? Где их отчёты?

«**Уоррен Баффет увеличил долю**» — это факт. Но Баффет — долгосрочный инвестор, его стратегия может не подходить обычному человеку. Также он мог купить акции по другим причинам (дивиденды, хеджирование).

«**Цены на нефть растут в условиях напряжённости**» — правда, но не вся. Они также могут резко падать. И зависимость не линейная.

Вывод «инвестиции выглядят разумно» — не следует из приведённых аргументов как единственно возможный. Может быть и «инвестиции рискованны».

— *То есть AI не врёт, но создаёт иллюзию обоснованности,* — заметила Ястребина.

— *Именно. Он подбирает факты, подтверждающие тезис, и игнорирует контраргументы. Это нарушение закона непротиворечия — неявно, но факты подобраны однобоко. А мы, читатели, проглатываем.*

Тактика защиты: три шага

Юрий нарисовал в воздухе схему.

Шаг 1. Вычленив тезис. Чётко сформулировать, что именно утверждает AI. Не «инвестиции разумны», а «конкретному человеку в текущих условиях стоит вложить X денег в акции Y».

Шаг 2. Отделить факты от их интерпретации. Где в ответе AI проверяемые данные (цифры, даты, цитаты)? А где — оценки («выглядят разумно», «многие считают»)?

Шаг 3. Проверить логическую связь между фактами и выводом. Даже если все факты верны, ведёт ли вывод из них с необходимостью? Или есть скрытые допущения?

— *И ещё одно,* — добавил Юрий. — *Если AI даёт совет, всегда задавайте ему встречный вопрос: «А какие аргументы против?» Большинство нейросетей с готовностью приведут контраргументы, если их попросить. Но сами они их не предлагают.*

Практическое задание

Выберите любой текст, сгенерированный нейросетью (можно взять ответ AI на ваш запрос, новость, сгенерированную AI, или рекламный пост).

Проанализируйте его на наличие логических ошибок из списка выше. Найдите хотя бы одну.

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.