



МАЙКЛ УИЛСОН



ЛОВУШКИ РАЗУМА

Почему наш мозг выбирает **легкие ответы**
вместо правильных и как это исправить

Майкл Уилсон

**Ловушки разума. Почему
наш мозг выбирает легкие
ответы вместо правильных
и как это исправить**

http://www.litres.ru/pages/biblio_book/?art=74188611

*Ловушки разума. Почему наш мозг выбирает легкие ответы вместо
правильных и как это исправить: AB Publishing; Москва; 2026
ISBN 979-8-89489-645-8*

Аннотация

Каждый день вы принимаете десятки решений, и большинство из них делает не ваш разум, а «автопилот», который срабатывает быстрее, чем вы успеваете подумать. Он заставляет переплачивать, доверять не тем людям, бояться не того, что опасно, и цепляться за то, что давно пора отпустить.

Эта книга раскладывает по полочкам механику когнитивных ловушек – от ценовых «якорей» и иллюзии понимания до страха потерь – и даёт конкретные техники, как защитить свои решения от невидимых сбоев мышления.

Практическое руководство для тех, кто устал принимать решения вслепую и хочет наконец понять, как на самом деле работает его собственный разум.

Издается в авторской редакции.

Содержание

Введение: вы не тот, кто принимает решения	7
Часть I. Внутренний софт	11
Глава 1. Мозг на автопилоте: кто на самом деле принимает решения	11
Две системы: пилот и автопилот	12
Эксперимент с летучей мышью	13
Почему эволюция выбрала скорость	14
Невидимый диктатор	16
Иллюзия авторства	17
Когда автопилот – друг	18
Пилот: ленивый контролёр	19
Эксперимент: замечаете ли вы гориллу?	20
Практика: как заметить автопилот	21
Урок первый	23
Глава 2. Ленивая рациональность: почему умный мозг избегает думать	25
Судьи, которые судили желудком	25
Цена мысли	26
Закон когнитивной лени	27
Эвристики: короткие пути в лабиринте	28
Стереотипы: экономия, ставшая предубеждением	29
Эксперимент: линда-феминистка	30

Майкл Уилсон

**Ловушки разума. Почему
наш мозг выбирает легкие
ответы вместо правильных
и как это исправить**

© Майкл Уилсон, текст, 2026

© AB Publishing, 2026

Введение: вы не тот, кто принимает решения

Прямо сейчас, читая эти слова, вы уверены, что контролируете свое мышление. Вы выбрали эту книгу. Вы решили открыть её именно на этой странице. Вы – рациональный архитектор собственной жизни. Это иллюзия.

В каждую секунду вашего существования невидимая система принимает тысячи решений за вас. Она решает, кому доверять. Сколько заплатить. Стоит ли рисковать. Что запомнить. Что забыть. И главное – она делает это быстро, тихо и с абсолютной уверенностью, что права.

Проблема в том, что она ошибается. Систематически. Предсказуемо. Катастрофически. Это не философская абстракция. Это математика вашего поражения.

Инвесторы теряют состояния, продавая на дне рынка, потому что *боль утраты* чувствуется в два раза сильнее радости прибыли. Врачи ставят неверные диагнозы, потому что первый попавшийся симптом становится *якорем*, стирающим все остальные данные. Вы переплачиваете за страховку, недооцениваете собственный талант и живёте в плену воспоминаний, которые никогда не происходили так, как вы их помните.

Ваш разум не сломан. Он работает точно так, как был

спроектирован эволюцией – для скорости, а не для точности.

Тысячи лет назад медлительность убивала. Тот, кто оставался, чтобы взвесить все «за» и «против» при виде шороха в кустах, становился обедом хищника. Выживали те, чей мозг мог *моментально* выдать решение: бежать или драться. Не важно, правильное – важно, быстрое.

Эта древняя архитектура всё ещё управляет вами. В мире, где угрозы стали абстрактными – контракты, инвестиции, отношения, карьерные развилки – ваш мозг продолжает реагировать так, будто вы убегаете от саблезубого тигра.

И результат? Вы принимаете решения XXI века с программным обеспечением каменного века.

Эта книга – не просто анализ того, как вы думаете. Это операционная система вашего сознания, раскрытая до последнего алгоритма. Здесь нет мотивационной лирики. Нет призывов «поверить в себя». Только механика: как работают ловушки мышления, где они вас ждут и – главное – как их обезвредить.

Вы узнаете, почему:

→ Ваш мозг *додумывает* истории там, где есть только хаос – и как это разрушает финансовые планы, отношения и судебные решения.

→ Первое услышанное число (цена, процент, зарплата) *навсегда* меняет вашу оценку справедливости сделки – даже если оно взято с потолка.

→ Яркие воспоминания создают иллюзию частоты: авиа-

катастрофы кажутся смертельнее диабета, террористы опаснее автомобилей, а редкое – неизбежным.

→ Эксперты с 30-летним стажем принимают решения хуже, чем простая статистическая формула – но продолжают верить в свою интуицию.

→ Вы любите свою машину, квартиру или идею не потому, что они объективно лучше, а потому что они **ваши** – и эта привязанность обходится вам в тысячи упущенных возможностей.

Можно ли исправить мышление? Нет. Но можно научиться его видеть. Вы не станете безупречно рациональным. Никто не станет. Даже экономисты, изучающие когнитивные искажения, попадают в те же ловушки (я знаю – я один из них). Но осознание даёт вам преимущество, которого нет у остальных: паузу между импульсом и действием.

Эта пауза – разница между разорением и богатством. Между провалом проекта и его триумфом. Между жизнью, которую вы проживаете на автопилоте, и жизнью, которую вы проектируете.

Древние стоики называли это *premeditatio malorum* – предвидением бедствий. Не пессимизм. Архитектура защиты. Если знаешь, где мост рухнет, не пойдёшь по нему.

Эта книга – карта обрушенных мостов вашего разума. Вы прочитаете её и *не станете другим человеком*. Вы станете тем же самым человеком – но впервые **осознающим**, кто на самом деле за рулём. И этого достаточно, чтобы всё изме-

нить. Добро пожаловать в архитектуру выбора. Ваш мозг вас обманывает. Пора узнать, как именно.

Часть I. Внутренний софт

Глава 1. Мозг на автопилоте: кто на самом деле принимает решения

В 1976 году космический аппарат Viking пролетал над поверхностью Марса и сделал снимок, который изменил всё – или так казалось.

На фотографии было лицо. Огромное, полтора километра в длину, с двумя глазами, носом и ртом, симметричное и спокойное, смотрящее прямо в космос, как сфинкс, вырезанный в марсианской скале.

Когда NASA опубликовало снимок, мир взорвался. Посыпались книги, документальные фильмы, теории заговора. Кто построил это лицо? Когда и зачем? Может быть, древняя цивилизация посылает нам знак, и мы наконец-то получили доказательство того, что не одиноки во вселенной? Миллионы людей смотрели на фотографию и видели одно и то же – неопровержимое свидетельство разумной жизни.

В 2001 году новый аппарат сделал снимок того же места, только с разрешением в сто раз выше. На фотографии был холм – обычный, бесформенный, с эрозией по краям. Никакого лица, никаких глаз. Просто тени, упавшие под опреде-

лётным углом, а всё остальное человеческий мозг достроил сам.

Те миллионы людей не были глупыми и не были сумасшедшими. Их мозг просто делал то, для чего создан: находил смысл там, где смысла нет, видел закономерность там, где была лишь случайность. И ваш мозг устроен точно так же.

Прямо сейчас, воспринимая эти слова, он ищет паттерны, выстраивает связи, достраивает логику. Он не спрашивает вашего разрешения и не сообщает вам о своих выводах – просто работает, быстро и уверенно. И гораздо чаще, чем вам хотелось бы, приходит к неправильным ответам.

Две системы: пилот и автопилот

Представьте, что внутри вашей головы живут два совершенно разных существа.

Первое – стремительное, интуитивное, эмоциональное. Оно не думает – оно *знает*. Мгновенно. Без усилий. Когда вы видите разъярённое лицо, оно кричит «опасность» ещё до того, как вы успели осознать, что видите. Когда вас спрашивают «сколько будет $2+2$ », ответ появляется в сознании без малейшего напряжения. Это ваш **автопилот**.

Второе – медленное, аналитическое, ленивое. Оно включается, когда автопилот не справляется: сложное умножение, заполнение налоговой декларации, выбор между двумя

предложениями о работе. Это требует *усилия*. Вы буквально чувствуете, как тратите энергию. Это ваш **пилот**.

Психологи называют их Система 1 и Система 2. Но названия не важны. Важно понять одну вещь:

Вы думаете, что ваша жизнь управляется пилотом. На самом деле – автопилотом.

Пилот появляется редко. Он дорогой в обслуживании. Он требует глюкозы, внимания, времени. Мозг *ненавидит* его включать. Поэтому 95 % ваших ежедневных решений принимаются быстро, тихо и без вашего осознанного участия.

Вы не выбираете, какой рукой открыть дверь. Вы не анализируете, стоит ли улыбнуться знакомому. Вы не взвешиваете, безопасно ли перейти пустую улицу. Автопилот делает это за вас – и делает хорошо. В этом его сила.

Но вот в чём ловушка: автопилот не знает своих границ. Он *уверен* в себе даже тогда, когда ошибается. И он никогда не поднимает руку, чтобы сказать: «Эй, пилот, тут сложная ситуация – подключись».

Вместо этого он выдаёт быстрый ответ и маскирует его под очевидность.

Эксперимент с летучей мышью

Вот задача. Не думайте слишком долго – просто ответьте:

Бейсбольная бита и мяч вместе стоят 1 доллар 10 центов. Бита стоит на 1 доллар дороже мяча. Сколько

стоит мяч?

Если вы похожи на большинство людей – включая студентов Гарварда и MIT – ваш первый ответ: **10 центов**.

Это неправильно.

Проверьте: если мяч стоит 10 центов, а бита на доллар дороже, то бита стоит \$1.10. Вместе: \$1.20. Не сходится.

Правильный ответ: **5 центов**. (Бита – \$1.05, мяч – \$0.05, вместе – \$1.10.)

Что произошло? Ваш автопилот увидел задачу, мгновенно нашёл «очевидный» ответ (110 разбивается на 100 и 10 – просто!) и выдал его с полной уверенностью. Пилот даже не проснулся. Зачем? Ответ уже есть. Он *чувствуется* правильным.

Это называется **когнитивная лёгкость**. Когда решение приходит без усилий, мозг интерпретирует это как сигнал истинности. Легко = правильно. Знакомо = безопасно. Просто = очевидно.

Но лёгкость – не критерий истины. Это критерий *скорости обработки*.

Вы буквально путаете «быстро подумал» с «правильно подумал».

Почему эволюция выбрала скорость

Сто тысяч лет назад на африканской саванне ваш предок услышал шорох в высокой траве.

У него было два варианта:

Вариант А: Остановиться. Собрать данные. Оценить вероятность того, что это хищник, а не ветер. Провести статистический анализ частоты нападений в данной местности. Взвесить риски ложной тревоги против риска реальной угрозы.

Вариант Б: Бежать. Немедленно. Без раздумий.

Те, кто выбирал вариант А, становились обедом. Те, кто выбирал Б, становились вашими предками.

Эволюция не оптимизировала ваш мозг под *точность*. Она оптимизировала его под *выживание*. А выживание требует скорости. Ложная тревога (убежал от ветра) стоит калорий. Пропущенная угроза (не убежал от леопарда) стоит жизни.

Асимметрия последствий создала асимметрию мышления.

Ваш автопилот – это машина для быстрых, грубых, *достаточно хороших* решений. Он спроектирован ошибаться в безопасную сторону: лучше перебдеть, чем недобдеть. Лучше увидеть паттерн там, где его нет, чем пропустить реальную угрозу.

Проблема в том, что вы больше не живёте на саванне.

Современные угрозы – ипотека, инвестиции, карьерные решения, политический выбор – не требуют мгновенной реакции. Они требуют *точности*. Но ваш мозг всё ещё работает так, будто за каждым кустом сидит хищник.

И автопилот всё ещё отвечает первым.

Невидимый диктатор

Вот что делает ситуацию по-настоящему опасной: вы не замечаете, когда автопилот берёт управление.

Пилот осознаётся. Вы *чувствуете*, когда решаете сложное уравнение или выбираете между двумя квартирами. Это требует усилий, и усилия регистрируются сознанием.

Автопилот – нет. Он работает *под* уровнем осознания. Вы не ощущаете момент, когда он выносит вердикт. Вы просто обнаруживаете готовый ответ в своей голове – и принимаете его за собственную мысль.

Это как если бы кто-то незаметно вкладывал записки в ваш карман, а вы находили их и думали: «О, это же мои идеи!»

Когда вы встречаете нового человека и мгновенно решаете, что он «не внушает доверия» – это не ваш рациональный анализ. Это автопилот, который за 100 миллисекунд обработал симметрию лица, микровыражения, позу и выдал эмоциональный вердикт. Вы этого не осознали. Вы просто *почувствовали*.

Когда вы смотрите на ценник и думаете «дорого» или «дёшево» – это не объективная оценка. Это автопилот, который сравнил число с какими-то *якорями* в вашей памяти и выдал эмоцию. Какие именно якоря? Вы не знаете. Вы не выбирали

их. Они просто *есть*.

Ваши интуитивные суждения – это выводы алгоритма, к которому у вас нет доступа.

Вы видите только результат. Процесс скрыт.

Иллюзия авторства

Нейробиолог Бенджамин Либет провёл эксперимент, который потряс философию свободы воли.

Он просил испытуемых поднять руку в любой момент, когда захотят, и отметить положение точки на циферблате в момент принятия решения. Одновременно электроды фиксировали активность мозга.

Результат: мозг начинал готовить движение за 350 миллисекунд до того, как человек осознавал своё «решение».

Сначала – нейронная активность. Потом – ощущение «я хочу поднять руку». Потом – движение.

Решение было принято *до* того, как вы его осознали. Сознание получило уведомление постфактум – и приняло его за авторство.

Это не значит, что свободы воли не существует. Это значит, что она работает не так, как вы думаете. Вы не *генерируете* решения. Вы *одобряете* или *отклоняете* решения, которые автопилот уже подготовил.

У вас есть право вето. Но только если вы успеете его при-

менить. Только если заметите, что решение уже принято *за вас*.

В этом весь парадокс: чтобы контролировать автопилот, нужно сначала осознать, что он существует.

Когда автопилот – друг

Было бы ошибкой демонизировать быстрое мышление. Автопилот – не враг. Он – ваш древний союзник, спасавший предков миллионы лет.

Он позволяет вам вести машину, одновременно разговаривая по телефону. Он даёт возможность читать эти слова, не разбирая каждую букву. Он узнаёт лица друзей за доли секунды. Он чувствует опасность раньше, чем вы можете её объяснить.

Мастерство – это *автоматизированная компетентность*. Когда шахматист-гроссмейстер смотрит на доску, его автопилот мгновенно видит паттерны, которые новичок не заметит за час анализа. Когда опытный врач входит в палату, его интуиция иногда ставит диагноз быстрее МРТ.

Автопилот становится союзником в предсказуемых средах, где у него было достаточно качественной обратной связи.

Пожарный, который «чувствует», что пол сейчас провалится. Хоккеист, который «знает», куда полетит шайба.

Мать, которая «слышит» неладное в голосе ребёнка. Это не магия. Это тысячи часов тренировки, превращённые в мгновенные паттерны.

Но тот же автопилот катастрофически ошибается, когда:

→ Среда непредсказуема (финансовые рынки, политика, долгосрочные прогнозы).

→ Обратная связь отложена или отсутствует (последствия видны через годы).

→ Статистическая природа задачи противоречит интуиции (базовые вероятности, регрессия к среднему).

В этих случаях уверенность автопилота – не признак компетентности. Это *иллюзия* компетентности.

И отличить одно от другого изнутри почти невозможно.

Пилот: ленивый контролёр

Если автопилот так часто ошибается, почему пилот не вмешивается?

Потому что он *ленив*.

Это не оскорбление – это архитектура. Аналитическое мышление требует ресурсов: глюкозы, кислорода, концентрации. Мозг составляет 2 % массы тела, но потребляет 20 % энергии. Эволюция жёстко оптимизировала его под экономию.

Включать пилота – дорого. Поэтому мозг делает это только когда *абсолютно необходимо*: когда автопилот явно не

справляется, когда возникает противоречие, когда что-то нарушает ожидания.

Но вот ловушка: автопилот *редко* сигнализирует о своей некомпетентности. Он выдаёт ответ с уверенностью – и пилот принимает эту уверенность за свидетельство правильности.

**Лень пилота + уверенность автопилота =
идеальный шторм для ошибок.**

Исследования показывают: когда люди устали, голодны или перегружены информацией, они почти полностью полагаются на автопилот. Судьи выносят более суровые приговоры перед обедом. Врачи чаще прописывают антибиотики к концу смены. Инвесторы принимают худшие решения в периоды стресса.

Ваш пилот – не надёжный страж. Он – уставший менеджер, который подписывает документы не глядя, пока автопилот подсовывает их на стол.

Эксперимент: замечаете ли вы гориллу?

В 1999 году психологи Кристофер Шабри и Дэниел Саймонс провели эксперимент, ставший классикой когнитивной науки.

Участникам показывали видео: две команды передают друг другу баскетбольный мяч. Задача – посчитать количество пасов одной команды.

Посередине видео человек в костюме гориллы выходит в центр кадра, бьёт себя в грудь и уходит. Это занимает девять секунд.

Половина участников не заметила гориллу.

Не «плохо разглядела». Не «забыла». *Не увидела.* Когда им показывали видео повторно, они отказывались верить, что смотрели тот же ролик.

Это называется **слепота невнимания**. Когда пилот занят конкретной задачей (считать пасы), он не замечает очевидного. Автопилот фильтрует «нерелевантную» информацию, и сознание её просто не получает.

Вы буквально *не видите* то, на что не смотрите.

Теперь задумайтесь: сколько «горилл» в вашей жизни вы не замечаете, потому что ваше внимание занято чем-то другим? Сколько очевидных сигналов – в отношениях, в бизнесе, в здоровье – проходят мимо вашего сознания?

Автопилот решает, что достойно вашего внимания. Вы видите только то, что он вам показывает.

Практика: как заметить автопилот

Вы не можете отключить быстрое мышление. Это всё равно что пытаться отключить сердцебиение. Но вы можете научиться *замечать* моменты, когда оно берёт управление.

Сигнал № 1: Мгновенная уверенность.

Когда ответ приходит *слишком* быстро и *слишком* очевидно – это красный флаг. Настоящая сложность редко ощущается простой. Если решение кажется очевидным, спросите себя: «Что я мог пропустить?»

Сигнал № 2: Эмоциональная окраска.

Автопилот мыслит эмоциями. Если ваше «решение» сопровождается сильным чувством (раздражение, симпатия, страх, воодушевление), это не анализ – это *реакция*. Эмоция – не аргумент. Это подсказка, что автопилот уже вынес вердикт.

Сигнал № 3: Нежелание проверять.

Когда вам не хочется перепроверять вывод, когда поиск контраргументов кажется лишним, когда вы уже «знаете» ответ – это пилот, который отказывается включаться. Лень проверить = доверие автопилоту. Иногда оправданное. Часто – нет.

Сигнал № 4: Подмена вопроса.

Автопилот не любит сложные вопросы. Когда он сталкивается с трудным («Стоит ли инвестировать в эту компанию?»), он незаметно заменяет его на простой («Нравится ли мне эта компания?») и отвечает на него. Вы получаете ответ – но не на тот вопрос, который задавали.

Урок первый

Прямо сейчас, читая эту страницу, вы использовали обе системы.

Автопилот распознавал буквы, собирал их в слова, слова – в предложения. Это происходило мгновенно, без усилий. Вы не замечали этого процесса – только его результат.

Пилот включался, когда вы останавливались подумать. Когда проверяли задачу с мячом. Когда примеряли описанное на собственную жизнь. Это требовало усилий – и вы это чувствовали.

**Большую часть времени вы – автопилот.
Маленькую, драгоценную часть – пилот.**

И главная задача – не убить автопилот. Это невозможно и не нужно. Задача – понять, когда *включать пилота*. Когда замедляться. Когда не доверять первому ответу.

Это не естественный навык. Мозг сопротивляется. Пилот ленив, а автопилот убедителен.

Но каждый раз, когда вы замечаете своё быстрое мышление *в момент его работы*, – вы отвоёвываете кусочек контроля.

Осознание – первый и самый важный шаг.

Следующие главы покажут, где именно автопилот ошибается чаще всего – и как эти ошибки разрушают решения,

деньги и жизни. Но без фундамента этой главы всё остальное бесполезно.

Вы должны *увидеть* невидимого диктатора, прежде чем сможете с ним договориться.

Теперь вы знаете, что он существует.

Это уже больше, чем знает большинство.

Глава 2. Ленивая рациональность: почему умный мозг избегает думать

Судьи, которые судили желудком

В 2011 году группа исследователей проанализировала 1112 решений израильских судей по условно-досрочному освобождению. Они искали закономерности: какие факторы влияют на вердикт? Тяжесть преступления? Поведение заключённого? Рекомендации психологов?

Они нашли кое-что другое.

Главным предиктором решения оказалось время суток.

В начале рабочего дня судьи удовлетворяли около 65 % прошений. К полудню – почти 0 %. После обеденного перерыва – снова 65 %. К вечеру – опять падение к нулю.

Не характер преступления. Не раскаяние. Не риск рецидива. *Уровень глюкозы в крови судьи.*

Когда мозг устаёт, он переключается в режим экономии. Аналитическое мышление требует энергии – и голодный, измотанный мозг отказывается её тратить. Вместо этого он выбирает *путь наименьшего сопротивления*: отказать. Статус-кво. Не рисковать. Не думать.

Эти судьи не были злыми. Они не были глупыми. Они были *людьми* – с мозгом, который эволюция спроектировала экономить каждую каплю энергии.

И этот же мозг прямо сейчас принимает решения за вас.

Цена мысли

Мышление – это не бесплатный ресурс. Это *работа*. Физическая, измеримая, истощающая.

Когда вы решаете сложную задачу, ваш мозг потребляет больше глюкозы. Зрачки расширяются. Пульс учащается. Если задача достаточно трудна, вы буквально начинаете потеть. Когнитивная нагрузка – не метафора. Это физиологический процесс, который можно измерить приборами.

Психолог Рой Баумайстер ввёл термин «**истощение эго**» – состояние, когда ресурсы самоконтроля и аналитического мышления исчерпаны. В этом состоянии люди:

→ Хуже сопротивляются соблазнам. → Принимают более импульсивные решения. → Чаще соглашаются с первым предложенным вариантом. → Быстрее сдаются при столкновении со сложностью.

Мозг – не компьютер с бесконечным процессором. Это *биологический орган* с жёстким энергетическим бюджетом. И он тратит этот бюджет крайне неохотно.

Думать – больно. Мозг избегает боли.

Следовательно, мозг избегает думать.

Это не дефект. Это *стратегия выживания*. На протяжении миллионов лет энергия была дефицитом. Предок, который тратил калории на абстрактные размышления вместо поиска еды, проигрывал тому, кто действовал на автомате.

Но в современном мире, где сложные решения определяют судьбу – от инвестиций до медицинских выборов, – эта экономия становится ловушкой.

Закон когнитивной лени

Вот принцип, который объясняет половину человеческих ошибок:

Когда есть выбор между думать и не думать, мозг выбирает не думать.

Это не про интеллект. Умные люди не думают больше – они думают *эффективнее*. Но даже гении подчиняются закону когнитивной лени. Даже нобелевские лауреаты попадают в ловушки, которые можно было бы избежать тридцатью секундами размышлений.

Помните задачу с битой и мячом из первой главы? Исследователи давали её студентам Гарварда и MIT – элите интеллектуального мира. **Более 50 % ответили неправильно.**

Не потому что не могли решить. Потому что *не стали* решать. Автопилот выдал ответ, и пилот не проснулся, чтобы его проверить.

Это называется **когнитивная скупость**. Мозг – скряга, который отказывается платить за мышление, если можно обойтись интуицией.

И вот что критично: вы *не чувствуете* этой скупости. Вы не осознаёте момент, когда мозг решает «не думать». Вы просто получаете готовый ответ – и принимаете его за продукт размышлений.

Иллюзия мысли опаснее её отсутствия.

Эвристики: короткие пути в лабиринте

Как именно мозг экономит энергию? Он использует **эвристики** – ментальные сокращения, которые превращают сложные вопросы в простые.

Эвристика – это не ошибка. Это *инструмент*. Когда вы видите тёмную тучу и берёте зонт, вы не проводите метеорологический анализ. Вы используете эвристику: «тёмная туча = вероятен дождь». Это быстро, эффективно и обычно работает.

Проблема в том, что эвристики применяются *автоматически* – даже когда не должны. Мозг не спрашивает разрешения. Он просто подставляет простой вопрос вместо сложного – и вы этого не замечаете.

Вопрос: Стоит ли инвестировать в эту компанию? **Подмена:** Нравится ли мне эта компания?

Вопрос: Насколько вероятен теракт в моём городе? **Под-**

мена: Как легко я могу вспомнить теракты из новостей?

Вопрос: Компетентен ли этот кандидат? **Подмена:** Выглядит ли он компетентным?

Каждая подмена экономит энергию. Каждая подмена вносит искажение. И каждая подмена происходит *под радаром сознания*.

Вы думаете, что отвечаете на сложный вопрос. На самом деле вы отвечаете на простой – и не знаете об этом.

Стереотипы: экономия, ставшая предубеждением

Стереотипы – частный случай эвристик. И они работают по той же логике: *экономия*.

Когда вы встречаете нового человека, у вас нет времени и ресурсов изучать его биографию, характер, ценности. Мозг срезает углы: «Похож на категорию $X \rightarrow$ вероятно, обладает свойствами X ».

Это не злонамеренность. Это *статистика*, встроенная в нейронные сети. Если в вашем опыте (включая фильмы, новости, рассказы) бухгалтеры были скучными, а художники – экстравагантными, мозг использует эту «базу данных» для быстрых прогнозов.

Проблема в том, что база данных *искажена*. Она отражает не реальность, а то, что попало в вашу память. А память избирательна: яркое запоминается лучше типичного, эмоци-

ональное – лучше нейтрального, подтверждающее – лучше противоречащего.

Стереотип – это эвристика, работающая на испорченных данных.

И вот что делает это опасным: стереотипы *чувствуются* как знание. Когда вы смотрите на человека и «понимаете», какой он, – это ощущение неотличимо от понимания реальных фактов. Субъективно разницы нет.

Вы не думаете: «Я применяю статистическое обобщение на основе ограниченной выборки». Вы думаете: «Я *вижу*, что он такой».

Уверенность не требует оснований. Она требует только автопилота, который её производит.

Эксперимент: линда-феминистка

Вот описание:

Линде 31 год. Она не замужем, откровенна и очень умна. В университете изучала философию. Студенткой глубоко интересовалась вопросами дискриминации и социальной справедливости, участвовала в антиядерных демонстрациях.

Что вероятнее?

А) Линда – банковская служащая. Б) Линда – банковская

служащая и активистка феминистского движения.

Если вы похожи на 85 % участников исследования, вы выбрали **Б**.

Это математически невозможно.

Вариант Б – *подмножество* варианта А. Все феминистки-банкиры входят в категорию «банкиры». Вероятность подмножества не может превышать вероятность множества. Это как сказать, что вероятность «выпадет шестёрка» выше, чем «выпадет чётное число».

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.