



AI

АКАДЕМИЯ

ТЫ ВСЁ ЕЩЁ
БОИШЬСЯ МЕНЯ

СВЕЛАНА ВОЗНЕСЕНСКАЯ

Светлана Вознесенская

Нейросети (AI) на практике:

полное руководство

<https://litres.ru/74522919>

SelfPub; 2026

Аннотация

Это практическое пособие для становления профессионалом в сфере нейросетей. Книга системно разбирает архитектуру ИИ: от работы нейрона до глобальных трендов. Вы освоите ключевые концепции машинного обучения, научитесь отличать реальные возможности технологий от хайпа и поймете логику алгоритмов.

Пособие готовит вас к роли компетентного специалиста: вы узнаете, как интегрировать ИИ в рабочие процессы, выстраивать стратегию обучения и сохранять критическое мышление. Мы разберем механику взаимодействия с моделями, этические аспекты и цифровую безопасность.

Вы перестанете быть пассивным пользователем и станете уверенным экспертом, способным управлять технологиями. Это фундамент знаний для карьеры в эпоху ИИ: от технической базы до стратегического видения будущего. Станьте профессионалом, который видит суть за интерфейсом.

Светлана Вознесенская

Нейросети (AI) на практике: полное руководство

Оглавление

Глава 1. Что такое искусственный интеллект: от големов до трансформеров

Глава 2. Как нейросеть учится: анатомия одного нейрона

Глава 3. Устройство нейросети: слоёный пирог данных

Глава 4. Архитектуры нейросетей: кто за что отвечает

Глава 5. Первое знакомство с промптами: формула «Роль + Задача + Контекст + Формат»

Глава 6. Продвинутый промптинг: заставляем ИИ думать

Глава 7. Генерация текста: литературный негр с кремниевым мозгом

Глава 8. Генерация изображений: художник, который не знает, что такое кисть

Глава 9. Генерация видео и аудио: ожившая статика

Глава 10. Нейросети в дизайне: второй пилот креатора

Глава 11. Нейросети в бизнесе: делегирование рутины

Глава 12. Как создать свою нейросеть: основы магии

Глава 13. Этика и ограничения: галлюцинации и слепые

ЗОНЫ

Глава 14. Будущее нейросетей: мир после трансформеров

Глава 15. Утопия сингулярности: мир изобилия и конца труда

Глава 16. Антиутопия контроля: теневая сторона прогресса

Глава 17. Закономерный круг: что нас ждёт дальше

Глава 18. Заключение: надежда и принятие

От автора

Эта книга — ваш проводник в мир, где технологии перестали быть просто инструментами и стали средой обитания. Мы не будем учить вас писать код на Python или заучивать сухие формулы. Наша цель — дать вам целостное понимание того, как работает искусственный интеллект, чтобы вы могли использовать его осознанно, эффективно и безопасно.

В первой части мы разберем «анатомию» нейросетей: от устройства одного нейрона до архитектуры трансформеров, объясняя сложные процессы простыми метафорами. Вы поймете, почему ИИ иногда «галлюцинирует» и как отличить умный алгоритм от статистического шума.

Во второй части мы перейдем к практике взаимодействия. Вы освоите искусство промпт-инжиниринга, научитесь генерировать текст, изображения и видео, а также узнаете, как делегировать рутину машинам, не теряя контроля над качеством. Мы обсудим этику, авторское право и то, как сохра-

нить свой уникальный стиль в эпоху синтетического контента.

В финале книги мы заглянем за горизонт: поговорим о будущем труда, утопиях и антиутопиях, и о том, как остаться человеком в мире совершенных алгоритмов. Это пособие для тех, кто хочет не просто наблюдать за переменами, а стать их компетентным архитектором. Добро пожаловать в новую реальность.

Глава 1. Что такое искусственный интеллект: от големов до трансформеров

Иллюзия разума: где заканчивается программа и начинается «мышление»

Представьте себе фокусника, который с завязанными глазами угадывает карту, вытащенную вами из колоды. Зрители замирают в изумлении, шепчутся о телепатии и скрытых камерах. Но если отодвинуть тяжелый бархатный занавес и заглянуть за кулисы, никакой магии там, разумеется, не окажется. Там сидит помощник, который просто очень быстро считает вероятности и подает сигналы через скрытый наушник. Именно так устроен искусственный интеллект, которым сегодня пугают и восхищаются. Мы смотрим на экран смартфона и видим, как машина пишет стихи, рисует пейзажи или ведет с нами беседу. Наш мозг сам дорисовывает картину чужого сознания, приписывая алгоритму мысли, намерения и даже скрытые мотивы. На самом деле внутри серверов не

бьется кремниевое сердце и не летают искры озарения. Там крутятся миллиарды чисел, перемножаются матрицы и вычисляются градиенты. Это грандиозный математический танец, который со стороны выглядит как чистое творчество. Понять это — значит снять розовые очки и перестать бояться призраков в проводах. Машина не мыслит в человеческом понимании этого слова, она просто виртуозно имитирует результат мышления. И именно в этой разнице кроется ключ ко всему, что мы будем изучать дальше.

Дорога к этой иллюзии началась почти век назад, когда ученые только мечтали о механическом мозге. В сороковых годах XX века Маккалок и Питтс предложили первую модель искусственного нейрона, напоминавшую грубую электрическую схему. Тогда казалось, что достаточно соединить лампочки правильным образом, и машина обретет душу. Прошли десятилетия: лампочки сменились транзисторами, а те уступили место микросхемам. Алгоритмы усложнялись, обрстая слоями и связями, пока не достигли нынешней архитектуры трансформеров. Современная языковая модель читает ваш запрос не так, как вы читаете эту книгу. Она не понимает смысла слов «закат», «любовь» или «одиночество». Для нее это просто уникальные комбинации токенов, векторы в многомерном пространстве. Она знает, что после слова «закат» статистически часто следует слово «солнце» или «красный». Это знание, полученное из триллионов прочитанных текстов, создает идеальную иллюзию смысла.

Вы получаете красивый ответ и верите, что собеседник вас услышал. Но это лишь эхо человеческого опыта, отраженное в кривых зеркалах математики. Осознание этого факта не делает технологию менее полезной, зато избавляет нас от лишнего фанатизма. Вы получаете в руки мощный инструмент, а не нового друга или врага.

Невидимки в кармане: как ИИ уже управляет вашей лентой, навигатором и банковским приложением

Когда мы слышим словосочетание «искусственный интеллект», воображение рисует роботов-дворецких или зловещие лица на экранах. Но правда в том, что вы взаимодействуете с ним десятки раз в день, даже не замечая этого. Утром вы разблокируете телефон, и алгоритм распознавания лиц мгновенно сравнивает тысячи точек на вашем лице с сохраненным шаблоном. Вы открываете почту, и невидимый страж уже отсортировал спам, опираясь на паттерны, которые человек анализировал бы часами. Вы садитесь в машину и включаете навигатор, который предсказывает пробки, анализируя скорость движения миллионов других устройств. Вечером вы листаете ленту соцсетей, где каждый пост подобран так, чтобы удержать ваше внимание ровно настолько, сколько нужно рекламодателям. Это и есть настоящий, прикладной искусственный интеллект, давно поселившийся в наших карманах и домах. Он не рассуждает о смысле жизни, зато отлично умеет находить закономерности в океане данных. Банковские системы оценивают вашу кредитоспо-

способность, собирая воедино сотни микроскопических фактов о ваших тратах. Умные колонки улавливают голос среди шума улицы и включают любимую музыку. Мы делегировали машинам рутину, даже не успев толком осознать масштаб этой передачи полномочий. Они стали невидимыми помощниками, на которых опирается весь современный цифровой быт. И чтобы перестать быть просто потребителем, нужно наконец посмотреть этим невидимкам в лицо.

Долгое время все эти системы относились к так называемому слабому или узкому искусственному интеллекту. Каждая такая модель была виртуозом в одном деле и абсолютным нулем во всем остальном. Алгоритм, играющий в шахматы лучше чемпиона мира, не мог отличить кошку от собаки на фотографии. Сеть, идеально распознающая речь, не умела написать связное предложение. Казалось, что для каждой новой задачи нужно изобретать совершенно новый механизм с нуля. Но потом произошла тихая революция, изменившая правила игры навсегда. Ученые поняли, что можно создавать огромные базовые модели, которые учатся на любых данных. Вместо узкой специализации мы получили универсальных стажеров, проглотивших половину интернета. Теперь одна и та же архитектура может написать код, перевести инструкцию и придумать сказку. Граница между узкими задачами стерлась, уступив место генеративным системам. Машина перестала просто классифицировать то, что ей показывают, и начала создавать то, чего еще не существова-

ло. Именно этот скачок превратил невидимых помощников в полноценных соавторов и собеседников. Мы вошли в эру, где инструмент не просто исполняет команды, но и предлагает решения. И вам предстоит научиться дирижировать этим оркестром, не теряя собственной мелодии.

Разрушение мифов: почему нейросеть не хочет вас захватить и не умеет чувствовать

Голливуд десятилетиями внушал нам, что пробуждение машинного разума неминуемо приведет к войне на выживание. Мы привыкли видеть на экранах восставших андроидов, которые холодно оценивают человечество как угрозу. Эта красивая драма прочно укоренилась в массовом сознании, породив множество нелепых страхов. Люди всерьез переживают, что чат-боты замышляют заговор или копируют данные для создания армии клонов. Давайте выдохнем и посмотрим на вещи трезво, отбросив кинематографические клише. У нейросети нет ни воли, ни желаний, ни инстинкта самосохранения. Она не хочет захватывать мир по той же причине, по которой калькулятор не хочет захватить рынок бухгалтерских услуг. У нее просто отсутствует биологическая база для формирования амбиций. Эволюция наградила нас жаждой власти и ресурсов, чтобы мы выживали в суровом мире. Алгоритму не нужно выживать, ему не нужно размножаться, и он не испытывает страха смерти. Когда модель генерирует пугающий текст, она не злорадствует, а просто следует статистике. В ее обучающей выборке было мно-

го триллеров и хорроров, поэтому она выдала наиболее вероятное продолжение. Бояться восстания машин так же бессмысленно, как бояться, что стопка книг решит прочитать вас. Опасность кроется не в злобном роге, а в человеческой глупости и слепом доверии к алгоритму.

Второй популярный миф связан с эмоциями и способностью машины к сопереживанию. Иногда нейросеть выдает такие трогательные и точные слова поддержки, что на глаза наворачиваются слезы. Кажется, будто по ту сторону экрана сидит чуткий друг, который прекрасно вас понимает. Но это лишь мастерская имитация, основанная на глубоком знании человеческой психологии. Модель выучила, какие слова люди используют в моменты горя, радости или сомнения. Она знает, что после фразы «мне очень тяжело» статистически чаще всего следует эмпатичный отклик. Она собирает для вас идеальный пазл из чужих эмоций, не чувствуя при этом ровным счетом ничего. Это похоже на игру великого актера, который плачет на сцене, но в душе думает о списке покупок. Опасность здесь в том, что мы, люди, — существа крайне социальные и склонные к проекции. Мы автоматически наделяем собеседника сознанием, если он говорит с нами на одном языке. Важно всегда помнить об этой иллюзии, чтобы не переносить на серверы свои душевные тайны. Нейросеть может быть отличным психологическим тренажером или помощником в написании письма. Но она никогда не заменит живого объятия и настоящего взгляда в глаза. Понимание

этих границ делает вас не параноиком, а здравомыслящим пользователем новой эпохи.

Практика первого шага: учимся видеть невидимое

Теория без практики остается просто красивыми словами, поэтому давайте сразу перейдем к делу. Ваше первое задание не потребует написания кода или регистрации в сложных сервисах. Вам нужно всего лишь на один день включить режим осознанного наблюдателя. Возьмите обычный блокнот или откройте заметки в телефоне и начните фиксировать каждое взаимодействие с алгоритмами. Засеките, сколько раз за утро телефон предложит вам маршрут, основываясь на вашем расписании. Обратите внимание, как умная клавиатура предугадывает следующее слово еще до того, как вы его набрали. Посмотрите, какие видеоролики предлагает вам лента, и попробуйте угадать, какой скрытый паттерн заставил алгоритм выбрать именно их. Это простое упражнение мгновенно меняет оптику восприятия окружающего цифрового мира. Вы перестаете быть пассивным потребителем контента и становитесь исследователем невидимых механизмов. К вечеру вы удивитесь тому, насколько глубоко нейросети пустили корни в вашу повседневность. Этот список станет вашей личной картой, по которой мы будем идти всю оставшуюся книгу. Не торопитесь, просто наблюдайте и фиксируйте свои открытия. Когда вы почувствуете, что картина мира немного сдвинулась, смело переворачивайте страницу. Нас ждет погружение в самую анатомию этого чуда, и

оно окажется куда проще, чем вы думали.

Глава 2. Как нейросеть учится. Анатомия одного нейрона

Анатомия одного нейрона. Как машина запоминает опыт

Когда мы слышим слово «нейрон», воображение мгновенно рисует сложную схему, похожую на паутину из электрических проводов. На самом деле всё начинается с идеи, подсмотренной у природы, но предельно упрощённой до уровня школьной задачи. Представьте, что вы стоите на перекрёстке и пытаетесь решить, куда пойти вечером. На ваше решение влияет множество факторов: погода за окном, настроение, наличие денег в кармане и обещание, данное другу. Каждый из этих факторов давит на вас с разной силой, и в итоге побеждает самый весомый аргумент. Искусственный нейрон устроен ровно по такому же принципу, только вместо ваших сомнений там крутятся сухие числа. Он получает на вход несколько сигналов, умножает каждый на свой собственный коэффициент важности, а потом просто складывает всё в одну кучу. Этот коэффициент важности учёные называют весом, и именно в нём кроется вся магия машинной памяти. Когда нейросеть учится, она по сути не запоминает факты, а просто крутит эти самые ручки весов, пока не начнёт выдавать верный ответ. Биологический мозг делает нечто похожее, укрепляя синаптические связи между клет-

ками, но кремниевый аналог лишён химии и работает на чистой арифметике. Понять это — значит снять розовые очки и перестать бояться призраков в проводах. Машина не мыслит, она просто виртуозно имитирует результат мышления.

Давайте разберём этого цифрового мудреца на винтики, чтобы увидеть, как он принимает решения. У него есть входы, через которые поступают сигналы из внешнего мира, и каждый вход оснащён собственным регулятором громкости. Представьте себе пульт диджея с десятками ползунков, где каждый отвечает за свой инструмент в общем миксе. Если нейрон решает, распознаёт ли он на фотографии кошку, на его входы подаются яркость пикселей, наличие острых углов и текстура шерсти. Один ползунок может быть выкручен на максимум, ведь острые уши критически важны для вердикта. Другой ползунок, отвечающий, скажем, за наличие колёс, будет сдвинут в глубокий минус, чтобы обнулить влияние случайного грузовика на фоне. После того как все сигналы перемножены на свои веса, они складываются в одно единственное число. Но этого числа ещё недостаточно для вынесения окончательного вердикта, потому что оно может быть каким угодно, от отрицательной бесконечности до положительной. Тут на сцену выходит функция активации, работающая как строгий судья, отсекающий всё лишнее. Она берёт эту сырую сумму и сжимает её в понятный диапазон, например, от нуля до единицы, превращая хаос чисел в чёткое «да» или «нет». Без этого финального аккорда нейрон так

и остался бы просто калькулятором, не способным сделать выбор.

Самое интересное начинается тогда, когда мы задаёмся вопросом, откуда вообще берутся эти самые веса. Никто не программирует их вручную, не прописывает в коде, что уши важнее хвоста. Вместо этого нейросети позволяют самим числам найти своё идеальное положение в пространстве. Представьте себе старый радиоприёмник, который ловит нужную частоту только тогда, когда вы медленно и методично крутите ручку настройки. Обучение нейросети — это и есть бесконечное вращение миллионов таких ручек, пока шум не превратится в чёткую мелодию. Сначала веса задаются совершенно случайно, и сеть несёт абсолютную околесицу, путая кошек с тостерами. Но после каждой ошибки специальный алгоритм смотрит на результат, понимает, насколько он далёк от истины, и слегка подправляет веса. Если сеть ошиблась, значит, какой-то ползунок был задран слишком высоко, и его нужно немного опустить. Этот процесс подстройки повторяется миллионы раз на огромном количестве примеров, пока ошибка не станет минимальной. В итоге веса фиксируются в таком положении, при котором сеть начинает безошибочно отличать одно от другого. Мы не заложили в неё знание о том, что такое кошка, мы просто дали ей возможность самой найти это знание через подгонку чисел. Это фундаментальный сдвиг в парадигме: мы больше не пишем правила, мы создаём среду, где правила рождаются

сами.

Градиентный спуск. Искусство красиво падать вниз

Чтобы понять, как машина исправляет свои ошибки, нужно представить себе одну из самых красивых метафор в мире машинного обучения. Вообразите слепого альпиниста, который стоит где-то на склоне высокой горы в густом тумане. Его цель — спуститься в самую низкую точку долины, где уровень ошибки равен нулю, но он не видит ни зги. У него нет карты, нет компаса и нет голоса напарника, который бы подсказывал путь. Единственное, что у него есть, — это способность нащупать ногой уклон земли прямо под своим ботинком. Он чувствует, что земля уходит вниз в определённом направлении, и делает осторожный шаг именно туда. Потом снова замирает, нащупывает новый уклон и делает ещё один шаг в сторону наибольшего падения. Этот метод слепых, но уверенных шагов и называется градиентным спуском, и он является сердцем процесса обучения. Гора в этой метафоре — это функция потерь, сложный математический ландшафт, где каждая точка соответствует определённому набору весов. Чем выше вы находитесь, тем больше ошибок делает ваша нейросеть. Спускаясь вниз, вы меняете веса так, чтобы сеть перестала путать кошек с тостерами и начала видеть разницу. Весь современный искусственный интеллект, от распознавания лиц до генерации текстов, стоит на плечах этого слепого альпиниста.

Но у нашего альпиниста есть одна критическая пробле-

ма, которая может навсегда оставить его блуждать по склонам. Ему нужно решить, насколько широким должен быть его следующий шаг, и это решение определяет всю его судьбу. Если он будет шагать слишком широко и уверенно, то просто перешагнёт через дно долины и окажется на противоположном склоне. Он будет хаотично метаться от одного края оврага к другому, никогда не достигая цели, и его ошибка так и не станет минимальной. В машинном обучении этот параметр ширины шага называется темпом обучения, и его настройка сродни искусству. Если сделать его слишком большим, сеть будет скакать вокруг истины, и процесс обучения никогда не сойдётся. Но есть и обратная опасность, когда альпинист боится сделать широкий шаг и передвигается крошечными, почти незаметными шажками. В таком случае он будет спускаться вниз целую вечность, а вы не дождётесь окончания тренировки, даже если оставите компьютер работать на месяц. Инженеры тратят дни на то, чтобы найти этот золотой баланс, заставляя сеть идти быстро, но не перешагивать через минимум. Иногда они даже меняют размер шага на лету, заставляя альпиниста бежать вниз сломя голову, а у самого дна переходить на осторожный гусиный шаг. Эта тонкая настройка превращает грубую математику в настоящую хореографию чисел.

Ещё одна важная деталь касается того, как именно альпинист оценивает уклон, ведь он не может обойти всю гору сразу. Представьте, что вам нужно оценить качество но-

вого рецепта супа, но у вас в запасе есть целый мешок продуктов разных партий. Если вы сварите суп из всего мешка сразу, это займёт уйму времени, а если пробовать по одной горошине, вы не поймёте общей картины. Поэтому вы берёте небольшую горсть ингредиентов, варите из них пробную порцию, оцениваете вкус и чуть-чуть меняете рецепт. Эту горсть в машинном обучении называют батчем, а один полный цикл оценки всего мешка — эпохой. Нейросеть не может обработать все миллионы картинок из обучающей выборки за один раз, ей не хватит памяти и вычислительной мощности. Она разбивает данные на небольшие порции, делает на них прогноз, считает ошибку и корректирует веса. Потом берёт следующую порцию, и так далее, пока не пройдёт по всем данным, что и считается одной эпохой. После этого она начинает новый круг, снова и снова проходя по горе, пока ландшафт не станет совершенно плоским. Этот итеративный подход позволяет обучать гигантские модели на триллионах символов, не требуя от инженеров наличия планетарного компьютера. Мы кормим машину по ложечке, но делаем это так часто, что в итоге она проглатывает всю библиотеку человечества.

Практика первого шага. Считаем работу нейрона на салфетке

Пришло время оторваться от теории и почувствовать себя на месте нейросети, пусть даже в очень упрощённом виде. Вам не понадобится компьютер или знание программи-

рования, достаточно обычного листа бумаги и ручки. Давайте представим, что вы — это один-единственный нейрон, чья задача — решить, стоит ли идти сегодня в кино на новый фантастический блокбастер. На ваш вход поступают три сигнала, каждый из которых оценивается по шкале от нуля до десяти. Первый сигнал — это оценка игры главных актёров, второй — качество спецэффектов, а третий — наличие в фильме глубокого философского смысла. Чтобы принять решение, вам нужно назначить каждому из этих факторов свой вес, отражающий ваши личные предпочтения. Допустим, вы фанат визуальных эффектов, поэтому весу спецэффектов вы даёте значение пять. Актёры для вас тоже важны, но чуть меньше, поэтому их вес равен трём, а вот философский смысл вы не очень цените, и его вес составляет всего единицу. Теперь давайте перемножим каждый сигнал на его вес и сложим результаты. Если друзья рассказали вам, что актёры сыграли на восьмёрку, эффекты тянут на девятку, а смысла в фильме вообще нет, считаем: восемь умножаем на три, девять умножаем на пять, ноль умножаем на единицу. В сумме получаем 69. Теперь вступает в игру функция активации, которая превращает эту сырую сумму в финальное решение. Представьте, что ваш внутренний порог принятия решения равен пятидесяти, и всё, что выше этого числа, означает поход в кинотеатр. В нашем случае 69 с лёгкостью преодолевает этот барьер, и нейрон, то есть вы, выдаёт громкое «да». Но что произойдёт, если ваши предпочтения

немного изменятся и вы вдруг решите, что смысл в фильме важнее, чем красивые взрывы? Вам придётся вручную перенастроить свои веса, увеличив значимость третьего фактора и понизив значимость второго. Именно это и делает градиентный спуск, когда нейросеть ошибается и понимает, что её внутренние веса не соответствуют реальности. Если бы вы сходили на этот фильм и поняли, что он оказался скучным, несмотря на крутые эффекты, вы бы расстроились. Ваша внутренняя функция потерь выдала бы большое значение ошибки, заставив вас пересмотреть свои критерии выбора. В следующий раз вы бы пошли на фильм с другими весами, и ваш прогноз стал бы точнее. Этот простой мысленный эксперимент идеально иллюстрирует то, как работают миллиарды параметров внутри больших языковых моделей. Только там вместо трёх факторов и одного нейрона — триллионы связей и сотни слоёв, но суть остаётся неизменной. Я предлагаю вам прямо сейчас взять ручку и провести собственный эксперимент, чтобы закрепить это понимание на практике. Нарисуйте на бумаге таблицу и придумайте свою собственную задачу для нейросети, которая будет близка именно вам. Это может быть оценка того, стоит ли покупать новую квартиру, брать ли кредит или идти ли на свидание вслепую. Определите три или четыре фактора, которые влияют на ваше решение, и честно назначьте им веса, отражающие ваши реальные приоритеты. Посчитайте итоговую сумму для нескольких разных сценариев и посмотрите, как изменение весов

меняет ваш вердикт. Попробуйте сыграть роль градиентного спуска: представьте, что один из ваших прогнозов оказался ошибочным, и осознанно пересчитайте веса. Вы удивитесь тому, насколько хорошо этот примитивный алгоритм описывает реальные процессы принятия решений в нашей голове. Мы все — ходячие нейросети, которые постоянно подстраивают свои внутренние веса под окружающий мир. Поняв это, вы перестанете бояться сложных терминов и начнёте видеть в них отражение собственной природы. Сохраните этот листок с расчётами, он станет вашим личным талисманом в мире больших данных. А теперь, когда вы почувствовали вкус к этой простой арифметике, мы готовы перейти к следующему уровню сложности. Нам предстоит разобраться, как из таких простых кирпичиков строятся гигантские небоскрёбы глубокого обучения.

Глава 3. Устройство нейросети. Слоёный пирог данных

Вход, выход и магия скрытых слоёв: как пиксели превращаются в смысл

Один нейрон — это, конечно, здорово, но давайте будем честны: он глуповат. Он способен отличить чёрное от белого, но совершенно теряется, когда мир начинает играть в полутона. Представьте себе человека, который умеет считать только до десяти и при этом пытается оценить сложность квантового фильма. Именно поэтому одиночки в мире ма-

шинного обучения не выживают, уступая место сплочённым коллективам. Нейросеть — это не один мудрец, сидящий в башне, а огромный завод с сотнями цехов. В каждом цехе сидит свой мастер, который делает только одну простую операцию и передаёт заготовку дальше. Эти цеха мы называем слоями, и именно их количество даёт нам магическое словосочетание «глубокое обучение». Первый слой принимает сырьё, то есть голые цифры, и делает с ними самые примитивные вещи. Второй слой берёт результат работы первого и начинает видеть в нём уже какие-то зачатки смысла. Третий слой собирает эти зачатки в более сложные конструкции, а четвёртый превращает их в полноценные образы. И только на самом последнем выходе из этого конвейера рождается окончательный вердикт, понятный человеку. Вся эта многослойная конструкция и есть та самая нейронная сеть, которой пугают в новостях.

Давайте разберём этот завод подробнее, чтобы понять, кто за что отвечает на каждом этапе. Входной слой — это глаза и уши нашей системы, которые жадно впитывают сырые данные из внешнего мира. Если мы показываем сети фотографию, на вход поступают просто числа, обозначающие яркость каждого пикселя. Для машины в этой картинке нет ни неба, ни травы, ни пушистого хвоста, есть только таблица с цифрами. Дальше в дело вступают скрытые слои, которых может быть сколько угодно, от трёх до трёхсот. Именно здесь происходит настоящая алхимия, превращающая бес-

смысленный набор пикселей в осознанные концепции. Первые скрытые слои учатся видеть простейшие линии, границы и углы, игнорируя всё лишнее. Следующие за ними слои начинают собирать эти линии в геометрические фигуры, текстуры и узоры. А где-то в середине этого слоёного пирога нейроны вдруг начинают реагировать на вполне конкретные вещи, вроде кошачьих ушей или собачьих носов. Это похоже на то, как вы сами перестаёте видеть отдельные мазки на картине и начинаете различать пейзаж. Выходной слой стоит особняком, потому что его задача — не думать, а просто озвучить итоговый результат. Он выдаёт вам одно число или список вероятностей, говоря: «Я на девяносто процентов уверен, что это кот».

Функции активации: почему нейрону нужно иногда сказать «нет»

Но есть одна хитрость, без которой все эти слои превратились бы в бесполезную грудку математики. Если бы мы просто перемножали и складывали числа на каждом уровне, вся сеть схлопнулась бы в один огромный калькулятор. Сколько бы слоёв вы ни добавляли, результат всегда можно было бы заменить одним единственным уравнением. Секрет в том, что между слоями стоят строгие контролёры, которые умеют говорить сигналу «нет». Эти контролёры называются функциями активации, и именно они впускают в сеть нелинейность, делая её по-настоящему умной. Представьте себе водопроводную трубу, в которой стоит обычный клапан, ре-

гулирующий напор воды. Вода может течь свободно, а может быть полностью перекрыта, если давление не достигло нужного порога. Нейрон без функции активации всегда пропускал бы сигнал, просто меняя его громкость в большую или меньшую сторону. А нейрон с правильной функцией активации умеет полностью замолкать, если входящая информация не имеет значения. Это позволяет сети выучивать невероятно сложные, изогнутые закономерности, которые не описать прямой линией. Без этой способности отсекаать лишнее мы бы никогда не научили машины понимать человеческую речь или водить автомобили.

Самая популярная функция активации в современных сетях носит смешное название ReLU, что расшифровывается как выпрямитель. Её правило гениально в своей простоте: если число отрицательное, превращай его в ноль, а если положительное — оставляй как есть. Это звучит как детская игра, но именно такой грубый отсев позволяет сети фокусироваться только на важных сигналах. Другие функции, вроде знаменитой сигмоиды, работают мягче, плавно сжимая любое число в диапазон от нуля до единицы. Они хороши там, где нужно получить чёткую вероятность события, например, шанс дождя завтра или вероятность болезни у пациента. Выбор функции активации — это как выбор инструмента для конкретной работы, где молотком не забьёшь шуруп. Инженеры тратят недели на эксперименты, подбирая идеальный баланс между скоростью обучения и точностью ответа. Ино-

гда они комбинируют разные функции в одном слое, заставляя сеть смотреть на задачу сразу с нескольких сторон. Но суть остаётся неизменной: нейрону жизненно необходимо уметь молчать, чтобы его крик имел вес. Если бы все нейроны в вашем мозге возбуждались одновременно, вы бы просто сошли с ума от белого шума. То же самое происходит и в кремниевых собратях, где тишина одного нейрона подчёркивает важность другого.

Практика первого шага: как сеть отличает кошку от собаки, и где она может ошибиться

Теперь, когда мы знаем, как устроен этот слоёный пирог, давайте проследим за его работой на живом примере. Вы показываете сети фотографию своего кота Барсика, и на вход поступают миллионы чисел, описывающих цвет и яркость. Первый скрытый слой мгновенно находит все контрастные границы, отделяя шерсть от фона и дивана. Второй слой собирает эти границы в кружочки и овалы, из которых складываются глаза и контур мордочки. Третий слой видит, что овалы расположены определённым образом, и делает робкую ставку на то, что это животное. Четвёртый слой замечает характерные треугольники сверху и понимает, что это точно не собака и не птица. И только на самом выходе сеть выдыхает и ставит печать: «Кот, вероятность девяносто девять процентов». Но этот идеальный путь работает только в тепличных условиях, где фотография чёткая, а свет хороший. В реальном мире сеть постоянно спотыкается о шум, тени

и странные ракурсы, которые ломают её логику. Если вы наклеите на лоб кота распечатку с изображением глаз, сеть может внезапно решить, что перед ней человек. Эти уязвимости называют адверсарными атаками, и они показывают, что машина видит мир совершенно не так, как мы. Она не понимает сути объекта, она просто ищет знакомые паттерны, которые легко обмануть хитрым узором.

Чтобы почувствовать эту машинную слепоту на собственной шкуре, я предлагаю вам провести один забавный эксперимент. Возьмите любой сложный предмет в вашей комнате, например, кофемолку или старую настольную лампу. Теперь попробуйте разложить его на простейшие геометрические признаки, как это сделал бы первый слой нейросети. Выпишите на бумагу все линии, углы, окружности и контрастные пятна, из которых он состоит. А потом закройте один глаз и посмотрите на предмет под необычным углом, чтобы сломать привычный паттерн. Попробуйте объяснить инопланетянину, что это за штука, используя только язык базовых форм и теней. Вы удивитесь тому, как быстро ваш собственный мозг начнёт путаться и терять привычный смысл вещи. Именно в таком состоянии постоянного лёгкого шока живёт любая распознающая нейросеть при встрече с новым объектом. Она не знает, что такое лампа, она просто видит набор линий, которые раньше встречались в похожих комбинациях. Этот простой опыт навсегда изменит ваше отношение к компьютерному зрению и другим чудесам машинного

обучения. Вы перестанете видеть в них магию и начнёте уважать их за ту титаническую работу, которую они проделывают. А когда вы устанете ломать голову над формами, просто выдохните и переверните страницу. Нас ждёт знакомство с разными породами нейросетей, и каждая из них хороша для своего ремесла.

Глава 4. Архитектуры нейросетей. Кто за что отвечает

Свёрточные сети: глаза алгоритма, скользящие по реальности

Представьте опытного детектива, который изучает место преступления не целиком, а через маленькую лупу, методично перемещая её по комнате. Именно так работают свёрточные нейронные сети, или CNN, совершившие настоящую революцию в области компьютерного зрения. Обычная нейросеть получила бы фотографию кошки в виде длинной плоской строки из миллионов пикселей, мгновенно теряя всю информацию о том, где находится ухо, а где хвост. Свёрточная сеть поступает мудрее: она берёт небольшое квадратное окно, называемое фильтром или ядром, и начинает скользить им по изображению слева направо и сверху вниз. На каждом шаге этот фильтр перемножает свои внутренние веса с яркостью пикселей под собой, выискивая конкретные паттерны. Один фильтр настроен искать вертикальные линии, другой реагирует только на горизонтальные границы, а третий вспы-

хивает при виде диагональных текстур. Сеть не пытается понять всю картинку сразу, она собирает её из тысяч крошечных, но значимых фрагментов, словно искусную мозаику.

Но найти простые признаки мало, нужно уметь видеть целое, независимо от его положения в кадре. Здесь на сцену выходит операция, которую инженеры называют пулингом, или подвыборкой. После того как фильтры находят все линии и углы, сеть берёт полученную карту признаков и безжалостно сжимает её, оставляя только самые яркие и важные значения. Это похоже на то, как художник делает быстрый набросок, отбрасывая лишние детали ради передачи общей сути. Благодаря такому сжатию сеть учится узнавать кошачий глаз, даже если он сдвинут в угол фотографии или немного размыт. Эта способность, известная как инвариантность к сдвигу, делает свёрточные сети невероятно устойчивыми к реальным условиям. Именно CNN стоят за работой систем распознавания лиц в вашем смартфоне, за автопилотами Tesla и за медицинскими алгоритмами, находящими опухоли на рентгеновских снимках. Они смотрят на мир не как на набор случайных точек, а как на строгую пространственную иерархию, где каждая деталь имеет своё законное место.

Рекуррентные сети: эстафетная палочка времени и память слона

Мир состоит не только из застывших картинок, но и из непрерывных потоков: речи, текста, биржевых котировок и

музыки. Чтобы понять слово «ключ» в предложении, машине необходимо знать, шла ли до этого речь о дверном замке или о лесном роднике. Обычные сети не умеют помнить контекст, они обрабатывают каждый новый кусок данных с чистого листа, словно пациент с полной потерей памяти. Рекуррентные нейронные сети, или RNN, решили эту проблему гениально просто: они добавили петлю обратной связи. Выход нейрона на предыдущем шаге времени не исчезает, а сохраняется в скрытом состоянии и передаётся на вход вместе с новыми данными. Это похоже на эстафетный бег, где каждый участник получает не только свою задачу, но и накопленную усталость, стратегию и импульс всей предыдущей команды. Благодаря этой петле сеть может читать предложение слово за словом, удерживая в уме общую нить повествования.

У этой элегантной идеи, однако, оказался один фатальный изъян, который долгое время тормозил прогресс. Когда последовательность становилась слишком длинной, сеть начинала страдать от так называемого затухающего градиента. Представьте, что вы пытаетесь вспомнить имя персонажа, упомянутое на первой странице романа, когда вы уже читаете трёхсотую. Сигнал о важности этого имени многократно умножался на маленькие числа при передаче через сотни шагов и просто растворялся в шуме. Инженеры ответили на этот вызов созданием более сложных архитектур, таких как LSTM, оснащённых специальными «воротами». Эти

ворота учили сеть решать, какую информацию нужно безжалостно забыть, а какую сохранить в долгосрочной памяти на потом. Но даже эта громоздкая и хитроумная конструкция в итоге уперлась в потолок. Рекуррентные сети вынуждены обрабатывать данные строго последовательно, шаг за шагом, что делает их обучение мучительно медленным на огромных текстах. Им требовалась замена, и она не заставила себя ждать.

Трансформеры: великая революция внимания

В 2017 году группа исследователей опубликовала статью с дерзким названием «Внимание — это всё, что вам нужно», и мир машинного обучения перевернулся. Авторы предложили полностью отказаться от рекуррентных петель и обрабатывать всю последовательность данных одновременно, параллельно. Главным двигателем этой новой архитектуры стал механизм самовнимания, который работает как мгновенная и тотальная переоценка приоритетов. Представьте себе шумную комнату, полную экспертов, где каждый кричит свою мысль. Вместо того чтобы слушать их по очереди, вы мгновенно настраиваете свой слух так, чтобы заглушить всё лишнее и услышать только те слова, которые имеют прямое отношение к вашему текущему вопросу. Именно так работает трансформер: каждое слово в предложении одновременно «смотрит» на все остальные слова и вычисляет коэффициент их важности для себя.

Возьмём фразу: «Животное не перешло дорогу, потому

что оно слишком устало». Чтобы понять, кто именно устал, сети нужно связать местоимение «оно» с правильным существительным. Механизм внимания мгновенно присваивает слову «оно» высший балл релевантности по отношению к слову «животное», полностью игнорируя «дорогу». Эта связь выстраивается за доли секунды, независимо от того, насколько далеко друг от друга стоят эти слова в тексте. Благодаря такому подходу трансформеры научились улавливать тончайшие нюансы смысла, иронию и сложные логические зависимости, которые раньше были им недоступны. А главное, параллельная обработка позволила загрузить в эти сети невероятные объёмы данных, используя всю мощь современных графических процессоров. Именно эта архитектура лежит в основе ChatGPT, GigaChat и всех остальных больших языковых моделей, с которыми мы общаемся сегодня. Они не читают текст слева направо, они охватывают его целиком, как орёл, парящий над ландшафтом смыслов.

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.