

ВЕДЬ ЭТО ТАК ПРОСТО!



Искусственный интеллект

для
Чайников[®]

Издательство ДИАЛЕКТИКА

Искусственный
интеллект и общество

Искусственный интеллект
роботов, дронов и беспилотных
автомобилей

Пределы возможностей
искусственного
интеллекта

Джон Пол Мюллер
Лука Массарон



Искусственный интеллект

для
чайников[®]



Artificial Intelligence

By John Paul Mueller and Luca Massaron

**for
dummies[®]**
A Wiley Brand



Искусственный интеллект

Джон Пол Мюллер и Лука Массарон

для
чайников®



Москва ♦ Санкт-Петербург
2019

ББК 32.973.26-018.2.75

М98

УДК 681.3.07

ООО “Диалектика”

Зав. редакцией *С.Н. Тризуб*

Перевод с английского и редакция *В.А. Коваленко*

По общим вопросам обращайтесь в издательство “Диалектика” по адресу:

info@dialektika.com, <http://www.dialektika.com>

Мюллер, Джон Пол, Массарон, Лука.

М98 Искусственный интеллект для чайников. : Пер. с англ. — СПб. : ООО “Диалектика”, 2019. — 384 с. : ил. — Парал. тит. англ.

ISBN 978-5-907114-57-9 (рус.)

ББК 32.973.26-018.2.75

Все названия программных продуктов являются зарегистрированными торговыми марками соответствующих фирм.

Никакая часть настоящего издания ни в каких целях не может быть воспроизведена в какой бы то ни было форме и какими бы то ни было средствами, будь то электронные или механические, включая фотокопирование и запись на магнитный носитель, если на это нет письменного разрешения издательства John Wiley & Sons, Inc.

Copyright © 2019 by Dialektika Computer Publishing.

Original English edition Copyright © 2018 by John Wiley & Sons, Inc., Hoboken, New Jersey.

All rights reserved including the right of reproduction in whole or in part in any form. This translation is published by arrangement with John Wiley & Sons, Inc.

No part of this publication may be reproduced, stored in a retrieval system or transmitted in any form or by any means, electronic, mechanical, photocopying, recording, scanning or otherwise without the prior written permission of the Publisher.

Научно-популярное издание

Джон Пол Мюллер, Лука Массарон

Искусственный интеллект для чайников

Подписано в печать 27.02.2019. Формат 70х100/16.

Усл. печ. л. 30,96. Уч.-изд. л. 24,7.

Тираж 500 экз. Заказ № 1903.

Отпечатано в АО “Первая Образцовая типография”

Филиал “Чеховский Печатный Двор”

142300, Московская область, г. Чехов, ул. Полиграфистов, д. 1

Сайт: www.chpd.ru, E-mail: sales@chpd.ru, тел. 8 (499) 270-73-59

ООО “Диалектика”, 195027, Санкт-Петербург, Магнитогорская ул., д. 30, лит. А, пом. 848

ISBN 978-5-907114-57-9 (рус.)

ISBN 978-1-119-46765-6 (англ.)

© 2019, ООО “Диалектика”

© 2018 by John Wiley & Sons,
Inc., Hoboken, New Jersey

Оглавление

Введение	17
Часть 1. Введение в искусственный интеллект	23
Глава 1. Знакомство с искусственным интеллектом	25
Глава 2. Определение роли данных	43
Глава 3. Вопросы использования алгоритмов	65
Глава 4. Первенство специализированных аппаратных средств	83
Часть 2. Использование искусственного интеллекта в обществе	99
Глава 5. Искусственный интеллект в компьютерных приложениях	101
Глава 6. Автоматизация наиболее популярных процессов	115
Глава 7. Применение искусственного интеллекта в медицине	125
Глава 8. Искусственный интеллект в человеческом общении	147
Часть 3. Программно-ориентированные приложения искусственного интеллекта	159
Глава 9. Анализ данных для искусственного интеллекта	161
Глава 10. Машинное обучение в искусственном интеллекте	179
Глава 11. Усиление искусственного интеллекта глубоким обучением	203
Часть 4. Работа с искусственным интеллектом в аппаратных приложениях	231
Глава 12. Разработка роботов	233
Глава 13. Полеты с дронами	249
Глава 14. Автомобиль, управляемый искусственным интеллектом	263
Часть 5. Будущее искусственного интеллекта	283
Глава 15. Причины неудач приложений	285
Глава 16. Искусственный интеллект в космосе	301
Глава 17. Новые профессии	321
Часть 6. Великолепные десятки	339
Глава 18. Десять профессий, недоступных искусственному интеллекту	341
Глава 19. Десять достижений искусственного интеллекта, наиболее полезных для общества	351
Глава 20. Десять областей, в которых искусственный интеллект терпит неудачу	361
Предметный указатель	371

Содержание

Об авторе	15
Введение	17
Часть 1. Введение в искусственный интеллект	23
Глава 1. Знакомство с искусственным интеллектом	25
Определение термина “искусственный интеллект”	26
Распознавание интеллекта	26
Выработка четырех способов определения искусственного интеллекта	31
История искусственного интеллекта	35
Символическая логика в Дартмуте	35
Экспертные системы	36
Преодоление зим искусственного интеллекта	37
Области применения искусственного интеллекта	38
Не обмануться в ожиданиях от искусственного интеллекта	40
Объединение искусственного интеллекта с базовым компьютером	41
Глава 2. Определение роли данных	43
Поиск данных, вездесущих на данном этапе	44
Понятие значений Мура	45
Данные используются повсюду	47
Приведение алгоритмов в действие	48
Успешное использование данных	50
Источники данных	51
Получение надежных данных	51
Повышение надежности ввода данных человеком	52
Автоматизированный сбор данных	53
Маникюр данных	54
Как справиться с отсутствием данных	55
Учет рассогласования данных	56
Отделение полезных данных от других	57
Пять недостоверностей данных	57
Усердие	58
Умолчание	59

Точка зрения	59
Предубежденность	61
Недопонимание	61
Установление пределов сбора данных	62
Глава 3. Вопросы использования алгоритмов	65
Понятие роли алгоритмов	66
Что означает алгоритм	67
Планирование и ветвление	68
Состязательные игры	71
Использование локального поиска и эвристики	73
Понятие обучения машины	77
Работа экспертных систем	78
Введение в машинное обучение	81
Достижение новых высот	81
Глава 4. Первенство специализированных аппаратных средств	83
Стандартные аппаратные средства	85
Понятие стандартных аппаратных средств	85
Недостатки стандартных аппаратных средств	86
Использование графических процессоров	88
Понятие узкого места фон Неймана	89
Определение GPU	90
Причины хорошей работоспособности GPU	91
Создание специализированной среды обработки	92
Увеличение возможностей аппаратных средств	93
Добавление специализированных сенсоров	95
Разработка методов взаимодействия с окружающей средой	96
Часть 2. Использование искусственного интеллекта в обществе	99
Глава 5. Искусственный интеллект в компьютерных приложениях	101
Наиболее популярные типы приложений	102
Использование искусственного интеллекта в типичных приложениях	103
Области применения искусственного интеллекта	104
Аргумент китайской комнаты	104
Как искусственный интеллект делает приложения более дружелюбными	106

Автоматическое исправление	107
Виды исправлений	108
Преимущества автоматических исправлений	108
Почему автоматические исправления не срабатывают	109
Создание рекомендаций	110
Получение рекомендаций на основании предыдущих действий	110
Получение рекомендаций на основании групп	110
Получение неправильных рекомендаций	111
Учет ошибок искусственного интеллекта	112
Глава 6. Автоматизация наиболее популярных процессов	115
Решения для скуки	116
Как сделать задачи интереснее	117
Как помочь работать эффективнее	118
Как искусственный интеллект борется со скукой	118
Как искусственный интеллект не может бороться со скукой	119
Работа в промышленных условиях	120
Уровни автоматизации	120
Больше чем просто роботы	121
Не автоматизацией единой	122
Безопасность окружающей обстановки	123
Роль скуки в происшествиях	123
Как искусственный интеллект помогает избежать проблем безопасности	123
Почему искусственный интеллект не может устранить проблемы безопасности	124
Глава 7. Применение искусственного интеллекта в медицине	125
Носимый терапевтический монитор	127
Ношение полезных мониторов	127
Доверие к критически важным мониторам	128
Использование носимых мониторов	128
Расширение возможностей людей	130
Использование игр для терапии	132
Использование экзоскелетов	133
Решение специфических задач	135
Программно-ориентированные решения	135
Доверие к аппаратному усилению	136
Искусственный интеллект эндопротеза	137

Новые способы анализа	138
Новые хирургические технологии	139
Выработка хирургических рекомендаций	139
Помощь хирургу	140
Замена хирургии наблюдением	141
Автоматизация решений	142
Работа с медицинскими записями	142
Предсказание будущего	143
Повышение безопасности процедур	144
Создание лучших медикаментов	144
Объединение роботов и медицинских специалистов	145
Глава 8. Искусственный интеллект в человеческом общении	147
Новые способы коммуникации	148
Создание новых алфавитов	149
Автоматизация перевода	150
Включение жестикуляции	151
Обмен идеями	153
Установление связи	153
Улучшение коммуникаций	154
Определение тенденций	154
Использование средств массовой информации	155
Улучшение сенсорного восприятия человека	156
Смещение спектра данных	156
Усиление человеческих органов чувств	157
Часть 3. Программно-ориентированные приложения искусственного интеллекта	159
Глава 9. Анализ данных для искусственного интеллекта	161
Анализ данных	162
Почему анализ столь важен	165
Пересмотр значения данных	166
Машинное обучение	168
Как работает машинное обучение	169
Преимущества машинного обучения	171
Будешь полезным — станешь обыденным	173
Пределы машинного обучения	173

Обучение на основе данных	175
Контролируемое обучение	176
Неконтролируемое обучение	177
Обучение с подкреплением	177
Глава 10. Машинное обучение в искусственном интеллекте	179
Способы обучения	180
Пять основных подходов к обучению искусственного интеллекта	180
Три наиболее перспективных подхода к обучению искусственного интеллекта	183
Ожидание следующего прорыва	184
Поиск правды в вероятностях	185
На что способны вероятности	186
Учет предыдущих знаний	188
Представление реального мира как графа	192
Рост деревьев, способных классифицировать	196
Предсказание результатов при разделении данных	196
Принятие решений на основании деревьев	199
Отсечение переросших деревьев	200
Глава 11. Усиление искусственного интеллекта глубоким обучением	203
Формирование нейронных сетей, подобных человеческому мозгу	204
Знакомство с нейронами	204
Сначала был чудесный перцептрон	206
Подражание учащемуся мозгу	208
Простые нейронные сети	208
Раскроем секреты коэффициентов	209
Роль обратного распространения ошибки	210
Знакомство с глубоким обучением	211
Различия в глубоком обучении	212
Поиск еще более умных решений	214
Выявление краев и форм образов	217
Распознавание символов	217
Как работает свертка	219
Соревнования с использованием образов	221
Обучение подражанию искусству и жизни	222
Запоминание важных последовательностей	223

Магия общения искусственного интеллекта	223
Соревнования между двумя искусственными интеллектами	226
Часть 4. Работа с искусственным интеллектом в аппаратных приложениях	231
Глава 12. Разработка роботов	233
Роли роботов	234
Разоблачение научно-фантастических представлений о роботах	236
Почему так трудно быть гуманоидом	239
Сотрудничество с роботами	242
Сборка простого робота	245
Компоненты	245
Восприятие мира	247
Управление роботом	247
Глава 13. Полеты с дронами	249
На грани искусства	250
Беспилотный вылет на задание	250
Знакомство с квадрокоптером	252
Области применения дронов	254
Дроны на невоенных ролях	255
Усиление дронов искусственным интеллектом	258
Проблемы регулирования	261
Глава 14. Автомобиль, управляемый искусственным интеллектом	263
Краткая история	264
Будущее мобильности	265
Восхождение на шестой уровень автономности	266
Пересмотр мнения о роли автомобилей	268
Введение в беспилотные автомобили	272
Объединение всех технологий	273
Появление на сцене искусственного интеллекта	274
Не искусственным интеллектом единым	275
Преодоление неопределенности восприятия	277
Знакомство с сенсорами автомобиля	278
Объединение обнаруженного	280

Часть 5. Будущее искусственного интеллекта	283
Глава 15. Причины неудач приложений	285
Использование искусственного интеллекта там, где он не будет работать	286
Определение границ искусственного интеллекта	286
Неправильное применение искусственного интеллекта	290
Мир нереалистичных ожиданий	291
Влияние зим искусственного интеллекта	292
Понятие зимы искусственного интеллекта	292
Причины зим искусственного интеллекта	293
Пересмотр ожиданий и новые цели	295
Решения в поиске задачи	297
Определение примочки	298
Реклама	299
Когда люди справляются лучше	299
Поиск простого решения	299
Глава 16. Искусственный интеллект в космосе	301
Наблюдение за Вселенной	302
Первый ясный взгляд	303
Поиск новых мест	304
Эволюция Вселенной	305
Новые научные принципы	305
Добыча ресурсов в космосе	307
Добыча воды	308
Добыча редкоземельных и других металлов	308
Поиск новых элементов	311
Улучшение коммуникаций	311
Исследование новых мест	312
Сначала зонды	313
Роботизированные миссии	314
Добавление человеческого элемента	316
Создание строений в космосе	317
Ваш первый отпуск в космосе	317
Проведение научных исследований	318
Промышленное пространство в космосе	318
Использование космоса для хранения	319

Глава 17. Новые профессии	321
Жизнь и работа в космосе	322
Строительство городов в опасной окружающей среде	324
Строительство городов в океане	324
Создание космических поселений	325
Строительство лунных баз	327
Повышение эффективности людей	329
Решение проблем в планетарном масштабе	331
Как устроен мир	332
Выявление потенциальных источников проблем	334
Поиск возможных решений	335
Контроль результатов решений	335
Повторная попытка	336
Часть 6. Великолепные десятки	339
Глава 18. Десять профессий, недоступных искусственному интеллекту	341
Общение с людьми	342
Обучение детей	342
Медицинский уход	343
Удовлетворение личных нужд	343
Решение проблем, связанных с развитием	344
Создание нового	345
Изобретения	345
Искусство	346
Воображение нереального	347
Интуитивное принятие решений	347
Расследование преступлений	347
Контроль ситуации в режиме реального времени	348
Распознавание фактов и вымысла	348
Глава 19. Десять достижений искусственного интеллекта, наиболее полезных для общества	351
Учет взаимодействий, специфических для человека	352
Изобретение активного протеза ноги	353
Постоянный мониторинг	353
Прием лекарств	354

Выработка промышленных решений	354
Искусственный интеллект и трехмерная печать	355
Передовые роботизированные технологии	355
Создание новых технологических сред	357
Разработка новых редких ресурсов	357
Увидеть невидимое	357
Сотрудничество с искусственным интеллектом в космосе	358
Доставка товаров на космические станции	358
Добыча внепланетных ресурсов	359
Исследование других планет	359
Глава 20. Десять областей, в которых искусственный интеллект терпит неудачу	361
Понимание	362
Интерпретация, а не анализ	363
Выход за пределы чистых чисел	364
Учет последствий	364
Открытия	365
Получение новых данных из старых	366
Видение вне шаблонов	366
Реализация новых чувств	367
Сочувствие	367
Войти в чужое положение	368
Выработка истинных отношений	368
Смена точки зрения	369
Вера в недоказуемое	369
Предметный указатель	371

Об авторе

Джон Мюллер — внештатный автор и технический редактор. Умение писать у него в крови, и на настоящий момент он является автором 108 книг и более чем 600 статей. Разнообразие тем распространяется от работы с сетями до искусственного интеллекта и от управления базами данных до автоматического программирования. В некоторых его книгах затрагиваются темы науки о данных, самообучения машин и алгоритмов. Его технические навыки редактирования помогли более чем 70 авторам усовершенствовать содержимое своих произведений. Джон оказывает услуги технического редактирования различным журналам, консультирует и пишет сертификационные экзамены. Обратитесь к блогу Джона по адресу <http://blog.johnmuellerbooks.com/>. Джон доступен в Интернете по адресу John@JohnMuellerBooks.com. У него также есть веб-сайт по адресу <http://www.johnmuellerbooks.com/>.

Лука Массарон — аналитик данных и директор по маркетинговым исследованиям, специализирующийся на многомерном статистическом анализе, обучении машин и изучении ожиданий потребителей, с более чем десятилетним опытом в решении реальных проблем и консультировании заинтересованных лиц на основании обоснованных доводов, статистики, данных и алгоритмов. Испытывая страсть ко всему, что касается данных, их анализа и демонстрации потенциальных возможностей управления данными как экспертам, так и не специалистам, Лука является соавтором, наряду с Джоном Мюллером, таких книг, как *Python for Data Science For Dummies*, *Machine Learning For Dummies* и *Алгоритмы для чайников*. Предпочитая простоту ненужной изощренности, он верит, что математику вполне можно объяснить простыми словами, а практика помогает в изучении любой дисциплины.

Посвящение Джона

Эта книга посвящается моим друзьям по библиотеке La Valle, добровольным сотрудником которой я являюсь. Каждую неделю я жду встречи с вами, поскольку вы делаете мою жизнь полнее.

Посвящение Луки

Эта книга посвящается семье Суда (Suda) из Токио: Ёсики (Yoshiki), Такайо (Takayo), Макико (Makiko) и Микико (Mikiko).

Благодарности Джона

Благодарю мою жену Ребекку. Хотя ее уже нет, ее дух присутствует в каждой написанной мной книге и в каждом слове на каждой странице. Она верила в меня, когда не верил никто.

Рус Маллен (Russ Mullen) заслужил благодарность за техническое редактирование этой книги. Он привнес точность и глубину в материал, который вы видите здесь. Рус всегда подбрасывает мне замечательные URL для новых произведений и идей. Он также осуществляет санитарную проверку моих работ.

Мэтт Вагнер (Matt Wagner), мой агент, заслуживает благодарности за то, что помог мне получить первый контракт и беспокоился о мелочах, которые большинство авторов действительно не учитывают. Я всегда ценю его помощь. Приятно знать, что кто-то хочет тебе помочь.

Многие люди прочитали всю эту книгу или ее часть, чтобы помочь мне усовершенствовать подход, проверить примеры и предоставить мне отзыв о том, что им не понравилось. Эти добровольные помощники помогли слишком во многом, чтобы упомянуть здесь все. Я особенно ценю усилия Евы Битти (Eva Beattie) и Освальдо Тельес Алмиралл (Osvaldo Téllez Almirall), которые предоставили общий отзыв, прочитав всю книгу, и самоотверженно посвятили себя этому проекту.

И наконец, я хотел бы поблагодарить Кэти Мор (Katie Mohr), Сьюзен Кристоферсен (Susan Christophersen) и остальных членов редакционного и технического коллектива.

Благодарности Луки

Моя первая и самая большая благодарность — моей семье, Юкико (Yukiko), Амелии (Amelia), за их поддержку, жертвы, любовь и терпение на протяжении долгих дней, ночей, недель и месяцев, пока я работал над этой книгой.

Я благодарю весь редакционный и технический коллектив издательства Wiley, в частности — Кэти Мор и Сьюзен Кристоферсен, за их высокий профессионализм и поддержку на всех этапах создания этой книги из серии *...для чайников*.

Введение

Вряд ли сегодня можно найти человека, который не слышал об искусственном интеллекте. Мы видим его в фильмах, в книгах, в новостях и в Сети. Искусственный интеллект — это часть роботов, беспилотных автомобилей, дронов, медицинских систем, сайтов интернет-магазинов и других технологий, влияющих на нашу повседневную жизнь очень многими способами.

Множество ученых мужей заваливают нас информацией (и дезинформацией) об искусственном интеллекте. Одни считают искусственный интеллект мягким и пушистым, а другие видят в нем потенциального массового убийцу рода человеческого. Проблема с захлестывающим валом информации в том, что мы из всех сил пытаемся отделить реальность от плода возбужденного воображения. По большей части причиной лжи об искусственном интеллекте являются чрезмерные и нереалистичные ожидания ученых и предпринимателей. Вам необходима эта книга, если вы чувствуете, что действительно мало что знаете о технологии, которая становится немаловажным элементом вашей жизни.

Если взять за отправную точку средства массовой информации, то складывается впечатление, что большинство полезных технологий весьма скучны. Конечно, никто им особенно не доверяет. То же самое относится к искусственному интеллекту: он столь вездесущ, что становится нудным. В некотором роде вы используете его даже сегодня; фактически вы, вероятно, полагаетесь на искусственный интеллект многими разными способами — вы просто не обращаете на это внимания, поскольку это так обыденно. Данная книга позволяет узнать о вполне реальном и насущном применении искусственного интеллекта. Интеллектуальный термостат в вашем доме не может показаться весьма захватывающим, но это невероятное практическое применение технологии, от которой некоторые люди в ужасе пускаются наутек.

Конечно, в этой книге рассматриваются и действительно замечательные способы использования искусственного интеллекта. Например, вы можете не знать, что существует система медицинского контроля, фактически способная прогнозировать возникновение кардиологических проблем, но такое устройство существует. Искусственный интеллект управляет беспилотными летательными аппаратами и автомобилями, а также всякого рода роботами. Сегодня вы уже видите применение искусственного интеллекта в приложениях для космических летательных аппаратов, он также фигурирует практически во всех космических приключениях людей в будущем.

В отличие от множества книг на эту тему, в данной книге представлена правда о том, где и как искусственный интеллект работать не может. Фактически он не может участвовать в определенных важных действиях ни сейчас, ни в далеком будущем. Некоторые люди попытаются уверить вас, что эти действия вполне возможны для искусственного интеллекта, но мы расскажем, почему он не сможет работать, а также развеем все мифы об искусственном интеллекте. Каждый вынесет из этой книги то, что люди всегда будут важны. На самом деле искусственный интеллект делает людей еще более важными, причем такими способами, которые вы даже не могли себе вообразить.

О книге

Эта книга поможет вам понять, что такое искусственный интеллект, как он должен работать и почему он терпел неудачи в прошлом. Вы узнаете о причинах некоторых проблем с искусственным интеллектом, а также о том, что в некоторых случаях их почти невозможно решить. Кроме того, вы узнаете о решении некоторых проблем и областях применения учеными искусственного интеллекта для поиска ответов.

Для выживания технология должна иметь набор солидных приложений, которые фактически работают. Она также должна обеспечивать инвесторам окупаемость и перспективу, чтобы они вкладывали в нее капитал. В прошлом искусственный интеллект не был в состоянии достичь существенного успеха, поскольку ему недоставало некоторых из этих средств. Искусственный интеллект страдал также опережением времени: чтобы фактически преуспеть, он нуждался в современных аппаратных средствах. Сегодня вы можете найти искусственный интеллект в различных компьютерных приложениях и автоматизированных процессах. Он широко используется в областях медицины и взаимодействия между людьми. Искусственный интеллект также тесно связан с анализом данных, машинным обучением и глубоким обучением. Эти термины не для всех очевидны, поэтому понимание взаимосвязи данных технологий может быть одной из причин для прочтения книги.

Сегодня у искусственного интеллекта действительно яркое будущее, поскольку он стал весьма важной технологией. В этой книге приведены также пути, которыми, вероятно, будет следовать искусственный интеллект в будущем. Различные тенденции, обсуждаемые в этой книге, построены на основании того, что люди фактически пытаются сделать сейчас. Новая технология еще не добилась успеха, но поскольку над ней сейчас работает множество людей, ее шансы на успех действительно хороши.

Соглашения, принятые в книге

Здесь используются соглашения, общепринятые в компьютерной литературе.

- » Новые термины в тексте выделяются *курсивом*. Чтобы обратить внимание читателя на отдельные фрагменты текста, также применяется *курсив*.
- » Текст программ, функций, переменных, URL веб-страниц и другой код представлены моноширинным шрифтом.
- » Все, что придется вводить с клавиатуры, выделено **полужирным моноширинным** шрифтом.
- » Знакоместо в описаниях синтаксиса выделено *курсивом*. Это указывает на необходимость заменить знакоместо фактическим именем переменной, параметром или другим элементом, который должен находиться на этом месте: `BINDSIZE= (максимальная ширина колонки) * (номер колонки)`.
- » Пункты меню и названия диалоговых окон представлены следующим образом: Menu Option (Пункт меню).

Текст некоторых абзацев выделен специальным стилем. Это примечания, советы и предостережения, которые помогут обратить внимание на наиболее важные моменты в изложении материала и избежать ошибок в работе.



СОВЕТ

Советы хороши тем, что позволяют сэкономить время и упростить решение некой задачи. Советы в этой книге описывают экономящие время методики или содержат указатели на ресурсы, с которыми имеет смысл ознакомиться, чтобы получить максимум пользы от изучения искусственного интеллекта.



ВНИМАНИЕ!

Мы не хотим походить на строгих родителей или каких-то маньяков, но вам не следует делать то, что отмечено данной пиктограммой. В противном случае вы можете распространить дезинформацию об искусственном интеллекте, которой сегодня пугают людей.



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Увидев эту пиктограмму, знайте, что это дополнительный совет (или методика). Вы могли бы найти его очень полезным или слишком скучным, но он может содержать необходимое решение для создания или использования искусственного интеллекта. Пропускайте эти разделы, если хотите.



ЗАПОМНИ!

Если вы не вынесли ничего из некой главы или раздела, то запомните хотя бы материал, отмеченный этой пиктограммой. Такой текст обычно содержит наиболее важную информацию, которую следует знать для успешного ознакомления с искусственным интеллектом.

Источники дополнительной информации

Эта книга — не конец вашего изучения искусственного интеллекта, а только начало. Чтобы она стала для вас максимально полезной, мы предоставляем дополнительные источники информации. Получая от вас письма по электронной почте, мы сможем ответить на возникшие у вас вопросы, а также подсказать, как нововведения в области искусственного интеллекта и связанных с ним технологий затрагивают содержание книги. Вы также можете использовать следующие замечательные источники.

- » **Шпаргалка.** Вы помните, как использовали в школе шпаргалки, чтобы получить лучшие оценки по контрольной? Да, это разновидность шпаргалки. В ней содержится ряд заметок о малоизвестных задачах, решаемых с помощью искусственного интеллекта. Шпаргалка находится в конце книги. В ней содержится действительно полезная информация, включая расшифровку всех аббревиатур и сокращений, связанных с искусственным интеллектом, машинным обучением и глубинным обучением.
- » **Обновления.** Рано или поздно все изменяется. Например, мы могли не заметить грядущих изменений, когда смотрели в свои хрустальные шары во время написания этой книги. Когда-то это просто означало, что книга устарела и стала менее полезной, но теперь вы можете найти ее обновления по адресу www.dummies.com, если будете искать по ее заголовку (*Artificial Intelligence For Dummies*). Кроме этих обновлений, имеет смысл посетить блог автора по адресу <http://blog.johnmuellerbooks.com/>, содержащий ответы на вопросы читателей и связанные с книгой полезные материалы.

Что дальше

Пришло время приступить к изучению искусственного интеллекта и узнать, чем он может быть для вас полезен. Если вы ничего о нем не знаете, начните с главы 1. Вы можете не захотеть читать каждую главу книги, но, начав с главы 1, вы освоите основы искусственного интеллекта, что необходимо для изучения других глав книги.

Если вы читаете эту книгу, чтобы узнать, где сегодня используется искусственный интеллект, начните с главы 5. Часть 2 поможет вам узнать, где искусственный интеллект используется сегодня.

Читатели, которые немного лучше знакомы с искусственным интеллектом, могут начать с главы 9. В части 3 этой книги содержится самый современный материал. Если вы не хотите знать, как искусственный интеллект работает на низком уровне (не как разработчик, а просто как человек, интересующийся искусственным интеллектом), то вполне можете пропустить эту часть книги.

Хорошо, вы хотите знать фантастические способы использования искусственного интеллекта сегодня и в будущем. Если это так, то начните с главы 12. В частях 4 и 5 демонстрируются невероятные способы применения искусственного интеллекта без необходимости иметь дело с горами лжи. Часть 4 сосредоточена на аппаратных средствах, реализующих искусственный интеллект, а в части 5 уделяется больше внимания футуристическим способам использования искусственного интеллекта.

Ждем ваших отзывов!

Вы, читатель этой книги, и есть главный ее критик. Мы ценим ваше мнение и хотим знать, что было сделано нами правильно, что можно было сделать лучше и что еще вы хотели бы увидеть изданным нами. Нам интересны любые ваши замечания в наш адрес.

Мы ждем ваших комментариев и надеемся на них. Вы можете прислать нам бумажное или электронное письмо либо просто посетить наш веб-сайт и оставить свои замечания там. Одним словом, любым удобным для вас способом дайте нам знать, нравится ли вам эта книга, а также выскажите свое мнение о том, как сделать наши книги более интересными для вас.

Отправляя письмо или сообщение, не забудьте указать название книги и ее авторов, а также свой обратный адрес. Мы внимательно ознакомимся с вашим мнением и обязательно учтем его при отборе и подготовке к изданию новых книг.

Наши электронные адреса:

E-mail: info@dialektika.com

WWW: <http://www.dialektika.com>



Введение в искусственный интеллект

В ЭТОЙ ЧАСТИ...

- » Что искусственный интеллект фактически может сделать для вас**
- » Как данные влияют на использование искусственного интеллекта**
- » Как искусственный интеллект полагается на алгоритмы при выполнении полезной работы**
- » Как специализированные аппаратные средства улучшают работу искусственного интеллекта**



Глава 1

Знакомство с искусственным интеллектом

В ЭТОЙ ГЛАВЕ...

- » Определение искусственного интеллекта и его история
- » Использование искусственного интеллекта для практических целей
- » Заблуждения по поводу искусственного интеллекта
- » Объединение искусственного интеллекта с компьютерной технологией

В прошлом у *искусственного интеллекта* (Artificial Intelligence — AI) было несколько фальстартов, частично из-за того, что люди действительно не понимают, что такое искусственный интеллект и на что он способен. Главная проблема в том, что кинематограф, телевидение и книги тайно замыслили давать ложные надежды о способностях искусственного интеллекта. Кроме того, люди имеют привычку наделять технологии *человеческими качествами*, а это создает впечатление, что искусственный интеллект способен на большее, чем есть на самом деле. Таким образом, лучше всего начать эту книгу с определения, чем искусственный интеллект является фактически, чем он не является и как он связан с компьютерами на сегодняшний день.



ЗАПОМНИ

Конечно, ожидания от искусственного интеллекта основаны на комбинации того, как вы его себе представляете исходя из имеющихся для его реализации технологий, а также задач, которые вы перед ним ставите. Следовательно, все представляют себе искусственный интеллект по-разному. В этой книге искусственный интеллект рассматривается с нескольких точек зрения. Книга не впадает в крайности, поддаваясь ложным обещаниям сторонников или запугиваниям противников, чтобы вы получили наилучшее возможное представление об искусственном интеллекте как о технологии. В результате вы можете найти, что ваши ожидания несколько отличаются от изложенного в этой книге, доходчиво демонстрирующей то, что технология фактически может сделать для вас, и не надеяться на что-то невозможное.

Определение термина “искусственный интеллект”

Прежде чем использовать термин любым осмысленным способом, необходимо иметь его определение. В конце концов, если не договориться о значении термина, он останется только набором букв. Определение идиомы (термина, смысл которого неоднозначно понятен из смысла составляющих его элементов) особенно важно для технических терминов, получивших более чем широкое освещение в прессе неоднократно и разными способами.



ЗАПОМНИ

Объяснение, что ИИ — это искусственный интеллект, в действительности не говорит ничего осмысленного, что является причиной многих разногласий и жарких обсуждений этого термина. Да, вы можете утверждать, что он искусственного происхождения, т.е. не имеет естественного первоисточника. Но в части интеллекта все в лучшем случае неоднозначно. Хотя вы и не обязаны соглашаться с определением искусственного интеллекта, данным в следующем разделе, в этой книге данный термин используется согласно этому определению, и его знание облегчает понимание остальной части текста.

Распознавание интеллекта

Люди определяют интеллект по-разному. Но вполне можно сказать, что интеллект подразумевает определенную умственную деятельность, состоящую из следующих действий.

- » **Обучение.** Способность получать и обрабатывать новую информацию.
- » **Рассуждение.** Способность манипулировать информацией разными способами.

- » **Понимание.** Осознание результата манипуляции информацией.
- » **Понятие правды.** Определение достоверности используемой информации.
- » **Учет взаимоотношений.** Прогноз взаимодействия достоверных данных с другими данными.
- » **Учет смысла.** Применение фактов к конкретной ситуации способом, совместимым с их взаимоотношениями.
- » **Отделение факта от веры.** Определение, подтверждаются ли данные достоверными источниками, регулярно демонстрирующими свою правдивость.

Этот список легко может стать весьма длинным, но даже в таком виде он по-разному интерпретируется теми, кто принимает его истинность. Однако, как можно заметить из этого списка, интеллект зачастую использует процесс, которому компьютерная система вполне может подражать в ходе симуляции.

1. Установить цель на основании необходимости или желания.
2. Оценить значение любой известной на настоящее время информации, имеющей отношение к цели.
3. Собрать дополнительную информацию, способную иметь отношение к цели.
4. Обработать данные так, чтобы они приняли форму, совместимую с существующей информацией.
5. Установить истинность фактов и отношения между существующей и новой информацией.
6. Определить достижимость цели.
7. Сменить цель в свете новых данных и их влияния на вероятность успеха.
8. При необходимости повторять пп. 2–7 до тех пор, пока не будет найден путь достижения цели (успех) или пока возможные пути не исчерпаются (неудача).



ЗАПОМНИ

Несмотря на возможность создания алгоритмов и предоставления доступа к данным, необходимым для поддержки этого процесса в пределах компьютера, возможность компьютера реализовать интеллект жестко ограничена. Например, компьютер не способен к пониманию чего бы то ни было, поскольку он полагается только на машинные процессы манипулирования данными, используя чистую математику строго механическим способом. Аналогично компьютеры не способны легко отличать правду от лжи (как описано в главе 2). Фактически никакой компьютер не может полностью реализовать ни одно из умственных действий, приведенных в списке описания интеллекта.

В процессе обсуждения того, что же фактически представляет собой интеллект, будет полезна категоризация интеллекта. Решая задачи, люди не используют только один тип интеллекта, они скорее полагаются на несколько типов. Говард Гарднер (Howard Gardner) из Гарварда определил несколько типов интеллекта (см. подробности на <http://www.pz.harvard.edu/projects/multiple-intelligences>), их знание поможет связать их типы с видами задач, которые компьютер способен моделировать как интеллект. Модифицированная версия списка с дополнительным описанием приведена в табл. 1.1.

Таблица 1.1. Понятие видов интеллекта

Тип	Возможность симуляции	Человеческие инструменты	Описание
Визуально-пространственный	Средняя	Модели, графика, диаграммы, фотографии, рисунки, трехмерные модели, видео, телевидение и средства массовой информации	Интеллект физической среды, используется такими людьми, как моряки и архитекторы (и многими другими). Чтобы двигаться вообще, люди нуждаются в понимании своей физической среды, включая ее размерности и характеристики. Эта способность нужна интеллекту каждого робота или ноутбука, но ее зачастую трудно моделировать (как с самоуправляемыми автомобилями) или сделать достаточно точной (как с пылесосом, который, обходя препятствия, создает впечатление осмысленного движения)
Телесно-кинестетический	От средней до высокой	Специализированное оборудование и реальные объекты	Движения тела, как у хирурга или балерины, требуют точности и понимания своего тела. Роботы обычно используют этот вид интеллекта для выполнения повторяющихся задач, зачастую куда точнее, чем люди, но иногда с меньшим изяществом. Следует делать различие между средством усиления возможностей человека, таким как хирургическое устройство, улучшающее физические способности хирурга, и истинным независимым движением. Выше сказанное — это просто демонстрация математической возможности, в которой все зависит от действий хирурга

Тип	Возможность симуляции	Человеческие инструменты	Описание
Творческий	Нет	Художественная деятельность, новый образ мыслей, изобретения, новые виды музыкальных произведений	Творчество — это действие по выработке нового образа мыслей, приводящее к уникальному результату в форме произведения искусства, музыки или литературы. Действительно, новое произведение — это результат творчества. Искусственный интеллект способен моделировать существующий образ мыслей и даже комбинировать несколько, чтобы создать уникальное представление того, что уже существует, но в действительности является только математически ориентированной версией существующего образа. Для творчества искусственный интеллект должен обладать самосознанием, которое требует внутриличностного интеллекта
Межличностный	От низкой до средней	Телефон, звуковая конференц-связь, видеоконференц-связь, записи, компьютерная конференц-связь, электронная почта	Взаимодействие с другими происходит на нескольких уровнях. Задача этой формы интеллекта — получать, предоставлять, обмениваться и манипулировать информацией на основании опыта других. Компьютеры способны отвечать на простые вопросы на основании ввода ключевых слов, а не потому, что они понимают вопрос. Интеллект действует, получая информацию, находя подходящие ключевые слова и предоставляя затем информацию на основании этих ключевых слов. Поиск терминов в таблице с перекрестными ссылками и последующими действиями согласно предоставляемой таблицей инструкциям демонстрирует логический интеллект, а не межличностный
Внутриличностный	Нет	Книги, творческие материалы, дневники, уединение и время	Взгляд внутрь себя, чтобы понять собственные интересы, и последующая выработка цели на основании этих интересов являются

Тип	Возможность симуляции	Человеческие инструменты	Описание
Лингвистический	Низкая	Игры, средства массовой информации, книги, устройства звукозаписи и произносимые слова	в настоящее время прерогативой только человеческого интеллекта. Как у машин у компьютеров нет никаких нужд, желаний, интересов или творческих способностей. Искусственный интеллект обрабатывает введенные числа, используя набор алгоритмов, и предоставляет вывод, но он не знает ни о том, что он делает, ни о том, зачем он это делает
			Работа со словами — это немало-важный инструмент общения, поскольку разговор и обмен письменной информацией происходит намного быстрее, чем в любой другой форме. Эта форма интеллекта подразумевает понимание речевого и письменного ввода, его обработку для составления ответа и предоставление результата в качестве понятного ответа. Во многих случаях компьютеры могут лишь анализировать ввод по ключевым словам, но вопрос фактически не могут понять вообще. Кроме того, предоставляемый ответ также может оказаться непонятным. У людей за разговорный и письменный лингвистический интеллект отвечают разные области мозга (http://releases.jhu.edu/2015/05/05/say-what-how-the-brain-separates-our-ability-to-talk-and-write/), а значит, обладатели высокого письменного лингвистического интеллекта могут не иметь столь же высокого разговорного лингвистического интеллекта. Компьютеры в настоящее время не разделяют письменную и разговорную лингвистические способности

Тип	Возможность симуляции	Человеческие инструменты	Описание
Логико-математический	Высокая	Логические игры, исследование, тайны и головоломки	Вычисление результата, осуществление сравнений, исследование шаблонов и выявление взаимозависимостей — все это области, в которых компьютеры сейчас превосходны. Когда вы видите победу компьютера над человеком на телевикторине, это практически единственная общеизвестная форма интеллекта из семи. Да, вы могли бы замечать небольшие фрагменты других видов интеллекта, но в фокусе этот. Оценка различий интеллекта человека и компьютера на базе только одной области — это отнюдь не лучшая идея

Выработка четырех способов определения искусственного интеллекта

Как упоминалось в предыдущем разделе, первая концепция, которую важно осознать, — это то, что искусственный интеллект действительно не имеет никакого отношения к человеческому интеллекту. Да, некоторые модели искусственного интеллекта созданы для симуляции человеческого интеллекта, но это только симуляция. Размышляя об искусственном интеллекте, обратите внимание на взаимосвязь между поиском цели, обработкой данных, используемой для достижения этой цели, и сбором данных для лучшего понимания цели. Чтобы достичь результата, искусственный интеллект полагается на алгоритмы, и результат может иметь некое отношение к человеческим задачам или методам их решения, а может и не иметь. С учетом этого искусственный интеллект можно классифицировать четырьмя способами.

- » **Действующий по-человечески.** Когда компьютер действует, как человек, он лучше проходит тест Тьюринга, при котором компьютер по возможности имитирует человека (см. детали на <http://www.turing.org.uk/scrapbook/test.html>). К этой категории относятся все, что средства массовой информации выдают за искусственный интеллект. Такие технологии используются для обработки текстов на естественном языке, для представления знаний,

автоматизации формулирования логических выводов и машинного обучения (все четыре способны проходить тест). Первоначальный тест Тьюринга не подразумевал физического контакта. Более новый, полный тест Тьюринга действительно подразумевает физический контакт в форме перцепционной способности взаимодействия, а значит, для успеха компьютеру требуются машинное зрение и робототехника. Современные технологии склоняются к идее достижения цели, а не полного подражания людям. Например, братья Райт (Wright Brothers) добились успеха в создании самолета не точным копированием полета птиц; птицы, скорее, подсказали идеи, которые привели к созданию аэродинамики, которая в конечном счете привела к полету человека. Цель в том, чтобы летать. И птицы, и люди решают одинаковую задачу, но используют разные подходы.

» **Думающий по-человечески.** Когда компьютер думает, как человек, он решает задачи, требующие интеллекта (а не механических процедур), схожего с человеческим, такие как вождение автомобиля. Чтобы удостовериться, думает ли программа, как человек, необходим некоторый метод выяснения того, как думают люди, определив подход когнитивного моделирования. Эта модель полагается на три методики.

- **Самоанализ.** Выявление и документирование методик, используемых для достижения цели при контроле собственных мыслительных процессов.
- **Психологическая проверка.** Наблюдение за поведением человека и внесение результатов в базу данных подобных поведения других людей при подобном стечении обстоятельств, задаче, ресурсах и условиях окружающей среды (помимо всего прочего).
- **Нейровизуализация.** Контроль мозговой деятельности непосредственно, с использованием различных аппаратных средств, таких как компьютерная томография (Computerized Axial Tomography — CAT), позитронно-эмиссионная томография (Positron Emission Tomography — PET), магнитно-резонансная томография (Magnetic Resonance Imaging — MRI) и магнитоэнцефалография (Magnetoencephalography — MEG).

После создания модели можно написать использующую ее программу. С учетом огромного разнообразия человеческих мыслительных процессов и трудностей их точного представления в программе результаты в лучшем случае носят экспериментальный характер. Эта категория, мышление по-человечески, зачастую используется в психологии и других областях, в которых моделируется человеческий мыслительный процесс для создания реалистичных симуляций.

Думающий рационально. Изучение процесса мышления людей с использованием некоего стандарта позволяет выработать правила, описывающие типичное человеческое поведение. Человека считают рациональным, когда он следует этим правилам поведения в пределах определенных уровней отклонения. Компьютер, который думает рационально, полагается на заранее заданные правила поведения для выработки направлений взаимодействия с окружающей средой на основании имеющихся данных. Цель этого подхода заключается в по возможности логическом решении проблем. Как правило, этот подход подразумевает создание базовой методики решения проблемы, которая будет затем модифицирована для фактического решения конкретной проблемы. Другими словами, решение проблемы в принципе зачастую отличается от ее решения на практике, но отправная точка все же нужна.

- » **Действующий рационально.** Изучение того, как люди действуют в неких ситуациях и при определенных условиях позволяет определить, какие методики эффективны. Компьютер, действующий рационально, полагается на заранее заданные действия при взаимодействии с окружающей средой на основании условий, внешних факторов и имеющихся данных. Как и при рациональном мышлении, рациональные действия зависят от решения в принципе, которое может и не оказаться полезным практически. Однако рациональное действие предоставляет базовую линию, относительно которой компьютер может начать договариваться об успешном завершении задачи.

ЧЕЛОВЕК ПРОТИВ РАЦИОНАЛЬНЫХ ПРОЦЕССОВ

Человеческие процессы отличаются от рациональных процессов. Процесс *рационален*, если на основании текущей информации он всегда делает правильные вещи с идеальной эффективностью. Короче говоря, если рациональные процессы руководствуются книгой, то они полагают книгу абсолютно правильной. Человеческие процессы задействуют инстинкт, интуицию и другие переменные, которые не обязательно учтут книгу и могут даже не рассматривать имеющиеся данные.

Например, рациональный способ управления автомобилем подразумевает неукоснительное следование законам. Однако дорожный трафик не рационален. Если вы будете точно следовать правилам дорожного движения, то закончите влипшим во что-нибудь, поскольку другие водители не следуют правилам абсолютно точно. Поэтому, чтобы быть успешным, беспилотный автомобиль должен действовать по-человечески, а не рационально.

Категории, используемые при определении искусственного интеллекта, предоставляют способ учета различных способов его применения. Некоторые из систем, используемых обычно для классификации искусственного интеллекта по типам, не очень точны и по большей части произвольны. Например, некоторые группы разделяют искусственный интеллект на сильный (обобщенный интеллект, способный адаптироваться к множеству ситуаций) и слабый (специфический интеллект, разработанный для хорошего выполнения конкретной специфической задачи). Проблема с сильным искусственным интеллектом в том, что он не справляется с задачами достаточно хорошо, в то время как слабый искусственный интеллект слишком специфичен, чтобы выполнять задачи независимо. Даже в этом случае, только при двух типах, классификация не решает своей задачи даже в общем смысле. Классификация из четырех типов, предлагаемая Арендом Хинтцем (Arend Hintze) (см. <http://theconversation.com/understanding-the-four-types-of-ai-from-reactive-robots-to-self-aware-beings-67616>), дает лучшее основание для понимания искусственного интеллекта.

- » **Реактивные машины.** Примерами реактивных машин являются компьютеры, выигрывающие у людей в шахматы или участвующие в телевикторинах. У них нет никакой памяти или опыта, чтобы обосновать решение. Для поиска каждого решения каждый раз они полагаются исключительно на вычислительную мощь и интеллектуальные алгоритмы. Это пример слабого искусственного интеллекта, используемого для специфической цели.
- » **Ограниченная память.** Беспилотный автомобиль или автономный робот не может позволить себе тратить время на поиск каждого решения с самого начала. Эти машины полагаются на небольшой объем памяти для хранения основанных на опыте знаний для различных ситуаций. Когда машина встречает ту же самую ситуацию, она может положиться на свой опыт и, сократив время реакции, предоставить больше ресурсов для принятия новых решений, которые еще не были приняты. Это пример нынешнего уровня сильного искусственного интеллекта.
- » **Теория сознания.** У машины, способной оценить и свои основные задачи, и потенциальные задачи других существ в том же самом окружении, есть своего рода понимание происходящего, что выполнимо до некоторой степени уже сегодня, но не в любой коммерческой форме. Однако, чтобы беспилотные автомобили стали действительно автономными, этот уровень искусственного интеллекта следует разработать полностью. Беспилотный автомобиль должен не только знать, что он должен прибыть из одной точки в другую, но и предугадывать потенциально противоречащие цели водителей вокруг себя и реагировать соответственно.

» **Самосознание.** Это разновидность искусственного интеллекта, который вы встречаете в кинофильмах. Однако он требует технологий, появление которых пока не предвидится, поскольку у такой машины должно быть и самоощущение, и самосознание. Кроме того, вместо простого предугадывания целей других на основании окружающей обстановки и реакции других существ этот тип машины был бы в состоянии предусмотреть намерения других на основании своего опыта.

История искусственного интеллекта

Предыдущие разделы этой главы помогли вам взглянуть на интеллект с человеческой точки зрения и увидеть, что современные компьютеры прискорбно неадекватны для моделирования такого интеллекта, не говоря уже о том, чтобы самим стать фактически интеллектуальными. Однако желание создать разумную машину (как идола в древние времена) так же старо, как и мир. Желание не быть одиноким во Вселенной, иметь нечто для общения без конфликтов с людьми весьма велико. Конечно, в одной книге нельзя рассмотреть всю историю человечества, поэтому в следующих разделах содержится очень краткий обзор истории современных попыток создания искусственного интеллекта.

Символическая логика в Дартмуте

Первые компьютеры были не более чем вычислительными устройствами. Они подражали человеческой способности манипулировать символами, чтобы выполнять простые математические задачи, такие как сложение. Впоследствии вполне резонно потребовалось добавить возможность выполнения математических рассуждений за счет сравнения (такого, как не больше ли одно значение другого). Тем не менее, чтобы выполнить вычисление, люди все еще должны были определить используемый алгоритм, предоставить необходимые данные в правильной форме, а затем интерпретировать результат. Летом 1956 года на семинаре в кампусе Дартмутского колледжа группа ученых попробовала сделать нечто большее. Они предположили, что машины, способные рассуждать так же эффективно, как люди, потребуют максимум появления следующего поколения. Они ошиблись. Только сейчас мы получили машины, способные выполнять математические и логические рассуждения так же эффективно, как и человек (а это значит, что компьютеры должны овладеть по крайней мере еще шестью типами интеллекта, прежде чем хоть как-то приблизятся к человеческому интеллекту).

Проблема, определенная Дартмутским колледжем и другими исследователями того времени, была связана с аппаратными средствами, а именно — с их

возможностью выполнять вычисления достаточно быстро, чтобы создать симуляцию. Но это еще не вся проблема. Да, аппаратные средства, конечно, тоже важны, но вы не можете моделировать процессы, которых не понимаете. Впрочем, даже в этом случае причиной высокой эффективности искусственного интеллекта в настоящем является то, что аппаратные средства, наконец, стали достаточно мощными, чтобы поддерживать необходимые объемы вычислений.



ВНИМАНИЕ!

Самая большая проблема первых попыток (и все еще значительная проблема нынешних) заключается в том, что мы не понимаем, как люди думают достаточно хорошо, чтобы создать хоть какую-нибудь симуляцию, если даже предположить, что такая симуляция вообще возможна. Вспомните упомянутые ранее проблемы управляемого полета: братья Райт добились успеха, не моделируя полет птиц, а скорее, выработав понимание процессов, используемых птицами, создав, таким образом, науку об аэродинамике¹. Следовательно, если некто говорит вам, что следующее большое усовершенствование искусственного интеллекта уже прямо за углом, а никаких конкретных научных диссертаций по задействуемым им процессам у него нет, то и прямо за углом ничего нет.

Экспертные системы

Экспертные системы появились в 1970-х годах, а в 1980-х с их помощью попытались снизить вычислительные требования к искусственному интеллекту за счет использования знаний экспертов. Создавалось множество вариантов экспертных систем, включая основанные на правилах (они используют операторы `if...then` для принятия решения на базе эмпирических правил), основанные на фреймах (они используют базы данных, организованные во взаимосвязанные иерархии обобщенной информации — фреймы) и основанные на логике (они полагаются на теорию множеств для установления отношений). Появление экспертных систем важно уже потому, что они представляют собой первые действительно полезные и успешные реализации искусственного интеллекта.



СОВЕТ

Экспертные системы используются еще до сих пор (хотя они так больше и не называются). Например, системы проверки правописания и поиска грамматических ошибок в вашем приложении — это разновидности экспертных систем. В частности, система проверки грамматики является примером экспертной системы на основании правил. Если оглядеться внимательнее, то можно найти и другие места, где экспертные системы все еще используются на практике в привычных приложениях.

¹ См. <https://ru.wikipedia.org/wiki/Газодинамика>. — *Примеч. ред.*

Проблема экспертных систем в том, что их иногда трудно создавать и поддерживать. Первые их пользователи должны были изучить специализированные языки программирования, такие как List Processing (LisP) или Prolog. Некоторые производители увидели возможность передать экспертные системы в руки менее опытных или начинающих программистов, используя такие продукты, как VPExpert (см. <https://www.amazon.com/exec/obidos/ASIN/155622057X/datanetprod00f-20/>), использовавшие подход на основании правил. Но эти продукты предоставляли чрезвычайно ограниченные функциональные возможности при использовании небольших баз знаний.

В 1990-х годах термин *экспертная система* начал исчезать. Сложилось мнение, что экспертные системы были ошибкой, но действительность такова, что на самом деле экспертные системы были настолько успешны, что прочно укоренились в составе приложений, для поддержки которых они и были разработаны. Используем для примера текстовый процессор Word. Когда-то вы должны были покупать отдельное приложение для проверки правописания, такое как RightWriter (<http://www.right-writer.com/>). Однако теперь Word имеет встроенную систему проверки грамматики, поскольку она доказала свою полезность (хотя и не абсолютную корректность) (см. подробнее на <https://www.washingtonpost.com/archive/opinions/1990/04/29/hello-mr-chips-pcs-learn-english/6487ce8a-18df-4bb8-b53f-62840585e49d/>).

Преодоление зим искусственного интеллекта

Термин *зима искусственного интеллекта* (AI winter) описывает период сокращения финансирования разработки искусственного интеллекта. Вообще, искусственный интеллект следовал пути, на котором сторонники преувеличивали его возможности, побуждая людей с деньгами, но абсолютно не знающих технологию, делать инвестиции. Затем был период критики, когда искусственный интеллект обманывал ожидания, и как следствие происходило сокращение финансирования. За прошлые годы произошло несколько таких циклов, и все они были губительными для истинного прогресса.

В настоящее время искусственный интеллект вновь находится в раздутой фазе благодаря технологии *машинного обучения* (machine learning), позволяющей компьютерам обучаться на основе данных. Обучение компьютера на основе данных означает независимость от задающего действия (устанавливающего задачи) человеческого программиста, он получает их непосредственно из примеров, демонстрирующих компьютеру необходимое поведение. Это похоже на обучение ребенка, когда ему на примерах показывают, как себя вести. У машинного обучения есть проблемы, поскольку при небрежном обучении компьютер может научиться делать нечто неправильно.

Над алгоритмами машинного обучения работают пять научных школ, каждая со своей точкой зрения (см. подробности в разделе “Не обмануться в ожиданиях

от искусственного интеллекта” далее в этой главе). Наиболее успешное решение на настоящий момент — это *глубокое обучение* (deep learning), технология, стремящаяся подражать человеческому мозгу. Глубокое обучение стало возможным благодаря доступности мощных компьютеров, более умных алгоритмов, больших наборов данных, созданных в результате оцифровки нашего сообщества, и огромных инвестиций от таких фирм, как Google, Facebook, Amazon и других, использующих этот ренессанс искусственного интеллекта в интересах своих корпораций.

Все говорят, что благодаря глубокому обучению зима искусственного интеллекта закончена, и это пока правда. Однако, оглядываясь на нынешние представления людей об искусственном интеллекте, вполне можно предугадать следующую фазу критики, которая неизбежно начнется, если сторонники не снизят градус риторики. Искусственный интеллект способен на удивительные вещи, но это вполне приземленный вид удивительного, как описано в следующем разделе.

Области применения искусственного интеллекта

Сегодня искусственный интеллект используется в очень многих приложениях. Единственная проблема в том, что технология работает настолько хорошо, что вы даже не подозреваете о ее существовании. Фактически вы можете быть крайне удивлены, обнаружив, как много устройств в вашем доме уже использует искусственный интеллект. Например, некоторые интеллектуальные термостаты самостоятельно создают расписания на основании того, как вы регулируете температуру вручную. Аналогично голосовой ввод, используемый для управления некоторыми устройствами, способен обучаться вашей речи, чтобы лучше понимать вас. Искусственный интеллект, определенно, присутствует в вашем автомобиле и почти наверняка — на рабочем месте. Фактически он используется миллионами экземпляров и вполне безопасен, раз он не виден, хотя и весьма драматичен по природе.

Вот лишь несколько сфер, в которых вы могли бы увидеть использование искусственного интеллекта.

- » **Обнаружение мошенничества.** Вы получаете из банка запрос, платили ли вы своей кредитной карточкой за определенную покупку. Банк не любопытен; он просто обеспокоен тем фактом, что некто мог сделать покупку, используя вашу карточку. Искусственный интеллект банка обнаружил незнакомый шаблон расходов и предупредил кого следует.

- » **Планирование ресурсов.** Многим организациям необходимо эффективное планирование использования ресурсов. Например, больницы, вероятно, придется решать, куда поместить пациента, исходя из его потребностей, доступности квалифицированного персонала и ожидаемого периода пребывания пациента в больнице.
- » **Комплексный анализ.** Люди нередко нуждаются в комплексном анализе, когда приходится учитывать слишком много факторов. Например, один и тот же набор симптомов может свидетельствовать о нескольких проблемах. Чтобы спасти пациенту жизнь, врачу или другому специалисту может понадобиться помощь в своевременной постановке диагноза.
- » **Автоматизация.** Любая форма автоматизации может извлечь пользу из применения искусственного интеллекта для реакции на непредусмотренные изменения или события. Проблема некоторых типов автоматизации сегодня заключается в том, что непредусмотренное событие, такое как нахождение объекта в неполюженном месте, фактически может привести к ее отказу. Добавление искусственного интеллекта к автоматизации может позволить справиться с непредвиденным событием и продолжить работу, как будто ничего не случилось.
- » **Клиентская служба.** На другом конце линии клиентской службы, в которую вы сегодня звонили, вполне могло и не быть человека. Автоматизация сегодня достаточно хороша, чтобы, используя различные сценарии и ресурсы, справиться с подавляющим большинством вопросов. При корректных интонациях голоса (также предоставляемых искусственным интеллектом) вы не сможете даже с уверенностью сказать, что говорили с компьютером.
- » **Системы безопасности.** Большинство систем безопасности современных машин различных видов полагается на искусственный интеллект, чтобы перехватить управление транспортным средством в критической ситуации. Например, многие автоматические тормозные системы полагаются на искусственный интеллект, чтобы остановить автомобиль исходя из всех исходных данных, которые может предоставить транспортное средство, например направления заноса.
- » **Эффективность машин.** Искусственный интеллект может помочь контролировать машину так, чтобы добиться максимальной эффективности. Искусственный интеллект контролирует использование ресурсов, чтобы система не превышала скорости и для других задач. Каждая унция (28,3495 г) мощности будет использована точно так, как необходимо, чтобы оказать желаемую услугу.

Не обмануться в ожиданиях от искусственного интеллекта

В этой главе сообщается, что слухи об искусственном интеллекте более чем раздуты. К сожалению, здесь даже поверхностно не затрагивается вся существующая ложь. Если вы смотрели такие фильмы, как *Она* (Her) (<https://www.amazon.com/exec/obidos/ASIN/B00H9HZGQ0/datacservip0f-20/>) и *Из машины* (Ex Machina) (<https://www.amazon.com/exec/obidos/ASIN/B00XI057M0/datacservip0f-20/>), то у вас могло создаться впечатление, что искусственный интеллект способен на очень многое. Проблема в том, что искусственный интеллект фактически находится в младенчестве, и любой вид его применения, такого как показано в фильмах, является плодом хорошего воображения.

Возможно, вы слышали нечто о сингулярности, послужившей причиной громких заявлений в средствах массовой информации и кинофильмах. *Технологическая сингулярность* (singularity) — это, по существу, *верховный алгоритм* (master algorithm), охватывающий все пять научных школ в области машинного обучения. Но чтобы достичь того, о чем вещают все эти источники, машина должна быть способна учиться, как человек, а это требует наличия всех семи видов интеллекта, обсуждавшихся ранее, в разделе “Распознавание интеллекта”.

Вот эти пять школ.

- » **Символисты.** Исходной точкой являются логика и философия. При решении проблемы эта группа полагается на обратную дедукцию.
- » **Коннекционисты.** Исходной точкой является неврология. При решении проблемы эта группа полагается на обратное распространение ошибки.
- » **Эволюционисты.** Исходной точкой является эволюционная биология. При решении проблемы эта группа полагается на генетическое программирование.
- » **Байесианцы.** Исходной точкой является статистика. При решении проблемы эта группа полагается на статистический вывод.
- » **Аналоги.** Исходной точкой является психология. При решении проблемы эта группа полагается на многоядерные машины.

Главная задача машинного обучения заключается в объединении технологий и стратегий, представляемых пятью школами для создания единого алгоритма (*верховного алгоритма* — master algorithm), который позволяет что-нибудь изучать. Конечно, до достижения этой цели еще далеко. Но несмотря на это, такие ученые, как Педро Домингос (Pedro Domingos) (<http://homes.cs.washington.edu/~pedrod/>), работают над этой задачей уже сейчас.

Для ясности: эти пять школ не могут предоставить достаточно информации, чтобы фактически решить проблему человеческого интеллекта, поэтому создание верховного алгоритма для всех пяти школ все еще никак не может привести к сингулярности. На настоящий момент вас должно поражать только то, насколько люди ничего не знают о том, как они думают или почему они думают определенным образом. Все слухи об искусственном интеллекте, распространяющиеся по миру, являются либо надеждами людей, либо просто ложью.

Объединение искусственного интеллекта с базовым компьютером

Чтобы увидеть искусственный интеллект в действии, необходимо иметь своего рода вычислительную систему, содержащую необходимое программное обеспечение, и базу знаний. Вычислительной системой может быть нечто с микросхемой внутри; фактически смартфон делает то же самое, что и настольный компьютер, до некоторой степени. Конечно, если речь идет о такой компании, как Amazon, предоставляющей рекомендации о покупке различным людям, то смартфона окажется маловато; для такого приложения понадобится действительно большая вычислительная система. Размер вычислительной системы прямо пропорционален ожидаемому объему работ, выполняемому искусственным интеллектом.

Приложения могут также отличаться по размеру, сложности и даже расположению. Например, занимаясь бизнесом и желая проанализировать клиентские данные, чтобы решить, как лучше поднять продажи, вы могли бы положиться на серверное приложение. С другой стороны, если вы клиент и хотите искать на Amazon товары для покупки, приложение может даже не располагаться на вашем компьютере; вы обращаетесь к нему через веб-ориентированное приложение, размещенное на серверах компании Amazon.

Базы знаний также различаются по расположению и размеру. Чем сложнее данные, тем больше вы можете из них получить, но и манипулировать ими также понадобится больше. Когда дело доходит до управления знаниями, бесплатного обеда вы не получите. Взаимосвязь между расположением и временем также важна. Сетевое соединение позволяет обращаться к большему объему знаний, но требует больше времени из-за необходимости передачи данных по сетевым соединениям. Однако, хотя локальные базы данных работают быстрее, им, как правило, недостает детальности.



Глава 2

Определение роли данных

В ЭТОЙ ГЛАВЕ...

- » Данные как универсальный ресурс
- » Получение и обработка данных
- » Поиск недостоверностей в данных
- » Определение пределов сбора данных

В данных нет ничего нового. Они связаны с каждым серьезным приложением, когда-либо написанным для компьютера. Данные существуют во многих формах — некоторые организованы, некоторые нет. Меняется только их объем. У нас теперь есть доступ к таким большим объемам данных, что охвачен почти каждый аспект жизни большинства людей, причем иногда на таком уровне подробностей, который люди не могут себе даже представить. Некоторые находят это почти ужасающим. Кроме того, использование передовых аппаратных средств и усовершенствованных алгоритмов делает данные универсальным ресурсом для искусственного интеллекта.

Для работы с данными их нужно сначала получить. Сегодня приложения собирают данные вручную, как в прошлом, а также автоматически, используя новые методы. Но это не вопрос только одной или двух методик сбора данных; спектр этих методов распространяется от полностью ручного до полностью автоматического.

Необработанные данные обычно малопригодны для анализа. Эта глава поможет вам осознать необходимость в манипулировании данными и формировании данных в соответствии с определенными требованиями. Вы также убедитесь в необходимости определения истинного значения данных, чтобы гарантировать соответствие результатов анализа набору задач приложения.

Как ни странно, существуют пределы для сбора обрабатываемых данных. В настоящее время не существует технологий для чтения чьих-либо мыслей телепатическим способом. Конечно, есть и другие пределы, о большинстве из которых вам, вероятно, уже известно, но не все из них очевидны.

Поиск данных, вездесущих на данном этапе

Большая революция данных — это намного больше, чем просто громкие слова производителей, предлагающих новые способы хранения и анализа данных; это уже повседневная действительность и движущая сила нашего времени. Возможно, вы уже слышали о больших объемах данных, упоминаемых во многих специализированных научных и деловых публикациях, и даже задавались вопросом, что именно означает этот термин. С технической точки зрения термин *большие данные* (big data) относится к большим и сложным объемам компьютерных данных, столь больших и запутанных, что приложения не могут с ними справиться за счет применения дополнительного хранилища или увеличения мощности компьютера.

Большие данные подразумевают революцию в хранении данных и манипулировании ими. Она затрагивает такое то, чего можно достичь, используя данные уже в качественных терминах (т.е. можно не только выполнять большие задачи, но и выполнять их лучше). Компьютеры хранят большие данные в различных форматах, с человеческой точки зрения, но компьютер видит данные как поток единиц и нулей (базовый язык компьютеров). Вы можете просматривать данные тем или иным способом в зависимости от того, как вы их получаете и используете. Структура одних данных проста (вы знаете точно, что они содержат и где искать каждую их часть), тогда как другие данные структуры не имеют (у вас есть представление о том, что они содержат, но вы не знаете точно, как их упорядочить).

Типичные примеры структурированных данных — это таблицы баз данных, в которых информация упорядочена в столбцы, и каждый столбец содержит определенный тип информации. Как правило, данные структурируются в соответствии с проектом. Вы собираете их выборочно и записываете в соответствующем месте. Например, вы могли бы расположить данные о количестве людей, покупающих определенный товар в определенном столбце определенной

таблицы в определенной базе данных. Как и в библиотеке, если вы знаете, какие данные необходимы, вы можете найти их немедленно.

Неструктурированные данные состоят из изображений, видео и звуковых записей. Неструктурированную форму можно использовать и для текста, чтобы можно было отметить его такими характеристиками, как размер, дата или тип содержимого. Обычно вы не знаете точно, где находятся данные в неструктурированном наборе, поскольку данные представляются как последовательности единиц и нулей, которые приложение должно интерпретировать или визуализировать.



ЗАПОМНИ

Преобразование неструктурированных данных в структурированную форму может потребовать много времени и сил, а также задействовать многих людей. Большинство данных большой революции данных не структурировано и хранится как есть, пока некто не визуализирует их и не структурирует.

Хранилища этих огромных и сложных данных не возникли вдруг в одночасье. На разработку технологий хранения таких объемов данных потребовалось время. Кроме того, ушло время на распространение технологий создания и доставки данных, а именно — компьютеров, датчиков, интеллектуальных мобильных телефонов, Интернета и его служб Всемирной паутины. Следующие разделы помогут понять то, что делает данные универсальным ресурсом.

Понятие значений Мура

В 1965 году Гордон Мур (Gordon Moore), соучредитель компаний Intel и Fairchild Semiconductor, написал в статье “Cramming More Components Onto Integrated Circuits” (<http://ieeexplore.ieee.org/document/4785860/>), что на протяжении следующего десятилетия количество компонентов в интегральных схемах будет удваиваться каждый год. Тогда в электронике доминировали транзисторы. Способность встроить больше транзисторов в *интегральную схему* (Integrated Circuit — IC) означала, что электронные устройства могут стать полезнее, получив больше возможностей. Этот процесс, интеграция, подразумевает миниатюризацию электроники (та же микросхема способна на большее). Современные компьютеры не намного меньше, чем компьютеры десятилетие назад, но все они решительно более мощные. То же самое относится к мобильным телефонам. Хотя их размер практически тот же, что и у предшественников, они способны выполнять куда больше задач.

Объявленное Муром в этой статье фактически было истиной на протяжении многих лет. В полупроводниковой индустрии это известно как закон Гордона Мура (см. подробности на <http://www.moorelaw.org/>). Удвоение

действительно происходило на протяжении первых десяти лет, как и было предсказано. В 1975 году Мур подправил свое утверждение, предсказав удвоение каждые два года. На рис. 2.1 демонстрируется эффект этого удвоения. Такая скорость удвоения все еще имеет место, хотя теперь все едины во мнении, что она не продержится дольше, чем до конца текущего десятилетия (примерно до 2020 года). Начиная с 2012 года начало наблюдаться несоответствие между ожидаемым увеличением скорости и тем, чего полупроводниковым компаниям удастся достичь в области миниатюризации.

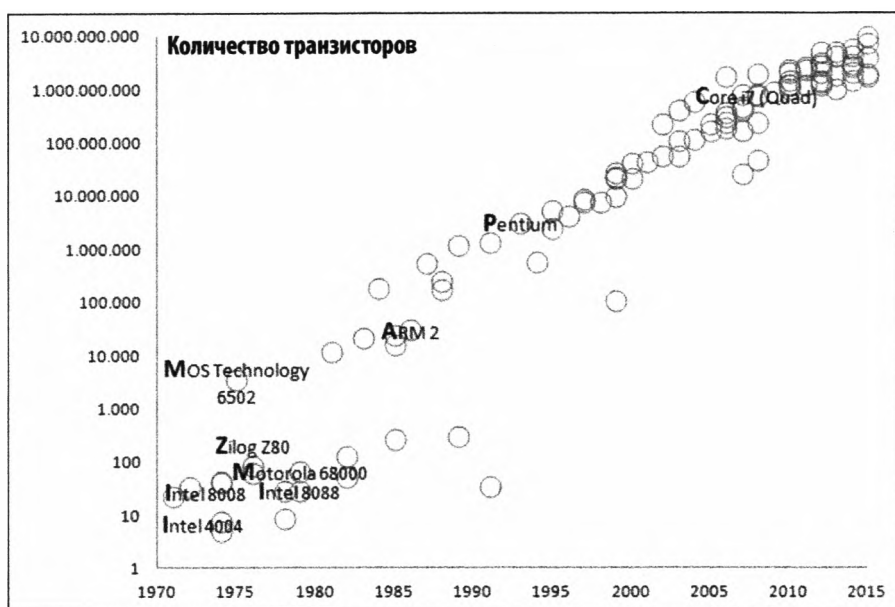


Рис. 2.1. Процессоры содержат все больше и больше транзисторов

Существуют физические препятствия интеграции большего количества элементов в микросхему при использовании существующих кремниевых технологий, поскольку размер компонентов уже практически предельно мал. Однако нововведения продолжают, как описано в статье <http://www.nature.com/news/thechips-are-down-for-moores-law-1.19338>. В будущем закон Гордона Мура может быть неприменим, поскольку промышленность перейдет на новые технологии (компоненты, например, могут создаваться на базе оптических лазеров, а не транзисторов; см. статью <http://www.extremetech.com/extreme/187746-by-2020-you-could-have-an-exascale-speed-of-light-optical-computer-on-your-desk> об оптическом компьютере). Тем не менее с 1965 года удвоение количества компонентов каждые два года возвестило большой скачок в развитии цифровой электроники, имевший далеко идущие последствия в сборе, хранении, манипулировании и управлении данными.

Закон Гордона Мура оказывает прямое влияние на данные. Начнем с более умных устройств. Чем умнее устройства, тем шире они распространены (как доказано электроникой, ставшей повсеместной). Чем шире становится распространенность, тем ниже становится цена — это бесконечный цикл, ведущий к повсеместному использованию мощных компьютеров и маленьких датчиков. Следствием доступности больших объемов машинной памяти и больших дисковых хранилищ для данных стало существенное увеличение доступности данных, включая веб-сайты, транзакционные записи, измерения, цифровые изображения и другие виды данных.

Данные используются повсюду

Для научных экспериментов ученые нуждаются в более мощных компьютерах, чем средний человек. Они имели дело с внушительными объемами данных за много лет до того, как некто выдумал термин “большие данные”. Тогда Интернет не предоставлял столь обширные объемы данных, как сегодня. Помните, что большие данные — не причуда, выдуманная производителями оборудования и программного обеспечения; они имеют основание во многих областях науки, таких как астрономия (космические миссии), спутники (наблюдение и контроль), метеорология, физика (ускорители частиц) и генетика (последовательности ДНК).

Хотя приложения искусственного интеллекта могут специализироваться в научной области (например, Watson от IBM имеет внушительные возможности в медицинской диагностике, поскольку способен изучать информацию из миллионов научных трудов о болезнях и медицине), у фактического утилитарного приложения искусственного интеллекта зачастую превалируют более мирские аспекты. Реальные приложения искусственного интеллекта ценят по большей части за то, что они способны распознавать объекты, находить пути или понимать то, что им говорят люди. Вклад данных в текущий ренессанс искусственного интеллекта поступил отнюдь не из классических источников научных данных.

Сегодня Интернет создает и распространяет новые данные в огромных количествах. На данный момент ежедневно создается приблизительно 2,5 квинтильона (число с 18-тью нулями) байтов данных, львиной долей которых являются такие неструктурированные данные, как видео и аудио. Все эти данные связаны с обычной человеческой деятельностью, чувствами, событиями и отношениями. Странствуя по этим данным, искусственный интеллект может легко учиться тому, как рассуждает и действует человек в различных ситуациях. Вот несколько примеров более интересных данных, которые можно найти.

- » Большое хранилище лиц и их выражений, полученных из фотографий и видео, опубликованных на социальных веб-сайтах, таких как Facebook, YouTube и Google. Предоставляет информацию о поле, возрасте, чувствах и, возможно, сексуальных предпочтениях, политической ориентации и показателе интеллекта (см. <https://www.theguardian.com/technology/2017/sep/12/artificial-intelligence-face-recognition-michal-kosinski>).
- » Частная медицинская информация и биопараметрические данные от смарт-часов, которые измеряют такие показатели, как температура и пульс, как во время болезни, так и в здоровом состоянии.
- » Наборы данных о взаимоотношениях между людьми и интересующих их вещах, полученные из таких источников, как социальные сети и поисковые механизмы. Например, исследование психометрического центра Кембриджского университета показало, что Facebook содержит множество данных о близких отношениях людей (см. <https://www.theguardian.com/technology/2015/jan/13/your-computer-knows-you-researchers-cambridge-stanford-university>).
- » Информация о том, как мы говорим, записывается мобильными телефонами. Например, функция OK Google на мобильных телефонах Android обычно записывает вопросы, а иногда и не только: <https://qz.com/526545/googles-been-quietly-recording-your-voice-heres-how-to-listen-to-and-delete-the-archive/>.

Каждый день пользователи подключают к Интернету все больше устройств, начинающих хранить новые личные данные. Теперь есть личные домашние помощники, такие как Amazon Echo и другие интегрированные интеллектуальные домашние устройства, позволяющие улучшать и регулировать домашнюю обстановку. Это только вершина айсберга, поскольку множество других бытовых приборов (от холодильников до зубных щеток) становятся взаимосвязанными и способными обрабатывать, записывать и передавать данные. *Интернет вещей* (Internet of Things — IoT) становится реальностью. Эксперты полагают, что к 2020 году подключенных вещей будет в шесть раз больше, чем людей, но исследовательские группы и исследовательские центры уже пересматривают эту картину (<http://www.gartner.com/newsroom/id/3165317>).

Приведение алгоритмов в действие

Человеческий род сейчас находится на невероятном перекрестке беспрецедентных объемов данных, создаваемых меньшими и все более мощными аппаратными средствами. Данные все так же обрабатываются и анализируются теми

же самыми компьютерами, что расширяет и углубляет процесс. Это утверждение может показаться очевидным, но данные стали настолько вездесущими, что их ценность больше не распространяется только на значение той информации, которую они содержат (как в случае данных, хранимых в базе компании, выполняющей свою повседневную работу); ценность, скорее всего, в ее использовании как средства создания новой информации; такие данные упоминаются как “новая нефть” (new oil). Это новое значение по большей части заключается в том, как приложения делают данным маникюр, хранят и возвращают их, а также в том, как вы фактически используете их с помощью интеллектуальных алгоритмов.

Алгоритмы и искусственный интеллект изменили игру с данными. Как упоминалось в предыдущей главе, для алгоритмов искусственного интеллекта опробовали различные подходы, от простых алгоритмов до символического рассуждения на основании логики, а затем перешли к экспертным системам. В последние годы стали увлекаться нейронными сетями, а впоследствии перешли к их более зрелой форме, глубокому обучению. Когда этот методологический пассаж произошел, данные превратились из информации, обрабатываемой предопределенными алгоритмами, в то, что формирует алгоритм, выполняющий нечто полезное для задачи. Данные превратились из просто сырья для принятия решения в самого автора решения, как показано на рис. 2.2.



Рис. 2.2. При существующих решениях для искусственного интеллекта больше данных означает больше интеллекта

Таким образом, фотография ваших котят становится куда как более полезной, и не просто из-за ее эмоционального значения (изображение ваших симпатичных котят), а потому, что может стать частью процесса обучения искусственного интеллекта, обнаруживая такие более общие концепции, как характеристики, определяющие кота, или помогая понять смысл определения “симпатичный”.

В значительно большем масштабе такие компании, как Google, кормят свои алгоритмы из источников общедоступных данных, таких как содержимое веб-сайтов или публично доступных текстов и книг. Программный паук Google ползает по Всемирной паутине, перепрыгивая с веб-сайта на веб-сайт, получая веб-страницы с содержащимся на них текстом и изображениями. Даже если Google возвращает часть этих данных пользователям в качестве результатов поиска, она извлекает из данных и другие виды информации, используя свои алгоритмы искусственного интеллекта, которые изучают ее для достижения других целей.

Алгоритмы обработки слов могут помочь системам искусственного интеллекта Google понять и предупредить ваши потребности, даже когда вы не выражаете их в наборе ключевых слов, а выражаетесь простым, невразумительным человеческим языком, на котором вы говорите каждый день (да, разговорный язык зачастую невразумителен). Если вы сейчас попытаете изложить вопрос, а не просто цепочку ключевых слов поисковому механизму Google, то заметите, что он склонен отвечать правильно. После 2012 года, с введением обновления Hummingbird (<http://searchengineland.com/google-hummingbird-172816>), система Google стала более приспособленной к пониманию синонимов и концепций, т.е. того, что обходит первоначально полученные данные, и это результат работы искусственного интеллекта. У Google есть и более передовые алгоритмы, например RankBrain, который непосредственно учится на многих миллионах запросов, поступающих каждый день, и может отвечать на неоднозначные или неясные поисковые запросы, даже выраженные на сленге или в разговорных терминах или просто написанных с ошибками. Алгоритм RankBrain обслуживает не все запросы, но обучается на основе данных тому, как лучше отвечать на запросы. Он уже обрабатывает 15 процентов запросов, а в будущем этот процент может достичь 100.

Успешное использование данных

Доступа к огромным объемам данных недостаточно для создания успешного искусственного интеллекта. Кроме того, алгоритм искусственного интеллекта не может извлечь информацию непосредственно из необработанных данных. Большинство алгоритмов полагается на внешнюю коллекцию и предварительную обработку перед анализом. Когда алгоритм собирает полезную информацию, он может и не представить то, что нужно. Следующие разделы помогут в общих чертах понять, как собрать и обработать данные, а также автоматизировать их сбор.

Источники данных

Используемые вами данные поступают из многих источников. Наиболее распространенный источник данных — это информация, введенная когда-то людьми. Даже когда система собирает данные сайта интернет-магазина автоматически, первоначально его информацию вводят люди. Человек щелкает на разных товарах, добавляет их в корзину покупателя, задает характеристики (например, размер) и количество, а затем просматривает. Затем, уже после покупки, человек дает оценку впечатлению от покупки, от товара и способа доставки, указывая их рейтинг, а также вводит комментарии. Короче говоря, каждое впечатление от покупки также становится источником для сбора данных.

Сегодня многие источники данных полагаются на сведения, собранные людьми. Люди также обеспечивают и их ввод вручную. Вы звоните или идете в некий офис, чтобы договориться о встрече со специалистом. Секретарь опрашивает вас, собирая информацию, необходимую для встречи. Эти вручную собранные данные в конечном счете лягут в основу набора данных, анализируемых где-нибудь в неких целях.

Данные собирают также сенсоры, и эти сенсоры могут быть почти любыми. Например, множество организаций специализируется на физическом сборе данных, таких как количество людей, смотревших на объект в витрине, или на обнаружении сотовых телефонов. Программное обеспечение распознавания лиц потенциально способно обнаруживать постоянных клиентов.

Сенсоры способны создавать наборы данных почти из чего угодно. Служба погоды полагается на наборы данных, созданные сенсорами, контролирующими погодные условия окружающей среды, такие как дождь, температура, влажность, облачный покров и т.д. Роботизированная система контроля позволяет исправить небольшие недостатки в работе робота за счет постоянного анализа собираемых сенсорами данных. Сенсор, объединенный с небольшим приложением искусственного интеллекта, мог бы подсказать вам, когда ваш обед будет совершенно готов сегодня вечером. Данные собирает сенсор, а приложение искусственного интеллекта использует соответствующие правила, чтобы точно определить, когда ваша еда будет готова.

Получение надежных данных

Определение термина *надежный* (reliable) кажется настолько же простым, насколько и трудным в реализации. Нечто считается надежным, если производимые им результаты ожидаемы и единообразны. Надежный источник данных предоставляет обыденные новости без неожиданностей, не потрясая никого сенсациями. В зависимости от вашей точки зрения это может быть и хорошо, когда люди смотрят новости не зевая, а потом не могут заснуть. Неожиданности

делают данные ценными для анализа и изучения. Следовательно, у данных есть аспект двойственности. Мы хотим надежных, обыденных, совершенно предсказуемых данных, которые просто подтверждают то, что мы уже знаем, но именно неожиданности делают сбор данных полезным.

Однако вам вовсе не понравятся совсем необычные данные, на первый взгляд пугающие. При получении данных необходим баланс. Данные должны находиться в определенных рамках (как упоминается далее, в разделе “Маникюр данных”). Они должны также соответствовать определенным критериям и по истинности значения (как описано далее, в разделе “Пять недостоверностей данных”). Данные должны также поступать в ожидаемые интервалы времени, и все поля поступающих записей должны быть заполнены.



ЗАПОМНИ

Защита данных также в некоторой степени влияет на их надежность. Целостность данных имеет несколько форм. Когда данные поступают, вы можете удостовериться в их принадлежности к ожидаемым диапазонам и соответствии конкретной форме. Но и после сохранения надежность данных может снизиться, если вы не гарантируете их хранение в надлежащей форме. Некто, поиграв с данными, способен повлиять на их надежность, сделав их подозрительными и даже потенциально непригодными для последующего анализа. Обеспечение надежности данных означает, что после поступления никто в них не вмешается и не приведет в соответствие с ожидаемыми пределами (сделав их в результате обыденными).

Повышение надежности ввода данных человеком

Люди делают ошибки — это свойство их натуры. Фактически неблагоприятно ожидать, что люди не будут делать ошибок. Все же проекты многих приложений предполагают, что люди так или иначе не будут делать никаких ошибок. Просто проект ожидает, что все будут следовать правилам. К сожалению, подавляющее большинство пользователей гарантированно не читает правил, поскольку они ленивы или им не хватает времени, когда дело доходит до вещей, которые не кажутся им действительно необходимыми.

Рассмотрим графу о штате в форме. Если вы предоставите просто текстовое поле, то некоторые пользователи введут название штата как, скажем, “Kansas” (Канзас). Конечно, некоторые пользователи сделают опечатку или ошибку в регистре букв, и получится “Kansus” или “kANSAS”. Кроме этих ошибок, у людей и организаций могут быть разные подходы к решению одних и тех же задач. Некто из издательской отрасли мог бы, руководствуясь стилем Ассошиэтед Пресс (Associated Press — AP), ввести “Kan”. Кто-то постарше, руководствуясь правилами Правительственной типографии США (Government

Printing Office — GPO), введет вместо этого “Kans.” Используются и другие сокращения. Почтовая служба США (U.S. Post Office — USPS) использует сокращение “KS”, а Береговая охрана США (U.S. Coast Guard) — “KA”. Тем временем Международная организация по стандартизации (International Standards Organization — ISO) ввела такую форму, как US-KS. Представьте, это только графа о штате, довольно простая, как можно было подумать до чтения этого раздела. Понятно, поскольку штаты не собираются менять свои названия часто, в форме вы можете просто предоставить раскрывающийся список, чтобы можно было выбрать штат в необходимом формате и устранить разночтения в использовании сокращений, избежать опечаток и ошибок регистра букв одним махом.



ЗАПОМНИ

Раскрывающиеся списки удивительно хорошо подходят для ввода самых разнообразных данных. Они гарантируют, что ввод данных человеком в эти поля будет чрезвычайно надежным, поскольку у человека нет другого выбора, кроме как использовать одно из стандартных значений. Конечно, человек вполне может выбрать неправильное значение. Вот здесь и играют свою роль перепроверки. Некоторые современные приложения сравнивают введенный почтовый индекс города и указанный штат, чтобы удостовериться в их соответствии. Если соответствия нет, пользователя просят повторить ввод. Эти перепроверки крайне раздражают, но пользователь вряд ли будет встречаться с ними часто, поэтому большой проблемой это стать не должно.

Даже при двойной проверке и статически заданных возможных значениях записи у людей все еще достаточно места для ошибок. Например, проблемой может стать ввод чисел. Когда пользователь должен ввести “2.00”, он может ввести “2” или “2.0”, или “2.”, или любой другой вариант. К счастью, анализ записи и ее переформатирование устранят проблему, и вы сможете решить эту задачу автоматически, без помощи пользователя.

К сожалению, переформатирование не исправит ошибочный числовой ввод. Частично вероятность таких ошибок можно снизить, применив контроль диапазона. Клиент не может купить –5 кусков мыла. Правильный способ позволить клиенту вернуть куски мыла — оформить это как возврат, а не как продажу. Однако пользователь может просто сделать ошибку, а вы можете предоставить сообщение, подсказывающее диапазон корректных вводимых значений.

Автоматизированный сбор данных

Некоторые полагают, что автоматизированный сбор данных решит все проблемы ввода данных человеком. Фактически автоматизированный сбор данных действительно имеет множество преимуществ.

- » Лучшая целостность
- » Повышенная надежность
- » Пониженная вероятность пропуска данных
- » Улучшенная точность
- » Пониженный разброс при периодическом получении исходных данных

К сожалению, нельзя сказать, что автоматизированный сбор данных решает все проблемы. Он все еще полагается на сенсоры, приложения и компьютерные аппаратные средства, разработанные людьми и предоставляющие доступ только к тем данным, которые выберут люди. В связи с ограничениями, налагаемыми людьми на характеристики автоматизированного сбора данных, его результат зачастую куда менее полезен, чем надеялись разработчики. Таким образом, автоматизированный сбор данных находится в постоянном процессе исканий, поскольку разработчики до сих пор пытаются решить все новые проблемы.

Автоматизированный сбор данных страдает также от недостатков оборудования и программного обеспечения, присущих любой вычислительной системе, но с более высокой вероятностью *мелких проблем* (soft issues), возникающих, когда система, безусловно, работает, но желательного результата не дает, чем проблем иного рода. Когда система работает, надежность ввода существенно превышает человеческие возможности. Но когда одолевают не острые проблемы, зачастую даже невозможно догадаться, что проблема в системе вообще существует, поскольку человек допустил бы больше ошибок. В результате набор данных окажется содержащим больше некачественных или даже неправильных данных, чем мог бы.

Маникюр данных

Когда люди говорят о данных, некоторые используют термин *манипуляция* (manipulation), создавая впечатление, что данные изменяются неким хитрым способом или шулерски подтасовываются. Возможно, лучшим окажется термин *маникюр* (manicuring), в результате которого данные получают хорошую форму и прелестный вид. Независимо от используемого термина, необработанные данные редко отвечают требованиям для анализа и обработки. Чтобы нечто извлечь из данных, им нужно сделать маникюр и привести в соответствие определенным требованиям. Следующие разделы посвящены задачам маникюра данных.

Как справиться с отсутствием данных

Чтобы ответить на заданный вопрос правильно, нужно иметь все факты. Вы можете предположить ответ на вопрос и без всех фактов, но этот ответ будет таким же вероятно правильным, как и неправильным. Зачастую некто принимает решение, по существу отвечает на вопрос, не имея всех фактов, а затем говорит, что пришел к такому заключению. Из-за недостающих данных в ходе анализа вы, вероятно, придете к большему количеству заключений, чем думаете. *Запись данных* (data record) — это один из элементов *набора данных* (dataset) (т.е. всех данных), состоящего из *полей* (field), которые содержат факты, используемые для ответа на вопрос. Каждое поле содержит один тип данных, соответствующий одному факту. Если это поле пусто, у вас нет необходимых данных для ответа на вопрос с использованием данной конкретной записи.



ЗАПОМНИ

Прежде чем начать решать проблему отсутствия данных, необходимо о ней узнать. Выяснение, что в вашем наборе данных отсутствует информация, может на практике оказаться весьма трудным, поскольку требует просмотра данных на низком уровне. Большинство людей не способны на это, но даже если у вас действительно есть необходимые навыки, то знайте, что эта работа весьма трудоемка. Зачастую первым признаком отсутствия данных являются нелепые ответы на ваши вопросы, выдаваемые алгоритмом и связанным с ним набором данных. Если используемый алгоритм правильный, значит, ошибка в наборе данных.

Если процесс сбора данных включает в набор не все данные, необходимые для ответа на некий вопрос, возможны проблемы. Фактически иногда выгоднее отбросить факт, чем использовать существенно искаженный факт. Если вы находите, что в некоем поле набора данных отсутствует 90 или более процентов его данных, то это поле, вероятно, бесполезно, и его необходимо удалить из набора данных (или найти способ получения всех недостающих данных).

У менее поврежденных полей данные могут отсутствовать по двум причинам. Случайно отсутствующие данные — это, как правило, результат ошибки сенсора или человека. Когда такое случается, в записях по всему набору данных есть отсутствующие элементы. Иногда причиной повреждений становится простой аппаратный сбой. Систематическое отсутствие данных происходит во время общего отказа некоторого типа. Пропадает весь сегмент записей в наборе данных, а значит, полученный результат анализа может быть существенно искажен.

Проще всего исправить случайное отсутствие данных. Как замену можно просто использовать среднее значение. Да, набор данных не будет абсолютно

точным, но работать он будет, вероятно, достаточно хорошо, чтобы получать вразумительные ответы. В некоторых случаях аналитики данных используют для вычисления недостающих значений специальный алгоритм, который позволяет сделать набор данных более точным.

Справиться с систематически отсутствующими данными значительно труднее, если вообще возможно, поскольку нет никаких окружающих данных, на основании которых можно было бы делать какие-либо предположения. Если вы можете найти причину отсутствия данных, то их иногда можно восстановить. Но если восстановление невозможно, такие поля лучше игнорировать. К сожалению, для некоторых ответов потребуются именно эти поля, а значит, проигнорировать, возможно, придется всю поврежденную последовательность записей данных, потенциально способную привести к неправильным выводам.

Учет рассогласования данных

Данные могут присутствовать в каждой записи набора данных, но они могут не согласоваться с данными в других имеющихся наборах. Например, числовые данные в поле одного набора данных могут иметь тип с плавающей точкой (с десятичной точкой), а в другом наборе данных — целочисленный тип. Прежде чем вы сможете объединить эти два набора данных, поля придется привести к одинаковому типу данных.

Возможны и другие виды рассогласований. Например, поля даты довольно часто имеют разные форматы. Для сравнения форматы дат следует сделать единообразными. Однако даты коварны своей склонностью выглядеть одинаково, но не быть таковыми. Например, даты в одном наборе данных могли бы использовать среднее время по Гринвичу (Greenwich Mean Time — GMT), а даты в другом — некий другой часовой пояс. Прежде чем вы сможете сравнивать время, его придется привести к единому часовому поясу. Все может стать куда интереснее, если даты в один набор данных поступают из мест, использующих летнее время, а в другой — из мест, его не использующих.

Даже если типы данных и формат совпадают, возможны и другие разновидности рассогласования. Например, поля в одном наборе данных могут не соответствовать полям в другом наборе данных. В некоторых случаях исправить эти различия довольно просто. В одном наборе данных первое и второе имена могут быть одним полем, а в другом — двумя отдельными полями (для первого и второго имен). Один из наборов данных придется изменить так, чтобы использовалось одно поле для обоих имен или два поля для первого и второго имен. К сожалению, распознать многие типы рассогласований в содержимом данных куда труднее. На практике зачастую оказывается совершенно невозможно понять, есть ли рассогласование вообще. Однако, прежде чем сдаваться, имеет смысл рассмотреть следующие потенциальные решения проблемы.

- » Вычислите отсутствующие данные из других доступных данных.
- » Найдите отсутствующие данные в другом наборе данных.
- » Объедините наборы данных так, чтобы получить один набор, содержащий единообразные поля.
- » Соберите дополнительные данные из разных источников, чтобы возместить отсутствующие.
- » Переопределите свой вопрос так, чтобы в отсутствующих данных не было больше нужды.

Отделение полезных данных от других

В некоторых организациях считается, что данных много не бывает, но избыток данных является такой же проблемой, как и их недостаток. Чтобы решать проблемы эффективно, искусственному интеллекту требуется только достаточное количество данных. При определении вопроса, на который вы хотите получить краткий полезный ответ, позаботьтесь об использовании правильного алгоритма (или набора алгоритмов). Конечно, главной проблемой наличия слишком больших объемов данных является то, что поиск решения (после перебора всей горы избыточных данных) отнимает слишком много времени, и иногда вы получаете невразумительные результаты, поскольку не можете видеть лес за деревьями.



ВНИМАНИЕ!

Создавая необходимый набор данных для анализа, делайте копию первоначальных данных, а не модифицируйте их сами. Всегда храните первоначальные необработанные данные неизменными, чтобы впоследствии их можно было использовать для другого анализа. Кроме того, для создания хорошего набора данных для анализа может потребоваться несколько попыток, поскольку может оказаться, что полученный результат не соответствует вашим ожиданиям. Главное, создать такой набор данных, который содержит только необходимое для анализа. Однако имейте в виду, что для получения желаемого результата могут понадобиться специфические виды усечения данных (выборка).

Пять недоверностей данных

Люди привыкли оценивать данные, выясняя, что они собой представляют, и вырабатывая мнение о них. Фактически люди иногда искажают данные до такой степени, что они становятся бесполезными, *недоверенными* (mistruth). Компьютер не способен различать правдивые и неправдивые данные; все, что он видит, — это данные. Одна из проблем, существенно затрудняющих, если

не делающих совсем невозможным, создание искусственного интеллекта, фактически думающего, как человек, заключается в том, что люди могут работать с недостоверными данными, а компьютеры — нет. Лучшее, чего можно надеяться достичь, — это выявить ошибочные данные по выбросам (outliers), а затем отфильтровать их, но данная методика не всегда решает проблему, поскольку использовать данные будет все еще человек и он будет пытаться установить факты на основании недостоверных данных.



ВНИМАНИЕ!

Общепринято мнение, что для создания менее замусоренных наборов данных не стоит позволять вводить данные людям, а по возможности полагаться на сбор данных сенсорами или другими средствами. К сожалению, сенсоры и другие механические технологии ввода отражают задачи разработавших их людей, и есть пределы того, что данная конкретная технология в состоянии обнаружить. Следовательно, даже данные, созданные машиной или сенсором, также склонны к недостоверностям, которые столь трудно обнаружить и преодолеть искусственному интеллекту.

В следующих разделах дорожно-транспортное происшествие используется в качестве базового примера для иллюстрации пяти типов недостоверностей, которые могут присутствовать в данных. Концепции для успешной попытки избежать происшествия (accident) не всегда могут присутствовать в данных, но могут и присутствовать, но не так, как нужно. Однако факт остается фактом: вам придется иметь дело с этим при просмотре данных.

Усердие

Недостоверность усердия (mistruth of commission) является результатом прямой попытки заменить правдоподобной информацией неправдоподобную. Например, при заполнении отчета о происшествии некто может заявить, что его на мгновение ослепило солнце, поэтому он не мог заметить пострадавшего. В действительности человека, возможно, отвлекло что-то другое либо он вообще думал вовсе не о дорожной обстановке (возможно, после хорошего обеда). Если никто не опровергнет эту теорию, человек может отделаться куда меньшим обвинением. Таким образом, данные могли бы быть искажены. В результате страховая компания будет принимать решение о выплате на основании ошибочных данных.



ЗАПОМНИ

Хотя может показаться, что недостоверности усердия легко преодолить, зачастую это не так. Человек говорит “мелкую безобидную ложь”, чтобы не затруднять других или справиться с проблемой наименьшими личными усилиями. Иногда причиной недостоверности

усердия является плохой почерк или слух. Фактически источников ошибок усердия (error of commission) так много, что действительно трудно придумать сценарий, в котором некто мог бы избежать их полностью. В результате недостоверностей этого типа иногда удается избежать, но, как правило, — нет.

Умолчание

Недостоверность умолчания (mistruth of omission) возникает, когда люди говорят правду по каждому запрошенному факту, но умалчивают важный факт, способный изменить восприятие происходящего в целом. Вернемся к отчету о происшествии. Скажем, некто сбивает оленя и наносит существенный урон своему автомобилю. Он правдиво свидетельствует, что дорога была мокрой, что приближались сумерки, что света было маловато, что он с опозданием нажал на тормоз, а олень просто выскочил из чащи прямо под колеса. Заключение: инцидент — это просто несчастный случай.

Однако этот человек умолчал об одном важном факте. В это время он писал SMS-сообщение. Если бы правоохранительные органы знали об SMS-сообщении, то это изменило бы причину происшествия на невнимательность при вождении. Водитель мог бы быть оштрафован, а страховая компания использовала бы другую причину при вводе инцидента в свою базу данных. Как и при недостоверности усердия, страховая компания, если не исправит ошибку в данных, будет принимать решение о выплате на основании ошибочных данных.



ЗАПОМНИ

Избежать недостоверностей умолчания почти невозможно. Да, кто-то может преднамеренно умолчать о факте в отчете, но не менее вероятно, что некто просто забыл включить в отчет все факты. В конце концов, большинство людей после происшествия сильно испуганы, в таком состоянии очень просто растеряться и сообщить только те факты, которые оставили наиболее значительное впечатление. Впоследствии человек может вспомнить дополнительные детали и сообщить их, но база данных все равно вряд ли будет содержать абсолютно полный набор фактов.

Точка зрения

Недостоверность точки зрения (mistruth of perspective) возникает, когда несколько сторон описывают происходящее с разных точек зрения. Рассмотрим пример дорожно-транспортного происшествия, когда был сбит пешеход. У человека, управлявшего автомобилем и совершившего наезд, и у свидетеля будут разные точки зрения на произошедшее. Офицер, берущий показания у каждого, понятное дело, получит от них разные факты, даже с учетом того, что все

говорят правду, известную именно ему. Фактически опыт доказывает, что записанное офицером в отчет почти всегда является средним из показаний обоих, с тенденциозным смещением согласно личному опыту. Другими словами, отчет будет близок к правде, но недостаточно близок для искусственного интеллекта.

Когда имеешь дело с точкой зрения, важно учитывать позицию наблюдателя. Водитель автомобиля мог видеть панель приборов и знать состояние автомобиля во время происшествия. Эта информация отсутствует у других сторон. Аналогично у человека, сбитого автомобилем, наилучшая позиция наблюдения, чтобы увидеть выражение лица водителя (его намерение). Свидетель мог бы находиться в наилучшей позиции, чтобы увидеть, пытался ли водитель затормозить или уклониться от столкновения и была ли у него такая возможность. Каждая сторона дает показания на основании замеченных данных, если не извлекает пользы из сокрытия фактов.



ВНИМАНИЕ!

Точка зрения является, вероятно, самой опасной из недостоверностей, поскольку любой попытавшийся выяснить правду в этом сценарии закончит в лучшем случае средним из всех показаний, которые никогда не будут полной правдой. Человек, просматривающий такую информацию, может полагаться на интуицию и инстинкт, чтобы как-то приблизиться к правде, но искусственный интеллект всегда будет использовать только среднее, а значит, всегда будет далек от истины. К сожалению, избежать недостоверности точки зрения невозможно, поскольку независимо от количества свидетелей события лучшее, на что можно рассчитывать, — это приближение к правде, а не абсолютная истина.

Есть и другая разновидность недостоверности точки зрения. Рассмотрим такой случай: вы — глухой человек в 1927 году. Каждую неделю вы идете в кинотеатр, чтобы посмотреть “немой фильм” и в течение часа или дольше чувствовать себя таким, как все. Вы воспринимаете фильм точно так же, как и все остальные, ничем от них не отличаясь. В октябре того же года вы видите объявление о модернизации кинотеатра звуковой системой, чтобы можно было показывать *звуковые фильмы* — фильмы со звуковой дорожкой. В объявлении сказано, что это прекрасная вещь, и почти все, казалось бы, с этим соглашались, кроме вас. Теперь глухой человек будет чувствовать себя второсортным гражданином, не таким, как все, и даже отлученным от кинематографа. В глазах глухого данное объявление является недостоверностью; добавление звука — это самое худшее, что могло произойти, а не самое лучшее. Дело в том, что очевидная истина не является истиной для всех. Идея всеобщей правды, т.е. истины для всех, является мифом. Ее не существует.

Предубежденность

Недостоверность предубежденности (mistruth of bias) происходит тогда, когда некто был в состоянии заметить факт, но упустил его из-за личных факторов или уверенности. В примере с ДТП водитель мог сосредоточить свое внимание на середине дороги, и олень на обочине стал невидимым. Следовательно, у водителя не было времени реагировать, когда олень внезапно выбежал на середину дороги.

Проблема с предубежденностью в том, что ее невероятно трудно классифицировать. Например, у водителя, не заметившего оленя, мог быть настоящий несчастный случай, олень мог быть скрыт за кустарником. Но водитель мог бы быть виновен и в невнимательном вождении. Водитель мог также отвлечься на мгновение. Короче говоря, тот факт, что водитель не видел оленя, не является вопросом; вопрос в том, почему он его не видел. Во многих случаях подтверждение предубежденности источника становится важным фактором при создании алгоритма, призванного избежать предубежденности источника.



ЗАПОМНИ

Теоретически избежать недостоверности предубежденности можно всегда. Но в действительности у всех людей есть предубеждения разных типов, и эти предубеждения всегда будут приводить к недостоверностям, искажающим наборы данных. Получить нечто, выглядящее, как факт, и фактом являющееся (т.е. сделать так, чтобы это запечатлелось в мозге человека), — довольно трудная задача. Люди фильтруют информацию, чтобы избежать информационной перегрузки, и эти фильтры — также источник предубеждения, поскольку они не позволяют людям видеть вещи реально.

Недопонимание

Система взглядов должна формироваться на основании только понимания, а не быть результатом ошибок из-за этих пяти недостоверностей. *Недостоверность недопонимания* (frame-of-reference mistruth) происходит, когда одна сторона описывает некое событие, а у второй стороны нет об этом никакого понимания. В результате подробности становятся запутанными или полностью непонятными. Многие комедии полагаются на ошибки недопонимания. Общеизвестный пример — эпизод *Кто на первой базе?* (Who's On First?) комедийной группы “Эбботт и Костелло”, доступный по адресу <https://www.youtube.com/watch?v=kTcRRaXV-fg>. Чтобы один человек понял то, что говорит другой, у первого человека должны быть опыт и знания по теме — система взглядов.

Другой пример недостоверности недопонимания — когда одна сторона не может понять другую. Например, у моряка есть опыт хождения в штормовом

море. Возможно, это просто муссон, но предположим на мгновение, что это серьезный шторм, возможно, опасный для жизни. Даже с помощью видео, интервью и симуляторов моряк не может передать опыт выживания в опасный для жизни шторм на море никому, кто не испытал такой шторм на себе; такой человек просто ничего не поймет.



ЗАПОМНИ

Наилучший способ избежать недостоверности недопонимания — обеспечить выработку всеми участвующими сторонами подобных систем взглядов. Для решения этой задачи, гарантированно точной передачи данных от одного человека к другому, все стороны должны иметь подобные опыт и знание. Ошибки недопонимания при работе с набором данных вполне возможны, когда у вероятного контролера нет необходимых опыта и знаний.

У искусственного интеллекта всегда будут проблемы с системой взглядов, поскольку у него отсутствует способность накопления опыта. Банк приобретенных знаний — это вовсе не то же самое, что и опыт. Банк данных содержал бы факты, но опыт на их основании — это не только факты, но и заключения, чего современные технологии повторить не могут.

Установление пределов сбора данных

Может показаться, что все собирают ваши данные безо всякой на то причины или задней мысли. И вы правы, так оно и есть. Фактически организации собирают, классифицируют и хранят всеобщие данные безо всякой видимой цели или намерения. Согласно сайту Data Never Sleeps (<https://www.domo.com/blog/data-never-sleeps-5/>), мир собирает данные со скоростью 2,5 квинтильона байтов в день. Эти ежедневные данные имеют разнообразные типы и формы, что подтверждают следующие примеры.

- » Google осуществляет 3 607 080 поисков.
- » Пользователи Twitter публикуют 456 000 сообщений.
- » Пользователи YouTube просматривают 4 146 600 видеофильмов.
- » Google Inbox получает 103 447 529 сообщений спама по электронной почте.
- » Weather Channel получает 18 055 555,56 запроса о погоде.
- » Giphy осуществляет 694 444 обмена файлами GIF.

Сбор данных стал наркотиком для организаций во всем мире, и некоторые думают, что организация, собравшая их больше всех, так или иначе получит

приз. Однако сбор данных сам по себе ничего не даст. Эта проблема во всей красе демонстрируется в книге Дугласа Адамса (Douglas Adams) *Автостопом по галактике* (*The Hitchhiker's Guide to the Galaxy*) (<https://www.amazon.com/exec/obidos/ASIN/1400052920/datacervip0f-20/>). В книге раса суперсозданий строит огромный компьютер, чтобы вычислить “Основной Ответ на Основной Вопрос”. Ответ 42 ничего не дает, поэтому создания жалуются, что сбор, классификация и анализ всех данных, использованных для ответа, не привел к пригодному для использования результату. Компьютер, без сомнений, разумный, и он дает людям действительно правильный ответ, но они должны знать, что сам вопрос для ответа имеет смысл. Сбор данных может происходить в неограниченных объемах, но сформулировать правильный вопрос, чтобы задать его, достаточно сложно, если вообще возможно.



ЗАПОМНИ

Основная проблема, которую должна решить любая организация по сбору данных, — какие вопросы задать и почему они важны. Организуя сбор данных для ответа на вопрос, следует позаботиться о том, чтобы ответ имел смысл. Например, если вы открываете магазин в городе, то, возможно, нуждаетесь в ответах на следующие вопросы.

- » Сколько людей проходит перед магазином каждый день?
- » Сколько людей останавливается и заглядывает в витрину?
- » Как долго они смотрят?
- » В какое время дня они смотрят?
- » Приводит ли смена содержимого в витрине к лучшим результатам?
- » Что из выставленного фактически заставляет людей заходить в магазин?

Список можно продолжить, но идея в том, что необходимо составить список вопросов, решающих конкретные задачи. После того как вы составили список, необходимо удостовериться, что каждый из вопросов фактически важен (т.е. решает конкретную задачу), а затем установить вид информации, необходимой для получения ответа на вопрос.



ВНИМАНИЕ

Конечно, попытка собрать все эти данные вручную была бы невозможна. Вот где играет свою роль автоматизация. Казалось бы, она обеспечивает надежный, воспроизводимый и единообразный ввод данных. Однако многие факторы автоматического сбора данных могут привести к сбору не особо полезных данных. Рассмотрим, к примеру, следующие проблемы.

- » Сенсоры могут собрать только те данные, для которых они разработаны, поэтому они могут пропустить данные, для которых они не предназначены.

- » Люди допускают ошибки в данных разными способами (см. раздел “Пять недостоверностей данных”), а значит, полученные данные могут быть ложными.
- » Данные могут быть искажены, когда условия для их сбора созданы неправильно.
- » Неправильная интерпретация данных также приводит к неправильным результатам.
- » Преобразование реального вопроса в понятный компьютеру алгоритм также является процессом, подверженным ошибкам.

Учесть необходимо и множество других проблем (хватит на целую книгу). Объединяя плохо собранные и плохо сформированные данные с алгоритмами, которые фактически не отвечают на ваши вопросы, вы получаете результаты, способные повести ваш бизнес в неправильном направлении, а затем искусственный интеллект обвинят в противоречивых или ненадежных ответах. Задайте правильный вопрос, предоставьте правильные данные, правильно их обработайте и проанализируйте — вот и все, что нужно сделать, чтобы сбор данных стал инструментом, на который можно полагаться.



Глава 3

Вопросы использования алгоритмов

В ЭТОЙ ГЛАВЕ...

- » Роль алгоритмов в искусственном интеллекте
- » Победа в играх с поиском в пространстве состояний и мини-максом
- » Анализ работы экспертных систем
- » Машинное обучение и глубокое обучение как часть искусственного интеллекта

Данные — это решающий фактор искусственного интеллекта. Недавние авансы по поводу искусственного интеллекта намекают, что для решения некоторых проблем правильный объем данных важнее правильного алгоритма. Например, в 2001 году два исследователя из Microsoft, Мишель Банко (Banko) и Эрик Брил (Brill), опубликовали незабываемую статью “Scaling to Very Very Large Corpora for Natural Language Disambiguation” (<http://www.aclweb.org/anthology/P01-1005>), в которой утверждали, что если вы хотите, чтобы компьютер создал модель языка, то вам не нужен самый умный алгоритм в городе. После ввода более чем одного миллиарда слов в пределах контекста любые алгоритмы начнут работать невероятно хорошо. Эта

глава поможет вам понять отношения между алгоритмами и данными, а также как они заставляют их выполнять полезную работу.

Однако, независимо от количества имеющихся данных, чтобы сделать их полезными, необходим алгоритм. Кроме того, чтобы данные правильно работали с выбранными алгоритмами, следует выполнить *анализ данных* (набор определенных этапов). Ничего пропускать здесь нельзя. Даже при том, что искусственный интеллект — это интеллектуальная автоматизация, иногда автоматизация должна держаться в тени анализа. Обучаемые машины находятся в далеком будущем. Сейчас вы не найдете машин, обладающих самодостаточностью и способных полностью избежать человеческого вмешательства. Вторая половина этой главы поможет понять роль экспертных систем, машинного обучения, глубокого обучения и таких приложений, как AlphaGo, для приближения будущих возможностей к современной действительности.

Понятие роли алгоритмов

Люди обычно распознают искусственный интеллект, когда инструмент представляет новый подход и взаимодействует с пользователем способом, подобным человеческому. Примерами являются такие цифровые помощники, как Сири (Siri), Алекса (Alexa) и Кортана (Cortana). Но некоторые другие общепринятые инструментальные средства, например навигатор GPS и специализированные планировщики (такие, как позволяющие избежать столкновений автомобилей, автопилоты в самолетах и составители рабочих планов), даже не выглядят, как искусственный интеллект, поскольку слишком распространены и считаются само собой разумеющимися, ведь они действуют на заднем плане приложений.

Это, безусловно, *эффект искусственного интеллекта* (AI effect), описанный Памелой Мак-Кордак (Pamela McCorduck), американской писательницей, автором известной истории искусственного интеллекта, в 1979 году. Эффект искусственного интеллекта заключается в том, что успешные интеллектуальные компьютерные программы быстро теряют благодарность людей и становятся тихими исполнителями, в то время как внимание смещается к тем проблемам искусственного интеллекта, которые все еще требуют решения. Люди перестают осознавать важность классических алгоритмов для искусственного интеллекта и начинают фантазировать об интеллекте, созданном по тайной технологии, или сравнивать его с последними футуристическими анонсами, таким как машинное и глубокое обучение.

Алгоритм (algorithm) — это процедура, представляющая собой последовательность операций (обычно выполняемых компьютером) и гарантирующая

нахождение правильного решения задачи за конечный промежуток времени или уведомляющая об отсутствии решения. Даже при том, что люди использовали алгоритмы вручную на протяжении буквально тысяч лет, в зависимости от сложности решаемой задачи они могут потребовать огромных периодов времени и множества числовых вычислений. Чем быстрее и проще алгоритмы находят решение, тем лучше. Алгоритмы являются продуктом интеллекта создавших их людей, поэтому любая работающая на алгоритмах машина не может не отражать человеческий интеллект, встроенный в такие алгоритмические процедуры.

Что означает алгоритм

Алгоритм всегда представляет последовательность этапов, но для решения проблемы не обязательно выполнять все этапы. Возможности алгоритмов невероятно велики. Операции могут задействовать хранение данных, их исследование и упорядочение в структуры данных. Вы можете найти алгоритмы, решающие задачи в науке, медицине, финансах, промышленном производстве и телекоммуникациях.

Все алгоритмы — это последовательности операций по нахождению правильного решения задачи за разумное время (или оповещению об отсутствии решения, если его нет). Алгоритмы искусственного интеллекта отличаются от обобщенных алгоритмов решения задач, являющихся обычно (или даже исключительно) результатом действия человеческого интеллекта. Алгоритмы искусственного интеллекта обычно имеют дело со сложными проблемами, зачастую являющимися частью NP-полной задачи (где NP — это недетерминированное полиномиальное время), а люди обычно имеют дело с комбинацией рационального подхода и интуиции. Вот лишь несколько примеров.

- » Планирование и распределение при недостатке ресурсов.
- » Поиск маршрутов в сложных физических или фигуративных пространствах.
- » Распознавание шаблонов в видимом изображении (не путать с восстановлением или обработкой изображения) или слышимом звуке.
- » Обработка языка (понимание текста и перевод на другой язык).
- » Игра (и победа) в соревновательных играх.



СОВЕТ

NP-полные задачи отличаются от других алгоритмических задач, поскольку поиск их решения за разумный период времени все еще невозможен. NP-полные задачи — это не те задачи, которые можно решить перебором всех возможностей и их комбинаций. Даже если бы были компьютеры, куда более мощные, чем современные, поиск

решения затянулся бы почти навсегда. Подобной проблемой в области искусственного интеллекта является *полный искусственный интеллект* (AI-complete).

Планирование и ветвление

Планирование помогает определить последовательность действий для достижения определенной цели. Это классическая задача искусственного интеллекта, ее примеры можно найти при планировании в промышленном производстве, распределении ресурсов и перемещениях робота в помещении. Искусственный интеллект определяет все возможные действия начиная с текущего состояния. Технически он *развивает* (expand) текущее состояние во многие будущие состояния. Затем он развивает все будущие состояния уже в их собственные будущие состояния и т.д. Когда дальнейшее развитие становится невозможным, искусственный интеллект останавливает процесс, создав *пространство состояний* (state space), состоящее из того, что могло бы случиться в будущем. Искусственный интеллект может использовать пространство состояний совсем не так, как возможный прогноз (фактически он прогнозирует все, хотя одни будущие состояния более вероятны, чем другие). Он может использовать пространство состояний для исследования решений, которые он может принять, чтобы достигнуть своей цели наилучшим способом. Это *поиск в пространстве состояний* (state-space search).

Работа с пространством состояний требует использования и специфических структур данных и алгоритмов. Общепринятыми структурами для базовых данных являются деревья и графы. Для эффективного исследования графов применяются такие популярные алгоритмы, как *поиск в ширину* (breadth-first search) и *поиск в глубину* (deep-first search).

Дерево данных очень похоже на настоящее дерево. Каждый добавляемый элемент является *узлом* (node). Узлы соединяются между собой связями. Комбинация узлов и связей формирует структуру, выглядящую, как дерево (рис. 3.1).



ЗАПОМНИ

Деревья имеют один корневой узел, точно так же, как и физическое дерево. *Корневой узел* (root node) — это отправная точка для выполняемой обработки. С корнем соединены ветви или листья. *Узел листа* (leaf node) — это конечная точка дерева. *Узлы ветвей* (branch node) поддерживают другие ветви или листья. Дерево на рис. 3.1 является бинарным, поскольку у каждого узла есть по меньшей мере два соединения (но у деревьев, представляющих пространства состояний, может быть какое угодно количество ветвей).

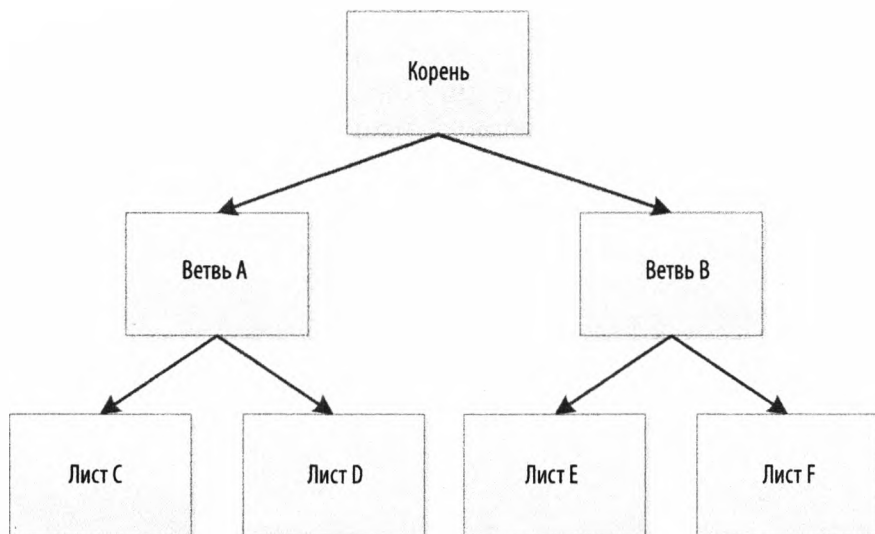


Рис. 3.1. Дерево может выглядеть как его физический прототип, а может располагаться корнем вверх

Если рассмотреть дерево, то ветвь В является *потомком* (child) корневого узла. Поэтому корневой узел является первым в списке. Листья Е и F оба являются потомками ветви В, что делает ветвь В *предком* (parent) листьев Е и F. Отношения между узлами важны, поскольку обсуждения деревьев нередко рассматривают родительски/дочерние отношения между узлами. Без этих терминов обсуждение деревьев может стать весьма невразумительным.

Граф (graph) — это своего рода модификация дерева. Подобно деревьям, здесь есть соединенные между собой узлы, формирующие отношения. Однако в отличие от бинарных деревьев, у узла графа может быть больше одной или двух связей. Фактически у узлов графа обычно есть множество соединений, а самые важные узлы могут иметь связи в любом направлении, а не только от предка к потомку. В простейшем виде граф выглядит как на рис. 3.2.

Графы — это структуры, представляющие множество узлов (или *вершин* (vertex)), соединенных множеством *ребер* (edge) или *дуг* (arc) в зависимости от представления. Граф можно считать структурой, подобной карте, на которой каждое место является узлом, а улицы — ребрами. Это отличается от дерева, на котором каждый путь завершается листом. Граф представлен на рис. 3.2. Графы особенно полезны при изучении состояний, представляющих своего рода физическое пространство. Например, системы GPS используют граф для представления мест и улиц.

Графы имеют также несколько новых подвохов, которые вы, возможно, не заметили. Например, граф может иметь концепцию *направленности*

(directionality). В отличие от дерева, у которого есть родительские/дочерние отношения, узел графа может соединяться с любым другим узлом в определенном направлении. Считайте их улицами в городе. Большинство улиц — с двухсторонним движением, но некоторые — с односторонним, по которым можно двигаться только в одном направлении.

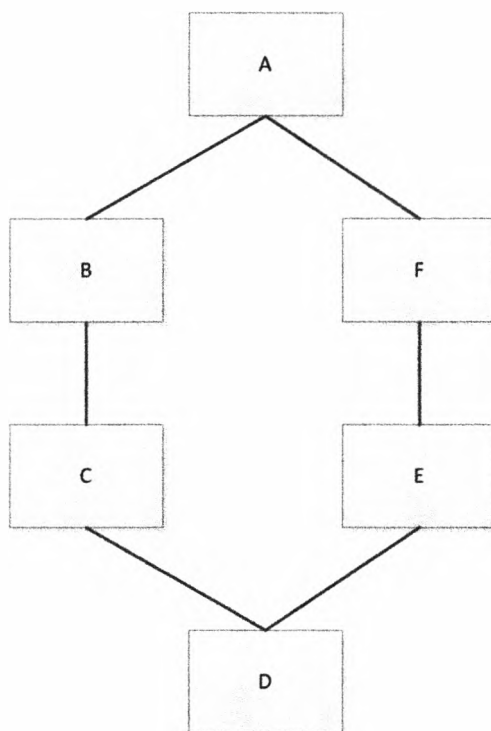


Рис. 3.2. Узлы графа могут соединяться между собой по-разному

Представление соединений графа может и не отражать реалии графа. Граф может определять коэффициенты для соединений. Коэффициент может определять дистанцию между двумя пунктами, время для прохождения маршрута или другие виды информации.



СОВЕТ

Дерево — это не более чем граф, в котором любые две вершины соединяются только одной связью, и это не допускает циклов (чтобы можно было вернуться к предку от любого потомка). Многие алгоритмы графов применимы только к деревьям.

Пересечение графа означает поиск (посещение) каждой вершины (узла) в определенном порядке. Процесс посещения вершины может включать его чтение и модификацию. По мере пересечения графа вы обнаруживаете вершины,

которые еще не были посещены. Вершина становится обнаруженной (поскольку вы ее только что посетили) или обработанной (поскольку алгоритм опробовал все следующие из нее ребра) после поиска. Порядок поиска определяет его вид: неосведомленный (поиск вслепую) и осведомленный (эвристика). При *неосведомленной* (uninformed) стратегии искусственный интеллект исследует пространство состояний без дополнительной информации, кроме структуры графа, которую он выясняет по мере его пересечения. В следующих разделах рассматриваются два популярных алгоритма поиска вслепую: поиск в ширину и поиск в глубину.

Поиск в ширину (Breadth-First Search — BFS) начинается с корня графа и исследует каждый присоединенный к нему узел. Затем поиск переходит на следующий уровень и далее по очереди, пока не будет достигнут конец. Следовательно, в примере графа поиск распространится от вершины А к В и С прежде, чем дойдет до вершины D. Поиск в ширину исследует граф систематически, просматривая вершины вокруг начальной вершины. Он начинается с посещения всех вершин от стартовой вершины, а затем переходит к следующей и т.д.

Поиск в глубину (Depth-First Search — DFS) начинается с корня графа, а затем переходит к каждому следующему узлу от корня вниз и до конца. Затем он следует в обратном порядке и начинает исследовать путь, отличный от предыдущего, пока снова не достигнет корня. В этот момент, если доступны другие пути от корня, алгоритм выбирает следующий и начинает тот же самый поиск снова. Идея в том, чтобы исследовать каждый путь полностью прежде, чем исследовать любой другой путь.

Состязательные игры

Самое интересное в поиске в пространстве состояний — это то, что он демонстрирует и нынешние и будущие возможности искусственного интеллекта. Так происходит в случае *состязательных игр* (игр, в которых один побеждает, а остальные проигрывают) или любой подобной ситуации, в которой участники преследуют цель, находящуюся в противоречии с целями других. Простая игра в крестики-нолики представляет собой совершенный пример игры поиска в пространстве, в которую вы, возможно, уже играли с искусственным интеллектом. В фильме *Военные игры* (WarGames), вышедшем в 1983 году, суперкомпьютер WOPR (War Operation Plan Response) играет против себя с потрясающей скоростью, но все же не может победить, поскольку игра действительно проста, и если использовать поиск в пространстве состояний, то проиграть невозможно.

В игре девять клеток, в которые каждый игрок записывает крестик или нолик. Первый, кто соберет метки в ряд (горизонтальный, вертикальный или диагональный), победил. При создании *дерева пространства состояний* каждый

уровень дерева представляет ход в игре. Конечные узлы представляют финальные состояния доски и определяют победу, ничью или поражение искусственного интеллекта. У каждого листа есть бал, у победного — высокий, у ничейного — средний, а у проигрышного — низкий или даже отрицательный. Используя суммирование, искусственный интеллект просчитывает балы к верхним узлам и ветвям до достижения стартового узла. Стартовый узел представляет текущую ситуацию. Для пересечения дерева используется простая стратегия: когда ход искусственного интеллекта и вы должны рассмотреть значения многих узлов, вы суммируете максимальное значение (очевидно, потому, что искусственный интеллект должен добиться максимального счета в игре); когда наступает очередь противника, вы суммируете вместо этого минимальное значение. В результате вы получаете дерево, ветви которого квалифицируются балами. Когда наступает очередь искусственного интеллекта, он выбирает свой переход на основании ветви с самым высоким значением, поскольку это подразумевает путь по узлам с самой высокой возможностью победить. Пример этой стратегии приведен на рис. 3.3.

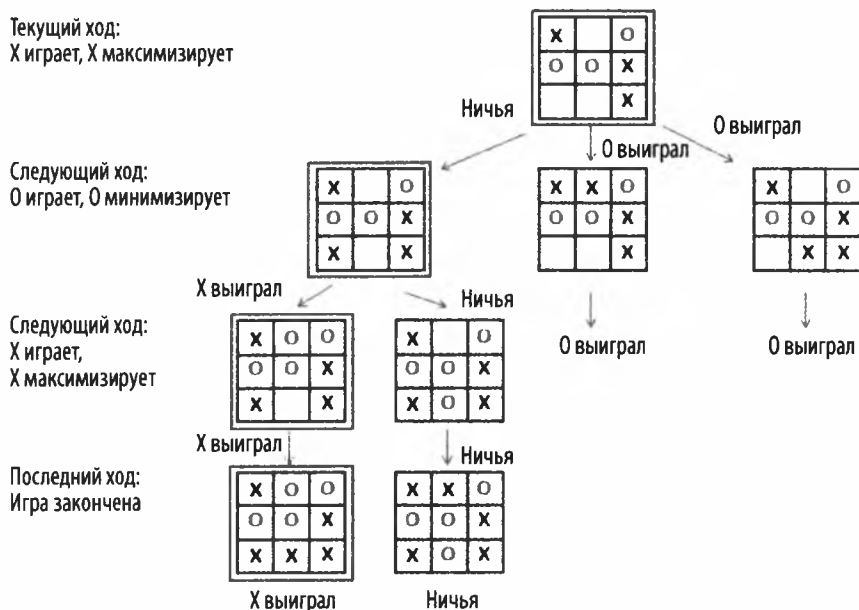


Рис. 3.3. Приближение к мини-максу в игре “крестики-нолики”

Это подход *приближения к мини-максу* (min-max approximation). Рональд Линн Ривест (Ronald Rivest) из лаборатории информатики Массачусетского технологического института опубликовал его в 1987 году (его статью можно прочитать по адресу <https://people.csail.mit.edu/rivest/pubs/Riv87c.pdf>). С тех пор этот алгоритм и его варианты легли в основу множества

соревновательных игр, включая недавно анонсированную игру AlphaGo от Google DeepMind, в которой используется подход, напоминающий приближение к мини-максу (в фильме *Военные игры* 1983 года представлен тот же самый подход).



СОВЕТ

Иногда в контексте приближения к мини-максу можно услышать термин *альфа-бета-отсечение* (alpha-beta pruning). Это интеллектуальный способ сокращения количества рассматриваемых узлов в древовидной иерархии сложных пространств состояний, существенно сокращающий количество вычислений. Не все игры имеют компактные деревья пространства состояний, но когда количества ветвей исчисляются миллионами, их отсечение необходимо для сокращения объемов вычислений.

Использование локального поиска и эвристики

Многое изменилось со времени появления подхода поиска в пространстве состояний. В конце концов, никакая машина независимо от ее мощности не сможет перебрать все возможности в любой ситуации. В этом разделе продолжается рассмотрение игр, поскольку они предсказуемы и имеют фиксированные правила, тогда как большинство реальных ситуаций непредсказуемы и никаких правил не имеют, что делает игры оптимальным выбором.

В шашках, относительно простой игре по сравнению с шахматами или Го, возможно 500 миллиардов миллиардов (500 000 000 000 000 000 000) позиций на доске. Это значение вычислено математиками из Университета Гавайев и соответствует количеству всех песчинок на Земле. Конечно, во время игры в шашки делается куда меньше ходов, но все же количество потенциальных вариантов при каждом ходе слишком велико для вычислений. Чтобы вычислить все 500 миллиардов миллиардов возможных ходов, мощным компьютерам потребуется 18 лет (<http://sciencenetlinks.com/science-news/science-updates/checkers-solved/>). Только представьте, как долго компьютер потребительского класса будет рассчитывать даже куда меньшее количество ходов. Чтобы процесс остался подконтрольным, это должно быть очень маленькое подмножество всех возможных ходов.

Оптимизация с использованием локального поиска и эвристики позволяет ограничить количество возможных вычислений, как при альфа-бета-отсечении, когда некоторые вычисления отбрасываются, поскольку ничего не дают для успеха. *Локальный поиск* (local search) — это общий подход к решению задач, реализуемый множеством алгоритмов, позволяющих избежать экспоненциального роста сложности многих NP-задач. Локальный поиск начинается с текущей ситуации или несовершенного решения задачи и перемещается от

нее шаг за шагом. Локальный поиск выясняет пригодность соседних решений, потенциально приводя к более совершенному решению на основании случайного выбора или осмысленной эвристики (а значит, никакого точного метода решения нет).



ЗАПОМНИ

Эвристика (heuristic) — это обоснованное предположение о решении, такое как эмпирическое правило, которое указывает направление к желаемому результату, но не способное подсказать точно, как его достичь. Это как будто, заблудившись в незнакомом городе, спрашивать людей, как добраться до своей гостиницы. Вам укажут направление, но без точных инструкций и примерного расстояния.

Алгоритмы локального поиска поэтапно улучшают состояние начиная от стартового и переходя шаг за шагом к соседним решениям в пространстве состояний, пока дальнейшее улучшение не окажется невозможным. Поскольку алгоритмы локального поиска просты и интуитивно понятны, разработать подход локального поиска для решения алгоритмической задачи нетрудно, а вот сделать его эффективным обычно более чем трудно. Главное — определить правильную процедуру.

1. Начните с существующей ситуации (это может быть текущая ситуация, случайное или известное решение).
2. Найдите набор новых возможных решений в пределах окрестностей текущего решения; они составят список кандидатов.
3. Определите решение, используемое вместо текущего на основании вывода эвристики, получившей на входе список кандидатов.
4. Продолжите выполнять пп. 2 и 3 до тех пор, пока не прекратится дальнейшее улучшение, означающее нахождение лучшего из решений.

Хотя разработка локального поиска проста, она зачастую не позволяет найти решение за разумное время (вы можете остановить процесс и использовать текущее решение) или предлагает решение, незначительно лучшее. Нет никакой гарантии, что локальный поиск найдет решение задачи, но есть шанс действительно улучшить первоначальное решение, если предоставить ему достаточно времени. Он остановится только после того, как не сможет найти дальнейший способ улучшить решение. Секрет в том, чтобы исследовать правильную окрестность. Если вы исследуете все, то скатитесь к полному перебору, подразумевающему взрывообразный рост вариантов для исследования и проверки.

Доверие к эвристике ограничивает, когда нужна эмпирика. Иногда эвристика — это случайность, и такое решение, несмотря на не интеллектуальность подхода, вполне может прекрасно работать. Немногие знают, что автономный робот-пылесос Roomba, созданный тремя специалистами из Массачусетского

технологического института, первоначально не планировал свою траекторию, он просто двигался в случайных направлениях. И все же владельцы считали его разумным устройством, ведь свою работу по уборке он выполнял превосходно. (Фактически интеллект заключался в идее использовать случайность для решения задачи, которая в противном случае оказалась бы слишком сложной).

Случайный выбор — не единственная доступная эвристика. При исследовании локальный поиск способен полагаться и на более осмысленные решения с использованием лучше обоснованной эвристики для получения направления поиска, такие как *оптимизация восхождением к вершине* (hill-climbing optimization) или *алгоритм возврата* (twiddle), и избежать ловушки принятия посредственных решений, как при *имитации отжига* (simulated annealing) и *поиске с запретами* (tabu search). Оптимизация восхождением к вершине, алгоритм возврата, имитация отжига и поиск с запретами — это все алгоритмы поиска, эффективно используемые эвристикой для получения направления.

Алгоритм восхождения к вершине сродни действию силы гравитации. Представьте, что, когда мяч катится по склону, он выбирает самый крутой спуск, а когда его нужно закатить на холм, имеет смысл выбрать кратчайший путь достижения вершины, которым также является самый крутой склон. Таким образом, задача искусственного интеллекта — спуститься во впадину или подняться на вершину, а эвристика подскажет направление, указав самый крутой подъем или спуск из возможных в пространстве состояний. Это весьма эффективный алгоритм, хотя в некоторых ситуациях, известных как плато (промежуточные точки минимума) и пики (локальные максимальные точки), он ошибается.

Алгоритм возврата или *покоординатного спуска* (coordinate descent) подобен алгоритму восхождения к вершине. Эвристический алгоритм возврата подразумевает исследование всех возможных направлений, но поиск концентрируется в направлении лучших окрестностей. По мере продвижения он калибрует свои шаги, замедляясь по мере приближения к трудным для поиска участкам, чтобы находить лучшие решения, пока процесс не достигнет остановки.

Термин *имитация отжига* произошел от металлургической технологии, в ходе которой металл нагревают, а затем медленно охлаждают, чтобы сделать его мягче для обработки в холодном состоянии, а также для устранения межкристаллических напряжений (см. <http://www.brighthubengineering.com/manufacturing-technology/30476-what-is-heat-treatment/>). Локальный поиск воспроизводит эту технологию при поиске решения, как будто меняется атомная структура, чтобы улучшить работоспособность. Температура — решающий фактор в процессе оптимизации. Подобно тому, как высокая температура ослабляет структуру материала (твердое плавится, а жидкое испаряется), высокая температура в алгоритме локального поиска снижает *критерий отбора* (objective function), позволяя предпочесть худшие решения лучшим.

Имитация отжига изменяет процедуру восхождения к вершине, сохраняя критерий отбора для вычисления соседнего решения и позволяя осуществлять выбор решения для поиска различными способами.

Поиск с запретами подразумевает запоминание подлежащих исследованию частей окрестности. Когда решение кажется найденным, предпринимается попытка вернуться к другим возможным путям, которые еще не опробовались, чтобы выявить лучшее из решений.

Использование направления (вверх или вниз), температуры (контролируемой случайности) или просто запрета или выделения части в искомом фактически является способом избежать перебора всех возможностей и сконцентрироваться на перспективных решениях. Рассмотрим, например, перемещение робота. Задача — провести робота по неизвестной окружающей обстановке, избегая препятствий, и достичь заданного места. Это сложная фундаментальная задача искусственного интеллекта. Для ориентации на местности роботы могут полагаться на лазерный инфракрасный дальномер (*лидар* — LIDAR) или *сонар* (использует звук для наблюдения окружающей обстановки). Но несмотря на техническое оснащение, роботы все еще нуждаются в надлежащих алгоритмах для следующего.

- » Поиск кратчайшего (или по крайней мере разумно короткого) пути к месту назначения.
- » Избегание препятствий на пути.
- » Выполнение специальных требований, таких как минимум поворотов или торможений.

Алгоритм поиска пути позволит роботу начать движение в одном месте и достичь места назначения по кратчайшему пути между этими двумя точками, избежав препятствий. (Реакции после столкновения со стеной недостаточно.) Поиск пути используется также при перемещении любого объекта к цели в пространстве, даже виртуальном, как в видеоигре или на веб-странице. Робот при поиске пути воспринимает пройденный путь как последовательность пространств состояний в границах его сенсоров. Если цель находится вне пределов их досягаемости, робот не будет знать, куда идти. Эвристика может направить его в правильную сторону (например, знание того, что цель находится в северном направлении), позволив избежать ненужных затруднений, связанных с необходимостью рассматривать все остальные возможные пути.

Понятие обучения машины

Все рассмотренные до сих пор примеры алгоритмов связаны с искусственным интеллектом, поскольку все они — интеллектуальные решения повторяющихся, жестко разграниченных, сложных задач, требующих интеллекта. Они требуют, чтобы архитектор изучил задачу и выбрал правильный алгоритм для ее решения. Однако смена задачи, ее коррекция или непредвиденные факторы могут стать реальной проблемой для успешного применения алгоритма. Это связано с тем, что изучение задачи и ее решение осуществляются раз и навсегда, пока алгоритм используется в программном обеспечении. Например, вы вполне можете создать программу с искусственным интеллектом для решения sudoku (популярная игра, требующая вписать числа в клетки поля согласно правилам: <https://www.learn-sudoku.com/what-is-sudoku.html>). Вы даже можете обеспечить алгоритму некую гибкость, которая позволит учесть больше правил или больший размер игрового поля. Питер Норвиг (Peter Norvig), директор по исследованиям корпорации Google, написал чрезвычайно интересное эссе на эту тему (<http://norvig.com/sudoku.html>), в котором демонстрирует, как разумное использование поиска в глубину позволяет сократить количество вычислений (в противном случае вычисления могут длиться вечно). Применение ограничителей позволяет вначале исследовать меньшие ветви, что делает решение sudoku вполне возможным.

К сожалению, не все проблемы могут полагаться на решения, подобные sudoku. Проблемы реального мира никогда не возникают в простых мирах совершенной информации и четких действий. Рассмотрим задачу поиска страхового мошенника или задачу диагностики заболеваний.

- » **Большой набор правил и возможностей.** Количество возможных мошенничеств невероятно велико; у многих болезней схожие симптомы.
- » **Недостаток информации.** Мошенники могут скрывать информацию; врачи зачастую полагаются на неполную информацию (анализы могут еще отсутствовать).
- » **Правила могут изменяться.** Мошенники изобретают все новые и новые способы надувательства; новые болезни возникают и обнаруживаются.

Для решения таких задач нельзя использовать predetermined подход. Чтобы справляться с любыми новыми ситуациями, подход должен быть гибким, необходимо накопление полезных знаний. Другими словами, для адаптации к изменениям в сложной окружающей обстановке необходимо продолжать учиться, как люди учатся весь свой век.

Работа экспертных систем

Экспертные системы были первой попыткой избежать жестко заданных алгоритмов и выработать более гибкие и интеллектуальные способы решения реальных задач. Лежащие в основе экспертных систем идеи были простыми и вполне подходящими на те времена, когда хранение и доступ к большим количествам данных в машинной памяти были дорогостоящими. Сегодня это может показаться странным, но в 1970-х годах такие ученые в области искусственного интеллекта, как Росс Квиллиан (Ross Quillian), вынуждены были демонстрировать построение рабочих языковых моделей на базе словаря только из 20 слов, поскольку машинная память тех времен могла содержать лишь столько слов. Немного возможностей доступно, если компьютер не может содержать все данные, и решение заключалось в работе только с ключевой информацией по проблеме и получению ее от людей, разбирающихся в ней лучше всех.



ЗАПОМНИ

Экспертные системы были экспертными не потому, что их знания базировались на результате собственного обучения, а скорее потому, что они получали свои знания от экспертов-людей, которые предоставляли предварительно подготовленную систему ключевой информации, полученной ими из книг либо от других экспертов или изобретенной самостоятельно. В основном это был интеллектуальный способ воплощения знаний в машине.

Примером одной из первых таких систем является MYCIN; она диагностировала болезни, связанные со свертываемостью крови, и такие бактериальные инфекции крови, как менингит (воспаление мембран, защищающих головной и спинной мозг). Система MYCIN рекомендовала правильную дозировку антибиотиков на основании более 500 правил, а при необходимости вызывала использующего систему врача. Когда информации не хватало, например отсутствовали результаты лабораторных анализов, система MYCIN начинала консультативный диалог, задавая корректные вопросы по существу, чтобы удостовериться в правильности диагноза и терапии.

Написанная на языке LisP в качестве докторской диссертации Эдвардом Хансом Шортлифом (Edward Shortliffe) из Стэнфордского университета, система MYCIN потребовала пяти лет труда, а по завершении работала лучше, чем любой начинающий врач, с перспективой достичь точности диагностики опытного врача. За несколько лет до этого в той же лаборатории была разработана система DENDRAL — первая из когда-либо созданных экспертных систем. Она была весьма сложным приложением, специализировавшимся на органической химии с алгоритмами, решавшими задачи “в лоб” и оказывающимися бесполезными при встрече с человеческой эвристикой на базе практического опыта.

Что касается успеха MYCIN, то возникли некоторые проблемы. Во-первых, были не ясны условия ответственности. (Если система поставила неправильный диагноз, то кто несет ответственность?) Во-вторых, у MYCIN была проблема удобства и простоты использования, поскольку во времена младенчества Интернета врач мог подключиться к MYCIN, только используя дистанционный терминал для мэйнфрейма в Стэнфорде, что было весьма трудно и медленно. Но все же система MYCIN доказала свою эффективность и полноценность в поддержке решений человека, и это проложило путь многим другим экспертным системам, которые распространились чуть позже, в 1970- и 1980-х годах.

Вообще, экспертные системы того времени состояли из двух великолепных компонентов: базы знаний и механизма логического вывода. *База знаний* (knowledge base) хранит знания как коллекцию правил в форме операторов *if...then* (где часть *if* включает одно или несколько условий, а часть *then* — выдаваемые рекомендации). Эти символы операторов, разные у разных систем, способны проверять одиночные события или факты, классы и производные классы, позволяя манипулировать ими, используя булеву логику или сложную логику первого порядка, которая учитывает все возможные операции.



СОВЕТ

Логика первого порядка (first-order logic) — это набор операций, идущих куда далее связанных с простыми комбинациями значений TRUE (истина) и FALSE (ложь). Например, здесь есть такие концепции, как FOR ALL (для всех) и THERE EXIST (существует), позволяющие иметь дело с операторами, которые могут быть истиной, но не могут быть подтверждены доказательством на данный момент. Куда больше об этой форме логики можно узнать из статьи <http://whatis.techtarget.com/definition/first-order-logic>.

Механизм логического вывода (inference engine) — это набор инструкций, указывающих системе, как манипулировать условиями на основании таких операторов булевой логики, как AND, OR или NOT. Этот логический набор, использующий TRUE (соответствующее или, технически, “сработавшее” правило) и FALSE (не соответствующее правило) как символьные условия, способен объединять их в сложные рассуждения.

Поскольку система была создана на базе последовательности операторов *if* (условия) и *then* (выводы) и имела вложенный структурированный по уровням характер, полученная вначале информация позволяла сразу исключить некоторые заключения, а также помочь системе, взаимодействуя с пользователем, выяснить информацию, способную привести к правильному ответу. При использовании с механизмом логического вывода для экспертных систем было характерно следующее.

- » **Прямая цепочка рассуждений.** Доступное доказательство делало на каждом этапе серию правил подходящими, а другие исключало. Система первоначально концентрировалась на правилах, которые могли бы быть подходящими и привести к заключению. Это явно подход, управляемый данными.
- » **Обратная цепочка рассуждений.** Система вычисляет каждое возможное заключение и пытается доказать каждое из них на основании доступного доказательства. Этот подход, управляемый целями, позволяет определить, какие вопросы стоит задать, и исключает целые наборы задач. Система MYCIN использовала обратную цепочку рассуждений — общепринятую стратегию развития от гипотезы назад, к доказанному медицинскому диагнозу.
- » **Разрешение конфликтов.** Если система сделала больше одного вывода за раз, она предпочитает заключение с определенными характеристиками (важность, риск или другие факторы). Иногда система консультируется с пользователем, и решение принимается на основании оценки пользователя. Например, система MYCIN использовала коэффициент уверенности, оценивавший вероятную точность диагноза.

Одним из главных преимуществ таких систем было представление знания в удобочитаемой для людей форме, процесс принятия решения был понятен и изменяем. Если система пришла к выводу, она оповещает об использованных при этом правилах. Пользователь может систематически наблюдать за работой системы и оценивать достоверность или ошибочность выводов. Кроме того, экспертные системы были просты в реализации на таких языках программирования, как LisP или ALGOL. Пользователи улучшали экспертные системы в течение долгого времени, добавляя новые или изменяя существующие правила. Они могли даже заставить ее работать в условиях неуверенности за счет применения *нечеткой логики* (fuzzy logic), своего рода многозадачной логики, в которой значение может содержать нечто среднее между 0, абсолютной ложью и 1, абсолютной истинной. Нечеткая логика избегает резких движений, когда правила выбираются на основании пороговых значений. Например, если утверждается, что правило применимо, если в комнате жарко, оно не будет применено при точном задании температуры, если она чуть-чуть не дотягивает до порогового значения. Закат экспертных систем наступил в конце 1980-х, их разработка остановилась главным образом по следующим причинам.

- » Логика и символика таких систем оказались ограниченными выражением правил, лежащих в основе решения, что привело к созданию индивидуальных систем, сведя все снова к жестко заданным правилам и классическим алгоритмам.

- » Экспертные системы для многих сложных задач стали настолько сложными и запутанными, что потеряли свою привлекательность с точки зрения реализуемости и цены.
- » Поскольку данные становятся все более распространенными и доступными, нет особого смысла тщательно опрашивать опытных специалистов, накапливать и систематизировать ценные знания, когда то же самое (или даже лучшее) можно просто найти в открытых данных.

Экспертные системы все еще существуют. Вы можете найти их используемыми при выдаче кредита, обнаружении мошенничества и в других областях, в которых ответ предоставляется наравне с правилами, лежащими в основе решения, и таким способом, которым пользователь системы находит приемлемым (как сделал бы настоящий эксперт).

Введение в машинное обучение

Решения, способные к самостоятельному обучению непосредственно на данных без их предварительного представления в качестве символов, возникли за несколько десятилетий до экспертных систем. Одни были по природе статистическими; другие по-разному подражали природе; а третьи пытались создать автономную символическую логику в форме правил для необработанной информации. Все эти выработанные различными школами решения сегодня известны под различными названиями, включая *машинное обучение* (machine learning). Машинное обучение — это целый раздел алгоритмов, хотя в отличие от многих других обсуждавшихся до сих пор алгоритмов они не являются набором предопределенных этапов, предназначенных для решения задачи. Как правило, машинное обучение имеет дело с задачами, которые люди не умеют подразделять на этапы, но сами их, естественно, решают. Примером таких задач является распознавание лиц на изображениях или слов в разговорной речи. Машинное обучение упоминается почти в каждой главе этой книги, но работе его основных алгоритмов посвящены только главы 9–11, включая глубокое обучение, вызвавшее новую волну технологий с применением искусственного интеллекта, заголовками о которых пестрят новости почти каждый день.

Достижение новых высот

Роль машинного обучения в новой волне алгоритмов искусственного интеллекта — частично заменить, а частично дополнить существующие алгоритмы, демонстрируя действия, требующие интеллекта, с человеческой точки зрения, поскольку их непросто формализовать в точную последовательность этапов. Хороший пример такой роли — демонстрация мастерства игры в Го, в которой

мастер сразу понимает угрозы и возможности обстановки на доске и интуитивно находит правильные ходы. (Читайте историю игры Го на <http://www.usgo.org/brief-history-go>.)

Го — невероятно сложная игра для искусственного интеллекта. В шахматах приходится просчитывать порядка 35 возможных ходов на доске, и игра обычно длится не более 80 ходов, в то время как в Го возможных ходов бывает до 140 и длится игра обычно более 240 ходов. Никакой вычислительной мощи в современном мире не хватит, чтобы создать полное пространство состояний для игры Го. Группа Google DeepMind в Лондоне разработала программу AlphaGo, победившую многих знаменитых мастеров Го (см. <https://deepmind.com/research/alphago/>). Программа не полагается на алгоритмический подход на базе поиска в огромном пространстве состояний, она использует следующее.

- » Метод интеллектуального поиска на основании проверки случайного возможного хода. Искусственный интеллект многократно применяет поиск в глубину, чтобы выяснить, является ли первый найденный результат лучшим или худшим (в неполном или частичном пространстве состояний).
- » Алгоритм глубокого обучения обрабатывает изображение доски (сразу) и получает как наилучший возможный ход в этой ситуации (алгоритм — *стратегическая сеть* (policy network)), так и оценку вероятности выигрыша искусственного интеллекта после данного хода (алгоритм — *оценочная сеть* (value network)).
- » Возможность учиться на примерах игр, сыгранных экспертами Го в прошлом, а также игра против себя, как WOPR в фильме *Военные игры* 1983 года. Последняя версия программы, AlphaGo Zero, способна учиться совершенно автономно, без человеческих примеров (см. <https://deepmind.com/blog/alphago-zero-learning-scratch/>). Эта способность известна как *обучение с подкреплением* (reinforcement learning).



Глава 4

Первенство специализированных аппаратных средств

В ЭТОЙ ГЛАВЕ...

- » Использование стандартных аппаратных средств
- » Использование специализированных аппаратных средств
- » Улучшение аппаратных средств
- » Взаимодействие с окружающей средой

В главе 1 упоминалось, что одной из причин первых неудач при создании искусственного интеллекта была нехватка подходящих аппаратных средств. Аппаратные средства просто не могли выполнять задачи достаточно быстро, чтобы быть полезными даже для обыденных потребностей, не говоря уже о чем-то столь сложном, как моделирование процесса мышления человека. Эта проблема довольно хорошо представлена в фильме *Игра в имитацию* (The Imitation Game) (<https://www.amazon.com/exec/obidos/ASIN/B00RY86HSU/datacservip0f-20/>), в котором Алан Тьюринг (Alan Turing) взломал код шифровальной машины “Энигма” догадавшись о наличии в конце каждого сообщения повторяющейся фразы “Хайль Гитлер”. Без данной догадки об упущении немецких операторов при использовании “Энигмы” довольно

медленное компьютерное оборудование Тьюринга просто не смогло бы решить задачу достаточно быстро (в фильме это никак не отражено). Если уж на то пошло, исторические документы (полностью рассекречено не так уж и много) свидетельствуют, что проблемы у Тьюринга были куда серьезнее, чем показано в фильме (см. <https://www.scienceabc.com/innovation/cracking-the-uncrackable-how-did-alan-turing-and-his-team-crack-the-enigma-code.html>). К счастью, сегодня стандартные общедоступные аппаратные средства вполне обладают скоростью, достаточной для решения многих задач, с которых начинается эта глава.



ЗАПОМНИ

Действительно, чтобы начать моделировать процесс мышления человека, нужны специализированные аппаратные средства, и даже лучшие из них на сегодня с задачей не справятся. Почти все стандартные аппаратные средства полагаются на *архитектуру фон Неймана* (<http://www.c-jump.com/CIS77/CPU/VonNeumann/lecture.html>), согласно которой память отделена от вычислений, что создает чудесную обобщенную среду обработки, в которой не работают хорошо только некоторые виды алгоритмов, поскольку низкая скорость передачи данных между процессором и памятью создает *узкое место Фон Неймана* (Von Neumann bottleneck). Во второй части этой главы представлены различные методы преодоления узкого места Фон Неймана, чтобы сложные, интенсивно использующие данные алгоритмы работали быстрее.

Даже с аппаратными средствами, специально разработанными для ускорения вычислений, машина, предназначенная для моделирования процесса мышления человека, способна работать только с такой скоростью, с какой позволяют ее механизмы ввода и вывода. Следовательно, необходимо создать лучшую среду, в которой смогут работать аппаратные средства. Это можно сделать множеством способов, но в данной главе рассмотрены только два из них: улучшение возможностей используемого оборудования и использование специализированных сенсоров. Такие улучшения среды работы аппаратных средств хороши, но ниже объясняется, что их все еще недостаточно для моделирования работы человеческого мозга.

В конечном счете аппаратные средства бесполезны, даже самые передовые, если полагающиеся на них люди не могут эффективно с ними взаимодействовать. Заключительный раздел этой главы посвящен методикам повышения эффективности взаимодействия. Эти взаимодействия — просто результат комбинации улучшенного вывода и грамотного программирования. Подобно тому, как Алан Тьюринг использовал уловку, чтобы заставить свой компьютер сделать куда больше того, на что он был способен, эти методики позволят современным

компьютерам выглядеть, как чудо. Фактически компьютер ничего не понимает; весь интеллект закладывают люди, создающие его программы.

Стандартные аппаратные средства

В большинстве проектов искусственного интеллекта предполагается, что их создание по крайней мере начнется со стандартных аппаратных средств, поскольку современные общедоступные компоненты фактически обладают весьма высокой мощностью обработки, особенно по сравнению с компонентами 1980-х годов, когда искусственный интеллект начал давать результаты, пригодные для практического применения. Следовательно, даже если стандартных аппаратных средств не хватит для создания полностью работоспособной системы, вы можете работать на них со своим экспериментальным кодом и получить рабочую модель, которая в конечном счете заработает с полным набором данных.

Понятие стандартных аппаратных средств

Архитектура (architecture), или структура, стандартного компьютера не менялась с тех пор, как Джон фон Нейман впервые предложил ее в 1946 году (см. статью https://www.maa.org/external_archive/devlin/devlin_12_03.html). Согласно истории (<https://lennartb.home.xs4all.nl/coreboot/col2.html>) в компьютерах 1981 года (и даже ранее) процессор уже соединялся с памятью и периферийными устройствами через шину. Во всех этих системах используется архитектура фон Неймана, поскольку она обеспечивает существенные преимущества и модульность. История свидетельствует, что эти устройства допускают модернизацию каждого компонента как индивидуальное решение, позволяя наращивать *возможности*. Например, в определенных пределах вы можете увеличить объем памяти или дискового хранилища любого компьютера. Вы можете также использовать лучшие периферийные устройства. Но все эти элементы подключаются через шину.



ЗАПОМНИ

Увеличение возможностей компьютера не отменяет фактов его базовой архитектуры. Так, у компьютера, используемого сегодня, та же архитектура, что и раньше; просто у него больше возможностей. Кроме того, форм-фактор устройства также не затрагивает архитектуру. Компьютеры в вашем автомобиле также полагаются на шину при подключении, что тоже является архитектурой фон Неймана. (Даже если вид шины другой, архитектура та же самая.) Что бы вы ни думали, но неизменной осталась архитектура любого устройства, достаточно взглянуть на блок-схему Blackberry по адресу <http://mobilesaudi>.

blogspot.com/2011/10/all-blackberry-schematic-complete.html. Она тоже полагается на схему фон Неймана. Следовательно, почти у любого устройства, которым можно разжиться на сегодняшний день, будет подобная архитектура несмотря на различие форм-факторов, типов шины и основных возможностей.

Недостатки стандартных аппаратных средств

У возможности создать модульную систему действительно есть существенные преимущества, особенно в бизнесе. Возможность вынимать и заменять отдельные компоненты существенно снижает стоимость, одновременно позволяя проводить улучшения по скорости и эффективности. Но как всегда, бесплатный сыр бывает только в мышеловке. Модульность архитектуры фон Неймана имеет ряд серьезных недостатков.

- » **Узкое место фон Неймана.** Из всех недостатков этот является самым серьезным с точки зрения требований таких дисциплин, как искусственный интеллект, машинное обучение и даже наука о данных. Более подробно этот недостаток обсуждается в разделе “Понятие узкого места фон Неймана” далее в этой главе.
- » **Единые точки отказа.** Любая потеря соединения с шиной обязательно означает отказ всего компьютера. Даже в системах с несколькими процессорами потеря одного из них приводит не просто к снижению производительности, а к полному отказу системы. То же самое происходит при потере других системных компонентов: вместо снижения функциональных возможностей система отказывает целиком. С учетом того, что искусственный интеллект зачастую требует непрерывной работы системы, это создает вероятность серьезных последствий и не позволяет полагаться на надежность аппаратных средств.
- » **Односмысловость (single-mindedness).** Шина фон Неймана способна получать либо инструкцию, либо данные, необходимые для выполнения инструкции, но она не может делать это одновременно. Следовательно, когда поиск данных требует нескольких циклов шины, процессор простаивает, а это существенно снижает его способность решать задачи искусственного интеллекта, связанные с интенсивным выполнением инструкцией.
- » **Управление задачами.** Когда мозг решает задачу, множество сигналов срабатывают одновременно, обеспечивая параллельное выполнение множества операций. Первоначальный проект фон Неймана допускал только одну операцию за раз и только после того, как система получала и необходимую инструкцию, и данные. Сегодня

у компьютеров, как правило, несколько ядер, что позволяет одновременно выполнять по одной операции в каждом ядре. Однако код приложения должен быть специально разработан так, чтобы затребовать эту возможность, поэтому зачастую она остается неиспользуемой.

ОТЛИЧИЯ ГАРВАРДСКОЙ АРХИТЕКТУРЫ

В своих аппаратных изысканиях вы можете встретить *Гарвардскую архитектуру* (Harvard Architecture), модифицированную форму которой некоторые системы используют для ускорения своей работы. И архитектура фон Неймана, и Гарвардская архитектура полагаются на шинную топологию. Однако при работе с архитектурой фон Неймана аппаратные средства полагаются на одну шину и одну область памяти как для инструкций, так и для данных, тогда как Гарвардская архитектура полагается на отдельные шины для инструкций и данных; они также способны использовать отдельные физические области памяти (см. <http://infocenter.arm.com/help/topic/com.arm.doc.faq/ka3839.html>). Использование индивидуальных шин позволяет системе с Гарвардской архитектурой получать следующую инструкцию, ожидая поступления из памяти данных для текущей инструкции. Таким образом, Гарвардская архитектура быстрее и эффективнее. Однако страдает надежность, поскольку теперь появляются две точки отказа для каждой операции: шина инструкций и шина данных.

Микроконтроллеры, такие как в микроволновой печи, нередко используют Гарвардскую архитектуру. Кроме того, ее можно встретить в некоторых необычных местах, по определенной причине. И iPhone, и Xbox 360 используют модифицированную версию Гарвардской архитектуры, в которой используется единая область памяти (а не две), но шины — отдельные. Причиной в данном случае являются *технические средства защиты авторских прав* (Digital Rights Management — DRM). Область памяти кода можно сделать доступной только для чтения, чтобы никто не мог ее изменить или создать новые приложения без разрешения. С точки зрения искусственного интеллекта это может быть проблематично, поскольку одна из его возможностей заключается в разработке новых алгоритмов (выполняемый код), когда необходимо справиться с непредвиденными ситуациями. Поскольку компьютеры редко реализуют Гарвардскую архитектуру в ее классической форме или как базовую конструкцию шины, в этой книге ей не уделяется много внимания.

Использование графических процессоров

После создания прототипа установки, выполняющего задачи, необходимые для моделирования процесса мышления человека в заданной области, вам, вероятно, понадобятся дополнительные аппаратные средства. Эти средства, обладающие достаточной вычислительной мощностью для работы с полным набором данных, необходимы для рабочей системы. Для обеспечения требуемой вычислительной мощности доступно множество путей, но, как правило, в дополнение к центральному процессору машины используют *графические процессоры* (Graphic Processing Unit — GPU). В следующих разделах описана предметная область применения GPU, что именно подразумевается под термином “GPU” и почему они ускоряют обработку.

ЭЛЕКТРОННО-МЕХАНИЧЕСКАЯ МАШИНА АЛАНА ТЬЮРИНГА

Машина Алана Тьюринга ни в коей форме не имела искусственного интеллекта. Фактически это даже не реальный компьютер. Она взломала криптографические сообщения “Энигмы”, и это все. Однако она действительно дала пищу для рассуждений Тьюринга, которые в конечном счете привели к публикации в 1950-х годах его работы “Вычислительные машины и разум” (Computing Machinery and Intelligence) (<http://www.loebner.net/Prize/TuringArticle.html>), как и в фильме *Игра в имитацию*. Однако сама машина фактически была создана на базе польской машины Bomba. Хотя некоторые источники и настаивают на том, что Алан Тьюринг работал сам, Bomba была произведением многих людей, в частности Гордона Уэлчмана (Gordon Welchman). Тьюринг также не возник из ничего, готовый крушить немецкие шифры. В его время в Принстоне находились такие великие люди, как Альберт Эйнштейн и Джон фон Нейман (продолживший саму концепцию программного обеспечения). Работа Тьюринга вдохновила этих и других ученых на эксперименты и доказала, что это возможно.

Пока ученые писали бумаги, ругали идеи друг друга, выдвигали собственные идеи и экспериментировали, продолжали возникать все новые и новые специализированные аппаратные средства всех видов. Когда смотришь фильм или читаешь книгу, претендующую на историческую правдивость, создается впечатление, что эти люди просыпались однажды утром и, сказав “О, сегодня я сделаю открытие!”, совершали нечто невероятное. На самом деле все новые достижения базируются на прежних достижениях. Таким образом, история важна, поскольку позволяет проследить пути развития, а также освещает другие перспективные пути.

Понятие узкого места фон Неймана

Узкое место фон Неймана (Von Neumann bottleneck) — это естественный результат использования шины для передачи данных между процессором, памятью, долговременным хранилищем и периферийными устройствами. Независимо от того, как быстро шина выполняет свою задачу, превзойти ее всегда невозможно (т.е. она создает узкое место, снижающее скорость). Со временем скорость процессора продолжала расти, а разработки памяти и других устройств сосредоточивались на плотности — возможности размещения большего в меньшем объеме. Следовательно, с каждым усовершенствованием узкое место становилось все большей проблемой, вынуждая процессор тратить все больше времени на ожидание.

Некоторые из проблем, связанных с узким местом Фон Неймана, можно преодолеть в разумных пределах и получить хоть и относительно небольшое, но существенное увеличение скорости выполнения. Вот наиболее распространенные решения.

- » **Кеширование.** Когда проблемы со скоростью получения данных из памяти в архитектуре фон Неймана стали очевидными, независимые разработчики аппаратных средств быстро ответили на них добавлением встроенной памяти, которая не требовала шинного доступа. Эта память является внешней по отношению к процессору, но является частью пакета процессора. Высокоскоростной кеш дорог, поэтому его размеры обычно невелики.
- » **Кеш процессора.** К сожалению, внешние кешы все еще не обеспечивают достаточной скорости. Даже использование самой быстрой оперативной памяти из доступных и полное исключение шинного доступа не обеспечат скорости, сравнимой с таковой у процессора. Следовательно, производители начали добавлять внутреннюю память — кеш, меньший, чем внешний кеш, но с очень быстрым доступом, уже как часть процессора.
- » **Предвыборка.** Проблема кеша в том, что он оказывается полезным только тогда, когда содержит правильные данные. К сожалению, процент удачных обращений к кешу низок в приложениях, использующих много данных и выполняющих широкое разнообразие задач. Следующий этап ускорения работы процессоров — это прогнозирование, какие данные потребуются приложению, и их загрузка в кеш прежде, чем приложение их затребует.
- » **Использование специальной оперативной памяти.** В неразбе-
рихе оперативной памяти можно утонуть, поскольку ее видов существует куда больше, чем большинство людей воображает. Каждый вид оперативной памяти способен решать по крайней мере часть

проблемы узких мест фон Неймана, и они действительно работают в определенных рамках. Как правило, усовершенствования сосредотачиваются на идее получения данных из памяти быстрее, чем через шину. На скорость влияют два главных (и множество второстепенных) фактора: *скорость памяти* (как быстро память обменивается данными) и *задержка* (как долго осуществляется поиск конкретной части данных). Больше о памяти и факторах, влияющих на ее скорость, можно узнать по адресу <http://www.computermemoryupgrade.net/types-of-computer-memory-common-uses.html>.



ВНИМАНИЕ!

Подобно многим другим областями технологии, обмен данными может стать проблемой. Например, *многопоточность* (multithreading), ставшая точкой преткновения для приложений, а также другой набор инструкций для раздельного выполнения процессором блоков кода, которые он может обрабатывать и по одному, зачастую рекламируются как средство преодоления узкого места фон Неймана, но фактически они не дают ничего, кроме дополнительной сложности (или проблем похуже). Многопоточность — это ответ на другую задачу: повышение эффективности приложений. Когда к узкому месту фон Неймана приложение добавляет собственные проблемы задержки, вся система замедляется еще больше. Многопоточность позволяет процессору не тратить время впустую на ожидание реакции пользователя или приложения, а заняться вместо этого чем-то другим. Задержка приложения присуща любой архитектуре процессора, не только архитектуре фон Неймана. Но даже в этом случае все ускоряющее общую работу приложения окажется ощутимым и для пользователя, и для системы в целом.

Определение GPU

Первоначально задача *графического процессора* (Graphics Processing Unit — GPU) заключалась в быстрой обработке данных изображений и их последующем отображении на экране. На ранней фазе развития компьютеров всю обработку выполнял процессор, а значит, графика могла отображаться очень медленно, пока процессор выполнял и другие задачи. В те времена компьютер обычно имел такое оборудование, как *видеоадаптер* (display adapter), обладавший (или не обладавший) некоторыми вычислительными возможностями. Все видеоадаптеры должны были преобразовывать компьютерные данные в визуальную форму. Фактически использование только одного процессора казалось почти неизбежным со времен появления компьютеров, когда дисплей

были только текстовыми или с чрезвычайно простой графикой на 16 цветов. Графические процессоры действительно не предвещали никакой революции в вычислительной сфере, пока людям не понадобился трехмерный вывод. Комбинация процессора и видеоадаптера на это просто не была способна.

Первым шагом в этом направлении стал Hauppauge 4860 (<http://www.geekdot.com/hauppauge-4860/>), системная плата которого включала центральный процессор и специальный графический процессор (в данном случае — 80860). Микросхема 80860 выполняла вычисления чрезвычайно быстро (см. <http://www.cpu-world.com/CPUs/80860/index.html>). К сожалению, эта многопроцессорная асинхронная система никак не оправдала возлагавшихся на нее надежд (хотя и была невероятно быстрой для своего времени) и оказалась чрезвычайно дорогой. Кроме того, были проблемы с написанием приложений, способных задействовать вторую микросхему. Эти две микросхемы также совместно использовали память (которой в системе было достаточно).

С появлением GPU обработка графики переместилась с системной платы на плату периферийного графического устройства. Центральный процессор может указать процессору GPU, выполнить некую задачу, а GPU сам определяет наилучший метод для этого и делает все независимо от центрального процессора. У GPU есть отдельная память и собственный очень широкий путь получения данных из шины. Кроме того, процессор GPU способен напрямую обращаться к оперативной памяти, чтобы получать необходимые для выполнения задачи данные и переносить результаты независимо от центрального процессора. Следовательно, такая схема сделала современные графические дисплеи возможными.



Действительно выдающимися GPU делает то, что они обычно содержат сотни ядер (см. <https://www.nvidia.com/en-us/about-nvidia/ai-computing/>), в отличие от лишь нескольких ядер центрального процессора. Центральный процессор обладает куда более универсальными функциональными возможностями, чем GPU, но GPU выполняет вычисления невероятно быстро, а их результат на дисплей передает даже еще быстрее. Эта способность делает специальные процессоры GPU критически важным компонентом современных систем.

Причины хорошей работоспособности GPU

Как и описанные в предыдущем разделе микросхемы 80860, современные GPU превосходны при выполнении специальных задач, связанных с обработкой графики, включая векторную. Все эти параллельно работающие ядра действительно способны ускорять вычисления для решения задач искусственного интеллекта.

В проекте Google Brain 2011 года (<https://research.google.com/teams/brain/>) искусственный интеллект обучали распознавать котов и людей при просмотре роликов на YouTube. Для решения этой задачи Google задействовал две тысячи процессоров в одном из своих гигантских центров данных. Немногим людям удастся разжиться подобным количеством ресурсов, чтобы повторить эксперимент Google.

Тем не менее Брайану Катанцаро (Bryan Catanzaro) и Эндрю Ингу (Andrew Ng) удалось повторить эксперимент Google, используя набор из 12 GPU NVidia (см. <https://blogs.nvidia.com/blog/2016/01/12/accelerating-ai-artificial-intelligence-gpus/>). После того как стало понятно, что платы GPU способны заменить кучу компьютерных систем, снабженных процессорами, появилась возможность запуска множества новых проектов искусственного интеллекта. В 2012 году Алекс Крижевски (Alex Krizhevsky) из Университета Торонто выиграл конкурс по компьютерному распознаванию образов, ImageNet, применив процессоры GPU. Фактически GPU теперь используют многие исследователи, причем с удивительным успехом (см. <https://adeshpande3.github.io/The-9-Deep-Learning-Papers-You-Need-To-Know-About.html>).

Создание специализированной среды обработки

Глубокое обучение и искусственный интеллект — это не фон-неймановские процессы, согласно многим экспертам, таким как Массимилиано Версаче (Massimiliano Versace), CEO из Neurala Inc. (<https://www.neurala.com/>). Поскольку выполняемая алгоритмами задача не соответствует используемому оборудованию, возникают всякого рода сложности, затрудняющие получение результата намного больше, чем хотелось бы. Поэтому разработка аппаратных средств, соответствующих программному обеспечению, весьма востребована. Управление перспективных исследовательских проектов Министерства обороны США (Defense Advanced Research Projects Agency — DARPA) реализовало один такой проект — Systems of Neuromorphic Adaptive Plastic Scalable Electronics (SyNAPSE). Его идея заключалась в копировании естественного подхода к решению задач за счет объединения памяти и процессора, а не их разделения. И они фактически создали эту систему (она была колоссальна), и вы можете прочитать о ней по адресу <http://www.artificialbrains.com/darpa-synapse-program>.

Проект SyNAPSE получил развитие. Копания IBM создала куда меньшую подобную систему, используя современные технологии; получилось невероятно быстро и эффективно (см. <http://www.research.ibm.com/>

cognitive-computing/neurosynaptic-chips.shtml). Единственная проблема была в том, что их никто не покупал. Многие люди утверждали, что Betamax лучше для хранения данных, чем VHS, но VHS победил по стоимости, легкости в эксплуатации и т.д. (см. <https://gizmodo.com/betamaxvs-vhs-how-sony-lost-the-original-home-video-1591900374>). То же самое относится и к предложению SyNAPSE от IBM, TrueNorth. Трудно было найти людей, желавших платить больше, хотя, конечно, были программисты, желавшие разрабатывать программное обеспечение, используя новую архитектуру и преимущества новой микросхемы. Таким образом, объединение процессора и GPU, даже с учетом его недостатков, продолжает побеждать.



ЗАПОМНИ

В конечном счете кто-то, вероятно, создаст микросхему, которая куда лучше напоминает биологический эквивалент мозга. Нынешние системы, вероятно, не смогут достичь необходимой вычислительной мощности. Такие компании, как Google, работают над альтернативами, например *тензорный процессор* (Tensor Processing Unit — TPU), который предполагается фактически использовать в таких приложениях, как Google Search, Street View, Google Photos и Google Translate (см. <https://cloud.google.com/blog/products/gcp/an-in-depth-look-at-googles-first-tensor-processing-unit-tpu>). Поскольку сейчас уже используются технологии для на самом деле крупномасштабных приложений, некоторые люди уже покупают новые микросхемы, некоторые программисты уже знают, как писать для них приложения, и уже существуют великолепные востребованные продукты. В отличие от SyNAPSE, процессоры TPU полагаются на хорошо известную технологию ASIC (Application Specific Integrated Circuit — *специализированная интегральная микросхема*), которая используется в бесчисленных приложениях, поэтому Google в действительности повторяет уже существующую технологию. В результате возможности микросхем этого вида, преуспевающих на рынке, намного выше, чем когда-то у системы SyNAPSE, которая полагалась на полностью новые технологии.

Увеличение возможностей аппаратных средств

Процессоры все еще хороши для бизнес-систем и приложений, требующих общей гибкости, когда факторы программирования перевешивают чистую вычислительную мощь. Однако сейчас GPU стали стандартом в науке о данных,

машинном обучении, искусственном интеллекте и глубоком обучении. Конечно, изыскания и разработки чего-то следующего и очень большого продолжаются. Центральный процессор и процессор GPU — это процессоры рабочего уровня. В будущем вместо этих двух видов процессоров может использоваться нечто иное.

- » **Специализированные ИМС (ASIC).** В отличие от процессоров общего назначения, производители создают схемы ASIC для вполне конкретной цели. Решения ASIC обеспечивают чрезвычайно высокую производительность, потребляя очень мало мощности, однако гибкость у них отсутствует. Примером решения ASIC является тензорный процессор Google (TPU), который используется только для обработки речи (см. <https://cloud.google.com/blog/products/gcp/an-in-depth-look-at-googles-first-tensor-processing-unit-tpu>).
- » **Программируемые пользователем вентильные матрицы (FPGA).** Подобно ASIC, производители обычно создают *программируемые пользователем вентильные матрицы* (Field Programmable Gate Array — FPGA) для вполне определенной цели. Однако в отличие от ASIC, матрицу FPGA можно запрограммировать так, чтобы изменить ее основные функциональные возможности. Примером решения FPGA является Brainwave от Microsoft для проектов глубокого обучения (см. <https://techcrunch.com/2017/08/22/microsoft-brainwave-aims-to-accelerate-deep-learning-with-fpgas/>).



ЗАПОМНИ

Битва между схемами ASIC и FPGA обещает накалиться, ведь у них появились победители — разработчики искусственного интеллекта. В настоящее время лидерство, похоже, принадлежит Microsoft и FPGA (см. <https://www.forbes.com/sites/moorinsights/2017/08/28/microsoft-fpga-wins-versus-google-tpus-for-ai/#6448980d3904>), но дело в том, что технологии изменчивы, как жидкость, и перемены не заставят себя ждать.

Производители также работают над совершенно новыми типами процессоров, которые могут заработать, как ожидается, а могут и не заработать. Например, Graphcore работает над процессором IPU (Intelligence Processing Unit), как описано по адресу <https://www.prnewswire.com/news-releases/sequoia-backs-graphcore-as-the-future-of-artificial-intelligence-processors-300554316.html>. С учетом всех прежних преувеличений, к подобным новостям о новых процессорах следует относиться с долей осторожности. Но когда видишь реальные приложения от таких больших компаний, как Google и Microsoft, то начинаешь чувствовать немного больше уверенности в будущем подобных технологий.

Добавление специализированных сенсоров

Важнейший компонент искусственного интеллекта — это его способность моделировать человеческий интеллект, используя полный набор чувств. Осмысление восприятия поможет людям в разработке различных видов интеллекта, как описано в главе 1. Человеческие органы чувств обеспечивают правильный способ восприятия для формирования интеллекта человека. Даже если предположить, что искусственный интеллект станет на это способен, для полной реализации всех семи видов интеллекта все еще потребуется правильный вид ввода, который сделает этот интеллект функциональным.

У людей обычно есть пять органов чувств для взаимодействия с окружающей средой: зрение, слух, обоняние, осязание и вкус.¹ Достаточно странно, но люди все еще полностью не понимают своих возможностей, поэтому нет ничего удивительного в том, что компьютеры отстают, когда дело доходит до восприятия окружающей обстановки таким же самым образом, как это делают люди. Например, до недавнего времени считалось, что на вкус различаются только соленое, сладкое, горькое и кислое. Теперь обнаружено еще два вкуса: умами и жирное (см. <https://fivethirtyeight.com/features/can-we-taste-fat/>). Аналогично некоторые женщины тетрахромны (<https://concettaantico.com/tetrachromacy/>), они способны различать 100 000 000 цветов, а все остальные — не более 1 000 000 (тетрахромами могут быть только женщины, из-за особенностей хромосом). Сейчас точное количество женщин, обладающих этой особенностью, неизвестно. (Согласно некоторым источникам их количество весьма велико, порядка 15 процентов; см. <http://sciencevibe.com/2016/12/11/the-women-that-see-100-million-colors-live-in-a-different-world/>).

Использование отфильтрованных статических и динамических данных позволяет сегодня искусственному интеллекту взаимодействовать с людьми определенными способами. Рассмотрим, например, устройство Amazon Alexa (<https://www.amazon.com/Amazon-Echo-And-Alexa-Devices/b?node=9818047011>), которое вас вполне очевидно слышит, а затем что-то отвечает. Даже при том, что Alexa фактически не понимает ничего из того, что вы говорите, общение с ней весьма увлекательно, и люди наделяют эти устройства человеческими чертами. Вообще, для работы Alexa требует доступа к специальному сенсору: микрофону, который позволяет ей слышать. Фактически у Alexa есть несколько микрофонов, чтобы слышать достаточно хорошо и создать иллюзию понимания. К сожалению, даже такая передовая вещь как Alexa, не может видеть, чувствовать, осязать или испытать нечто, что сделает ее хоть в чем-то отдаленно напоминающим человека.

¹ А также вестибулярный аппарат (чувства равновесия, положения в пространстве, ускорения, ощущение веса). — *Примеч. ред.*



СОВЕТ

В некоторых случаях люди хотят, чтобы органы чувств их искусственного интеллекта превосходили или отличались от человеческих. Искусственный интеллект, обнаруживающий движение ночью и реагирующий на него, мог бы полагаться на инфракрасное, а не обычное зрение. Фактически применение альтернативных органов чувств — это один из правильных способов использования искусственного интеллекта сегодня. Возможность работать в окружающей среде, в которой люди работать не могут, является одной из причин огромной популярности роботов некоторых типов, но работа в этих средах зачастую требует набора чувств (сенсоров), отличных от человеческих. Следовательно, тема сенсоров фактически относится к двум категориям (ни одна из которых не определена полностью): сенсоры, подобные человеческим, и альтернативные сенсоры окружающей обстановки.

Разработка методов взаимодействия с окружающей средой

Искусственный интеллект, замкнутый сам на себя и никак не взаимодействующий с окружающей средой, абсолютно бесполезен. Конечно, это взаимодействие имеет форму ввода и вывода. Традиционный метод обеспечения ввода и вывода — непосредственно через легко понятные компьютеру потоки данных, такие как наборы данных, текст, запросы и т.д. Но эти подходы едва ли дружелюбны человеку, для них требуются специальные навыки.



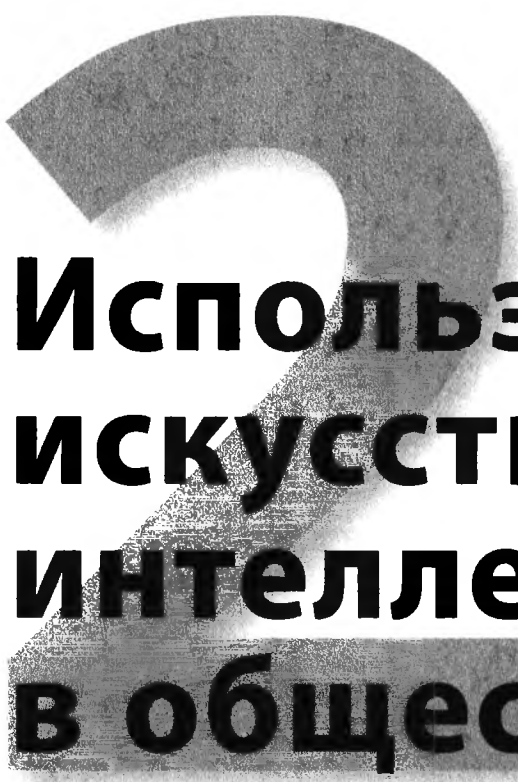
ЗАПОМНИ

Взаимодействие с искусственным интеллектом все чаще осуществляется способами, более приемлемыми для людей, чем непосредственный контакт с компьютером. Например, когда вы задаете Alexa вопрос, ввод осуществляется через набор микрофонов. Искусственный интеллект превращает ключевые слова вопроса в лексемы, которые он может понять. Затем эти лексемы инициализируют вычисления, в результате которых формируется вывод. Искусственный интеллект преобразует вывод в форму, понятную человеку: разговорную речь. Затем вы слышите сентенцию, проговариваемую Alexa через динамик. Короче говоря, для обеспечения полезных функциональных возможностей Alexa должна взаимодействовать с окружающей средой двумя разными способами, присущими людям, но фактически непонятными для самой Alexa.

Взаимодействия могут принимать множество форм. Фактически количество видов и форм взаимодействий непрерывно растет. Например, искусственный интеллект может теперь различать запах (см. <http://www.sciencemag.org/news/2017/02/artificial-intelligence-grows-nose>). Однако фактически компьютер никакого запаха не чувствует. Сенсоры преобразуют результаты химического анализа в данные, которые искусственный интеллект может затем использовать таким же образом, как и любые другие данные. Возможность анализировать химикаты не нова; способность преобразовывать результаты анализа химикатов в цифровую форму не нова; как не новы и алгоритмы, используемые для взаимодействия с вновь полученными данными. Новым является набор данных, используемый для интерпретации входящих данных как запаха, но эти наборы данных являются результатом человеческих исследований. У носа искусственного интеллекта есть множество возможных областей применения. Например, искусственный интеллект мог бы использовать нос при работе в очень опасной окружающей обстановке, например почувствовать запах при утечке газа, чего другие сенсоры не способны заметить.

Физические взаимодействия также развиваются. Роботы, работающие на сборочных линиях, уже привычны, но давайте рассмотрим роботы, способные управлять автомобилем. Это отличная область для применения физического взаимодействия. Предположим также, что реагировать искусственный интеллект может небольшим количеством способов. Хью Херр (Hugh Herr), например, использует искусственный интеллект для взаимодействия с бионическим протезом ноги (см. <https://www.smithsonianmag.com/innovation/future-robotic-legs-180953040/>). Этот бионический протез является превосходной заменой для людей, потерявших реальную ногу. Вместо статической обратной связи, предоставляемой человеку стандартным протезом, этот фактически обеспечивает активную обратную связь, которую люди привыкли получать от реальной ноги. Например, опора на ногу при подъеме по склону отличается от таковой при спуске. Аналогично подъем на бордюр требует обратной связи, отличной от таковой при спуске с него.

Дело в том, что поскольку искусственный интеллект становится все более и более способным выполнять сложные вычисления с большими наборами данных, его возможность выполнять действительно интересные задачи растет. Но человеческих категорий у задач, выполняемых искусственным интеллектом, сегодня быть не может. Вы никогда не сможете взаимодействовать с искусственным интеллектом, который понимает вашу речь, но вы можете положиться на искусственный интеллект, который спасет вам жизнь или по крайней мере сделает ее лучше.



Использование искусственного интеллекта в обществе

В ЭТОЙ ЧАСТИ...

- » Работа с искусственным интеллектом в компьютерных приложениях**
- » Использование искусственного интеллекта для автоматизации популярных процессов**
- » Как искусственный интеллект решает медицинские задачи**
- » Методы взаимодействия непосредственно с людьми**



Глава 5

Искусственный интеллект в компьютерных приложениях

В ЭТОЙ ГЛАВЕ...

- » Определение и применение искусственного интеллекта в приложениях
- » Использование искусственного интеллекта для исправлений и рекомендаций
- » Потенциальные ошибки искусственного интеллекта

Искусственный интеллект вы, вероятно, уже используете, в некоторой форме он присутствует в большинстве компьютерных приложений, на которые вы полагаетесь в своей работе. Например, разговор с вашим смартфоном требует искусственного интеллекта для распознавания речи. Аналогично искусственный интеллект отфильтровывает весь спам, поступающий в вашу папку “Входящие”. В первой части этой главы обсуждаются типы приложений искусственного интеллекта, большинство из которых вас удивит, а также области, в которых он обычно решает множество важных задач. Вы также узнаете об источнике ограничений на создание приложений, использующих

искусственный интеллект, это поможет вам понять, почему разумных роботов не будет никогда, по крайней мере при технологиях, доступных на настоящее время.

Но независимо от того, достигнет ли искусственный интеллект когда-нибудь самосознания, остается тот факт, что он действительно уже выполняет существенное количество практических задач. В настоящее время искусственный интеллект помогает людям двумя основными способами: это исправления и рекомендации. Вряд ли эти два термина требуют объяснений. Исправление — это необязательно ответ на ошибку, а рекомендация — это необязательно ответ на вопрос. Рассмотрим, например, помощь при вождении автомобиля (когда искусственный интеллект помогает водителю, а не заменяет его). Когда автомобиль находится в движении, искусственный интеллект может вносить небольшие корректировки, учитывающие дорожные условия и обстановку, включая пешеходов и множество других проблем еще до того, как ошибка станет фактом. Искусственный интеллект применяет превентивный подход к решению проблемы, которая может возникнуть, а может и нет. Аналогично искусственный интеллект может предложить управляющему автомобилем человеку определенный путь к цели поездки, наилучший на данный момент, а затем изменить рекомендацию на основании новых условий.

Во второй части главы исправления и рекомендации рассматриваются по отдельности.

В третьей части главы обсуждаются потенциальные ошибки искусственного интеллекта. Ошибка происходит всякий раз, когда результат отличен от ожиданий. Результат может быть успешным, но он остается неожиданным. Конечно, настоящие ошибки также происходят; искусственный интеллект не всегда может предоставить успешный результат. Результат может даже противоречить первоначальной задаче (и принести вред). Если вы понимаете, что приложения искусственного интеллекта имеют полутона, а не только черные или белые цвета, то это хорошо, и вы находитесь на пути к пониманию того, как искусственный интеллект изменяет типичные компьютерные приложения, которые действительно предоставляют абсолютно правильный или абсолютно неправильный результат.

Наиболее популярные типы приложений

Подобно тому, как единственным ограничением разнообразия типов процедурных компьютерных приложений является воображение программистов, приложения искусственного интеллекта применимы в любом месте и почти для любых целей, о которых многие даже не подозревают. Гибкость предложений

искусственного интеллекта фактически означают, что некоторые из них могут использоваться в таких местах, для которых они первоначально даже не предназначались. Фактически когда-нибудь программные средства систем искусственного интеллекта смогут описать собственное следующее поколение (см. <https://www.technologyreview.com/s/603381/ai-software-learnsto-make-ai-software/>). Однако, чтобы получить лучшее представление о том, что делает искусственный интеллект полезным в приложениях, имеет смысл рассмотреть наиболее популярные случаи использования искусственного интеллекта сегодня (а также возможные проблемы, связанные с этим), описанные в следующих разделах.

Использование искусственного интеллекта в типичных приложениях

Искусственный интеллект можно найти в таких местах, где его применение даже трудно предположить. Например, ваш интеллектуальный термостат, контролирующий температуру в доме, может содержать искусственный интеллект, если он достаточно сложен (см. <https://www.popsci.com/gadgets/article/2011-12/artificially-intelligent-thermostats-learns-adapt-automatically>). Использование искусственного интеллекта даже в таких сугубо специфических приложениях действительно имеет смысл, если он используется для вещей, в которых искусственный интеллект особенно хорош, например для отслеживания предпочтительных температур на протяжении долгого времени, чтобы автоматически выработать расписание температурного режима. Вот лишь некоторые из наиболее типичных случаев использования искусственного интеллекта.

- » Моделирование творчества
- » Компьютерное зрение, виртуальная реальность и обработка изображений
- » Диагностика (искусственный интеллект)
- » Распознавание лиц
- » Искусственный интеллект игр, боты компьютерных игр, теория игр и стратегическое планирование
- » Распознавание рукописного текста
- » Обработка текстов на естественном языке, перевод и виртуальные собеседники
- » Нелинейное управление и робототехника
- » Оптическое распознавание образов
- » Распознавание речи

Области применения искусственного интеллекта

Приложения определяют конкретные виды применения искусственного интеллекта. Искусственный интеллект может использоваться и в более общих экспертных областях. Вот список наиболее вероятных областей применения искусственного интеллекта.

- » Искусственная жизнь
- » Автоматизация формулирования логических выводов
- » Автоматизация
- » Бионика
- » Интеллектуальный анализ понятий
- » Интеллектуальный анализ данных
- » Фильтрация спама электронной почты
- » Гибридная интеллектуальная система
- » Интеллектуальный агент и интеллектуальное управление
- » Представление знаний
- » Судопроизводство
- » Поведенческая робототехника, процесс познания, кибернетика, эволюционная робототехника (эпигенетика) и инженерная робототехника
- » Семантически структурированная сеть

Аргумент китайской комнаты

В 1980 году Джон Роджерс Сёрл (John Searle) опубликовал в научном журнале *Behavioral and Brain Sciences* статью *Сознание, мозг и программы* (Minds, Brains, and Programs). Акцент в ней был сделан на опровержении теста Тьюринга, в котором, отвечая на серию вопросов, компьютер может уверить человека в том, что он является человеком, а не компьютером (см. <https://www.abelard.org/turpap/turpap.php>). Основная идея в том, что функционализм, или способность имитировать определенные характеристики человеческого разума, — это вовсе не то же самое, что и фактическое мышление.

Мысленный эксперимент “Китайская комната” полагается на два теста.

В первом тесте некто создает искусственный интеллект, способный получать китайские иероглифы и, используя набор правил, создавать на них ответ, а затем выводить ответ, также используя китайские иероглифы. Искусственный интеллект должен интерпретировать поставленный вопрос и ответить на него,

отражая фактический смысл, а не просто случайным образом. Искусственный интеллект настолько хорош, что никто вне комнаты не может догадаться, что эту задачу выполняет искусственный интеллект. Люди, говорящие на китайском языке, полностью вводятся в заблуждение, полагая, что искусственный интеллект действительно может читать и понимать китайский язык.

Во втором тесте человеку, не говорящему на китайском языке, предоставляются три вещи, как и компьютеру в предыдущем тесте. Первый — список, содержащий огромное количество китайских иероглифов, второй — история на китайском языке и третий — набор правил для соотнесения первого элемента со вторым. Некто передает набор вопросов, написанных на китайском языке, человек в комнате, используя свод правил, находит на основании интерпретации китайских иероглифов место в истории, содержащее ответ.

Полученный на основании правил ответ в виде набора китайских иероглифов вполне соответствует вопросу. Человек может настолько хорошо справляться с этой задачей, что никому и в голову не придет, что он совершенно не знает китайского языка.

Эти тесты призваны продемонстрировать, что с помощью формальных правил можно получить результат (синтаксис) и без фактического понимания того, что происходит (семантика). Сёрл сделал вывод, что синтаксиса недостаточно для семантики, но все же некоторые люди, реализующие искусственный интеллект, пытаются доказать обратное, когда дело доходит до создания различных механизмов на базе правил, таких как Script Applier Mechanism (SAM) (см. <https://eric.ed.gov/?id=ED161024>).

Основная проблема — в наличии сильного искусственного интеллекта, того, который фактически понимает то, что пытается делать, а слабый искусственный интеллект — это тот, который просто следует правилам. Весь искусственный интеллект сегодня — это слабый искусственный интеллект, который фактически ничего не понимает. То, что вы видите, является хитрой программой, имитирующей мышление за счет использования правил (rule) (включая неявные в алгоритмах).

Конечно, возникает большое идейное противоречие, согласно которому независимо от сложности будущих машин, фактически разработать мозг не получится, а значит, они никогда не будут ничего понимать.

Таким образом, Сёрл утверждает, что искусственный интеллект всегда останется слабым. Более подробная информация по этой теме приведена на сайте <http://www.iep.utm.edu/chineser/>. Читать аргументы и контраргументы на этом сайте интересно, поскольку они позволяют понять суть того, что действительно важно при создании искусственного интеллекта.

Как искусственный интеллект делает приложения более дружелюбными

Есть множество разных взглядов на вопрос дружелюбия приложений с поддержкой искусственным интеллектом. В самой простой форме искусственный интеллект может помогать при пользовательском вводе. Например, когда пользователь вводит несколько букв некоего слова, искусственный интеллект предлагает остальные буквы. Оказывая эту услугу, искусственный интеллект достигает нескольких целей.

- » Пользователь работает эффективнее, вводя меньше знаков.
- » Приложение получает меньше записей с ошибками из-за опечаток.
- » Пользователь и приложение, оба, участвуют в высокоуровневой коммуникации. Приложение узнает у пользователя правильные или улучшенные термины, которых пользователь мог бы и не помнить, а также позволяет избегать альтернативных терминов, которые компьютер не может распознать.

Искусственный интеллект может также обучаться на предыдущем пользовательском вводе, видоизменяя рекомендации таким способом, который лучше соответствует подходам пользователя к решению задач. Этот следующий уровень взаимодействия находится в пределах области рекомендаций, описанной далее, в разделе “Создание рекомендаций”. Рекомендации могут также включать подсказку пользователю идей, которые в противном случае пользователь, возможно, и не рассматривал бы.



ВНИМАНИЕ!

Даже в области рекомендаций у людей может сложиться мнение, что искусственный интеллект способен думать, но это не так. Искусственный интеллект просто выполняет расширенный поиск по шаблону и анализ вероятности уместности термина в текущем вводе. Различия между слабым искусственным интеллектом, встречающимся практически в каждом приложении, и сильным искусственным интеллектом, который в конечном счете, возможно, удастся получить, рассматривались ранее, в разделе “Аргумент китайской комнаты”.

Искусственный интеллект позволяет людям осуществлять и другие виды интеллектуального ввода.

Пример распознавания голоса избит уже до синяков, но остается одной из главных общепринятых методик интеллектуального ввода. Но даже если искусственный интеллект испытывает пока недостаток в полном диапазоне смыслов,

как описано в главе 4, он уже может поддерживать широкое разнообразие невербальных способов интеллектуального ввода данных. В качестве визуальных способов можно упомянуть распознавание лица владельца или выявление угрозы на основании выражения лиц.

Однако ввод может также подразумевать мониторинг. Вполне можно отслеживать показатели состояния организма пользователя для выявления потенциальных проблем. Фактически искусственный интеллект может использовать огромное количество методов интеллектуального ввода данных, большинство из которых еще даже не изобретено.

В настоящее время приложения обычно обеспечивают только три первых уровня дружелюбия.

По мере роста интеллекта искусственный интеллект становится немаловажным фактором действия *дружественного искусственного интеллекта* (Friendly Artificial Intelligence — FAI), совместимого с *общим искусственным интеллектом* (Artificial General Intelligence — AGI), оказывающим положительное влияние на человечество. У искусственного интеллекта есть задачи, но они могут не иметь ничего общего с человеческой этикой, и это потенциальное различие вызывает тревогу сегодня. Интеллект FAI обладает логикой, гарантирующей соответствие задач искусственного интеллекта человеческим, — это три закона роботехники Айзека Азимова (<https://www.auburn.edu/~vestmon/robotics.html>), книги которого обсуждаются подробнее в главе 12. Однако многие утверждают, что эти три закона — только хорошая отправная точка (<http://theconversation.com/after-75-years-isaac-asimovsthree-laws-of-robotics-need-updating-74501>) и что необходимы дополнительные гарантии.



СОВЕТ

Конечно, все это обсуждение законов и этики может оказаться весьма сомнительным и трудным в определениях. Но вот простой пример поведения FAI: он откажется разглашать личную информацию пользователя, если у спрашивающего нет необходимости ее знать.

Фактически FAI может пойти даже далее ввода данных человеком, поиска соответствия по шаблону, защиты его личной информации и уведомления пользователя о возможном вреде прежде, чем он опубликует некую информацию где-нибудь. Дело в том, что искусственный интеллект может существенно изменить точку зрения людей на сами приложения и взаимодействие с ними.

Автоматическое исправление

Люди постоянно все исправляют. И вовсе не потому, что нечто неправильно.

Это скорее вопрос постоянного усовершенствования (или по крайней мере попытки сделать это). Даже когда людям удастся достичь необходимого уровня правильности, в некий момент появляется новый опыт, подвергающий этот уровень сомнению, поскольку теперь у человека есть дополнительные данные для суждения о том, что есть правильно в конкретной ситуации. Для полного подражания человеческому интеллекту у искусственного интеллекта также должна быть эта возможность — постоянно исправлять достигнутые результаты, даже когда эти результаты положительны. В следующих разделах обсуждается проблема правильности, а также исследуется, почему автоматизированные исправления иногда терпят неудачу.

Виды исправлений

Когда большинство людей думает об искусственном интеллекте и исправлениях, они имеют в виду системы проверки правописания и грамматики. Человек совершает ошибку (или так полагает искусственный интеллект), и искусственный интеллект эту ошибку исправляет, чтобы документ получился настолько точным, насколько возможно. Конечно, люди делают много ошибок, поэтому идея использовать искусственный интеллект таким образом, безусловно, хороша.

Исправления могут принимать разные формы и необязательно означать, что ошибка произошла или произойдет в будущем. Например, автомобиль может помогать водителю, регулярно корректируя маршрут. Водитель вполне может и сам хорошо вести машину, но искусственный интеллект способен вносить микро-коррекции, чтобы помочь водителю оставаться в границах безопасности.

Рассмотрим такой случай: перед автомобилем с искусственным интеллектом внезапно останавливается другой автомобиль из-за оленя на дороге. Водитель первого автомобиля не совершил никакой ошибки. Однако искусственный интеллект успел среагировать быстрее и действовал так, чтобы остановить автомобиль настолько быстро и безопасно, насколько это возможно, чтобы избежать столкновения с автомобилем, остановившимся перед ним.

Преимущества автоматических исправлений

Когда искусственный интеллект видит необходимость в исправлении, он может попросить у человека разрешение внести исправление или сделать это автоматически. Например, некто использует распознавание речи для набора документа и делает грамматическую ошибку. Искусственный интеллект должен просить разрешение, прежде чем делать замену, поскольку человек,

возможно, именно это слово и имел в виду, а искусственный интеллект, возможно, неправильно понял сказанное.

Но иногда критически важно, чтобы искусственный интеллект обеспечивал процесс принятия решений, достаточно надежный для того, чтобы выполнять исправления автоматически. Например, в предыдущем случае экстренного торможения у искусственного интеллекта нет времени на просьбы о разрешении; он должен тормозить немедленно, или произойдет авария. У автоматических исправлений с искусственным интеллектом есть вполне определенная область применения — когда критически важные решения следует принимать быстро и искусственный интеллект достаточно надежен.

Почему автоматические исправления не срабатывают

Как упоминалось в разделе “Аргумент китайской комнаты”, фактически искусственный интеллект ничего понять не может. Без понимания этого нет никакой возможности справиться с непредвиденным обстоятельством. В данном случае под непредвиденным обстоятельством следует понимать событие, сценарий которого не предусмотрен заранее, когда искусственный интеллект не может накопить дополнительные данные или положиться на другие средства механического решения.

Человек может решить проблему, поскольку понимает ее суть, и обычно достаточно многие из окружающих событий подчинены шаблону, помогающему выработать решение. Кроме того, человек способен на новаторские и творческие решения, когда не подходит ни одно из очевидных. В настоящее время искусственный интеллект испытывает недостаток и в новаторском, и в творческом потенциале, поэтому он находится в невыгодном положении при решении задач в некоторых областях.

Чтобы получить представление об этой проблеме, рассмотрим случай с программой поиска опечаток. Человек вводит совершенно правильное слово, отсутствующее в словаре, используемом искусственным интеллектом при внесении исправлений. Искусственный интеллект нередко заменяет слово похожим, но другим словом. Даже после того, как человек, проверив документ и исправив неправильное слово, добавляет нужное в словарь, искусственный интеллект все еще склонен делать ту же ошибку. Например, искусственный интеллект может считать аббревиатуру *CPU* (процессор) отличной от *сри*, из-за разницы в регистре символов. Человек бы увидел, что эти два сокращения означают одно и то же самое и что сокращение во втором случае правильно, но, возможно, должно быть в верхнем регистре.

Создание рекомендаций

Рекомендация — это вовсе не команда. Хотя некоторые люди, кажется, совершенно упускают этот момент, рекомендация — это просто идея, выдвинутая как возможное решение задачи. Рекомендация подразумевает возможность и других решений, а принятие рекомендации не означает ее неизбежной реализации. Фактически рекомендация — это только идея; она может даже не сработать. Конечно, в совершенном мире все рекомендации хороши, по крайней мере возможные решения правильного вывода, что редко имеет место в реальном мире. В следующих разделах описана природа рекомендаций применительно к искусственному интеллекту.

Получение рекомендаций на основании предыдущих действий

Наиболее распространенный способ использования искусственного интеллекта для создания рекомендаций подразумевает сбор данных о действиях при предыдущих событиях и выработку новых рекомендаций на их основании.

Например, некто покупает Half-Baked Widget каждый месяц в течение трех месяцев. В начале четвертого месяца он решает, что имеет смысл купить другой. Действительно умный искусственный интеллект мог бы порекомендовать и нужное время месяца. Например, если пользователь делает покупку между третьим и пятым днем месяца в течение первых трех месяцев, он платит в третий день месяца, а затем переходит на что-то еще после пятого дня.

Решая задачи, люди отбрасывают огромные количества подсказок. В отличие от людей искусственный интеллект фактически обратит внимание на каждую из этих подсказок и запишет их единообразным способом. Во многих случаях единообразный набор данных о действиях позволяет искусственному интеллекту давать рекомендации на основании предыдущих действий с высокой степенью точности.

Получение рекомендаций на основании групп

Другой наиболее популярный способ создания рекомендаций основан на принадлежности к группе. В данном случае принадлежность к группе не должна быть формальной. Группа может состоять из свободной ассоциации людей, обладающей некоторой незначительной общей потребностью или деятельностью. Например, дровосеки, владельцы магазинов и диетологи все могут покупать детективные романы. Даже при том, что ничего иного общего у них нет, они даже не живут рядом, тот факт, что все они имеют детективные романы, делает их частью группы. Искусственный интеллект легко может определить шаблон, который люди могли бы и упустить. Таким образом, искусственный

интеллект способен давать хорошие коммерческие рекомендации на основании принадлежности даже к весьма свободно связанным группам.

Группы могут объединять как эфемерные, так и временные связи. Например, все люди, летевшие рейсом 1982 из Хьюстона в определенный день, могут быть группой.

Между этими людьми может не быть никакой связи вообще, за исключением того, что они оказались на некоем конкретном рейсе. Однако, зная эту информацию, искусственный интеллект может (осуществив дополнительную фильтрацию) выявить людей в пределах этого рейса, которым нравятся детективы. Дело в том, что искусственный интеллект может давать хорошие рекомендации на основании принадлежности к группам, даже когда группу очень трудно (если вообще возможно) выявить с человеческой точки зрения.

Получение неправильных рекомендаций

Любой, кто делал покупки в интернет-магазине, знает, что веб-сайты зачастую предоставляют рекомендации на основании различных критериев, таких как предыдущие покупки. К сожалению, эти рекомендации редко бывают правильными, поскольку у базового искусственного интеллекта отсутствует понимание. Когда некто делает покупку всей своей жизни и приобретает Супер-Мега-Штуку, он, вероятно, заранее знает, что такая покупка действительно бывает только раз в жизни, поскольку крайне маловероятно, что такого понадобится два. Однако искусственный интеллект этого факта не понимает. Поэтому, если программист не создаст специально правило, гласящее, что Супер-Мега-Штуку покупают только раз, искусственный интеллект может решить продолжать рекомендовать этот товар, поскольку его продажи, понятно, низки. Следуя второму правилу о продвижении товаров с низкими продажами, искусственный интеллект ведет себя согласно требованию, предусмотренному разработчиком, но рекомендации, которые он дает, абсолютно неправильны.

Помимо ошибок из-за правил или логических ошибок, рекомендации искусственного интеллекта могут быть ошибочны из-за проблем с данными. Например, GPS-навигатор может дать рекомендацию о наилучшем маршруте на основании имеющихся данных. Однако состояние дорог может сделать предложенный путь непроходимым, поскольку одна из дорог закрыта на ремонт. Конечно, многие приложения GPS учитывают состояние дорог, но они иногда не принимают во внимание другие проблемы, такие как внезапное изменение в ограничении скорости или погодные условия, делающие некий участок дороги опасным. Люди способны преодолевать недостаток данных новаторски, используя менее популярные дороги или поняв смысл надписи об объезде.

Когда искусственный интеллект сумеет преодолеть проблемы логики, правил и данных, он все равно будет иногда давать плохие рекомендации, поскольку он не осознает корреляции между определенными наборами данных, как человек. Например, искусственный интеллект не сможет предложить цвет краски после того, как человек купил комбинацию батарей и гипсокартона для ремонта дома. Необходимость покрасить гипсокартон и все вокруг после ремонта очевидна для человека, поскольку у него есть эстетический вкус, а у искусственного интеллекта его нет. Человек легко находит корреляцию между различными элементами, но она не очевидна для искусственного интеллекта.

Учет ошибок искусственного интеллекта

Прямая ошибка (outright error) происходит в результате ввода исходных данных, неправильных в любой форме. Результат не содержит правильный ответ на вопрос. Примеры подобных ошибок искусственного интеллекта найти совсем не трудно. Недавно сеть BBC News опубликовала сообщение о том, что изображения с различием в один пиксель ввели в заблуждение специальный искусственный интеллект (см. статью на <http://www.bbc.com/news/technology-41845878>). Вы можете узнать больше о влиянии враждебных атак на искусственный интеллект по адресу <https://blog.openai.com/adversarial-example-research/>. Статья лаборатории Касперского по адресу <https://www.kaspersky.com/blog/aifails/18318/> дает дополнительную информацию о ситуациях, когда искусственный интеллект не в состоянии дать правильный ответ. Дело в том, что у искусственного интеллекта все еще высок процент ошибок при некоторых обстоятельствах, и разработчики, работающие с искусственным интеллектом, обычно не могут с уверенностью указать их причину.

Ошибки искусственного интеллекта возникают по множеству причин. Но как уже упоминалось в главе 1, искусственный интеллект не обязан подражать всем семи формам человеческого интеллекта, поэтому ошибки не только возможны, но и неизбежны. Большая часть главы 2 сосредоточена на данных и их воздействии на искусственный интеллект, когда данные в некотором роде неверны. В главе 3 также упоминается, что даже у используемых искусственным интеллектом алгоритмов есть пределы. Глава 4 указывает, что искусственному интеллекту недоступны ни то же самое количество, ни те же самые типы чувств, что и человеку.

В статье на сайте TechCrunch по адресу <https://techcrunch.com/2017/07/25/artificial-intelligence-is-not-as-smart-as-you-or-elon-musk-think/> указывается множество, по-видимому, невозможных задач, которые искусственный интеллект пытается сегодня решать методами “в лоб”, а не используя нечто, даже близко напоминающее мышление.

Все более очевидной становится главная проблема — корпорации зачастую просто замалчивают или даже игнорируют проблемы искусственного интеллекта. Акцент смещается на использование искусственного интеллекта для снижения стоимости и улучшения продуктивности, что может быть и недостижимо. Компания Bloomberg опубликовала статью <https://www.bloomberg.com/news/articles/2017-06-13/the-limits-of-artificial-intelligence>, посвященную подробному обсуждению этой проблемы. Одним из наиболее интересных недавних примеров сущности корпораций, зашедших слишком далеко с искусственным интеллектом, является бот Tay от Microsoft (см. <https://www.geekwire.com/2016/microsoft-chatbot-tay-mit-technology-fails/>), который “набрался ума” из комментариев расиста, женофоба и любителя порнографии, а затем начал публиковать оскорбительные и провокационные сообщения через свою учетную запись в Twitter.



СОВЕТ

Ценнейший факт, который стоит извлечь из этого раздела, не в том, что искусственный интеллект ненадежен или неприменим. Фактически рядом с хорошо образованным человеком искусственный интеллект может выработать свой человеческий образ быстро и эффективно. Искусственный интеллект способен помочь людям избегать наиболее распространенных или повторяющихся ошибок. В некоторых случаях ошибки искусственного интеллекта способны даже вносить немного юмора в повседневную жизнь. Однако искусственный интеллект не думает, и он не сможет заменить людей во многих ситуациях сегодня. Он работает лучше, когда решение контролирует человек или когда окружающая обстановка достаточно статична, чтобы предсказуемость хороших результатов была высока (все хорошо, пока человек не решается обмануть искусственный интеллект).



Глава 6

Автоматизация наиболее популярных процессов

В ЭТОЙ ГЛАВЕ...

- » Применение искусственного интеллекта для нужд человека
- » Повышение эффективности производства
- » Разработка динамических протоколов безопасности с использованием искусственного интеллекта

В главе 5 рассматривалось применение искусственного интеллекта в *приложении*, когда человек взаимодействует с искусственным интеллектом неким осмысленным способом, даже если не осознает присутствия искусственного интеллекта. Задача заключается в том, чтобы помочь людям сделать нечто быстрее, проще и эффективнее либо удовлетворить некоторую другую потребность. *Процесс*, в котором участвует искусственный интеллект, зависит от конкретной задачи, поскольку искусственный интеллект должен либо помочь человеку в решении его задачи, либо работать самостоятельно, без прямого вмешательства человека. В первом разделе этой главы рассматривается, как процессы помогают людям. Возможно, это покажется скучным (только подумайте обо всех неприятностях, которые случаются, когда людям становится скучно), но здесь процесс искусственного интеллекта рассматривается именно с точки зрения скуки людей.

Один из самых давних процессов, применяющих искусственный интеллект, — это промышленное производство. Вспомните обо всех роботах, которыми теперь полны заводы во всем мире. Контролируемая искусственным интеллектом автоматика не только заменяет людей, она и делает труд людей безопаснее, выполняя задачи, обычно считающиеся опасными. Как ни странно, одной из наиболее серьезных причин несчастных случаев на производстве и многих других проблем является скука (см. <https://thepsychologist.bps.org.uk/volume-20/edition-2/boredom-work>). Роботы могут выполнить монотонные действия постоянно, и им это не надоест.

На случай, если вы еще достаточно не скучали, можете прочитать об этом в третьем разделе данной главы, в котором обсуждаются некоторые из самых новых областей, в которых искусственный интеллект просто великолепен, — как сделать окружающую обстановку во всех смыслах более безопасной. Фактически только в автомобильной промышленности можно найти бесчисленные способы применения искусственного интеллекта для улучшений (см. <https://igniteoutsourcing.com/publications/artificial-intelligence-in-automotive-industry/>).



ЗАПОМНИ

Основная идея этой главы — искусственный интеллект хорош в производственных процессах, особенно в тех, в которых люди обычно скучают и допускают ошибки в отличие от искусственного интеллекта. Конечно, искусственный интеллект не может устранить каждый источник потери эффективности, включая незаинтересованность и проблемы безопасности. С одной стороны, люди могут решить игнорировать помощь искусственного интеллекта, но природа ограничений имеет куда более глубокие корни, чем этот. Как обсуждалось в предыдущих главах (особенно в главе 5), искусственный интеллект ничего не понимает; он не может предоставить творческое или новаторское решение проблем, поэтому некоторые проблемы искусственным интеллектом не разрешимы, независимо от того, сколько усилий некто вложил в его создание.

Решения для скуки

Опросы зачастую показывают, что люди только думают, что они чего-то хотят, а на самом деле не хотят, но такие опросы все же полезны. Когда опрашивали недавних выпускников колледжа, о какой жизни они мечтают, ни один из них не сказал, что мечтает о скучной жизни (см. https://www.huffingtonpost.com/paul-raushenbush/what-kindof-life-do-you_b_595594.html). Фактически опросить можно практически любую группу, и никто не будет мечтать

о скуке. Большинству людей (сказав “всем”, я, вероятно, нарвусь на лавину электронных писем с опровержениями) скука не нравится. В некоторых случаях искусственный интеллект может, работая с людьми, делать их жизнь интереснее, по крайней мере для человека. Ниже обсуждаются решения для человеческой скуки, которые может обеспечить искусственный интеллект (а также те, которые не может).

Как сделать задачи интереснее

У любого занятия, персонального или группового, есть определенные привлекательные характеристики, из-за которых люди желают в нем участвовать. Вполне очевидно, что некоторые занятия, такие как забота о собственных детях, бесплатны, но удовлетворение от них может быть невероятно высоким. Аналогично работу бухгалтера могут оплачивать весьма хорошо, но морального удовлетворения в этой работе немного. Различные опросы (такие, как на <https://www.careercast.com/jobs-rated/jobs-rated-report-2016-ranking-200-jobs>) и статьи (такие, как на <http://www.nytimes.com/2010/09/12/jobs/12search.html>) свидетельствуют о балансе денег и удовлетворения, но их результаты зачастую выглядят сомнительно, поскольку основания для таких выводов неоднозначны. Однако большинство этих источников сходится во мнении, что после того, как человек получает определенную сумму денег, ключом для поддержания интереса к занятию становится удовлетворение (причем независимо от занятия). Конечно, практически невозможно выяснить, в чем заключается удовлетворение от работы, но интерес остается в списке первым. Удовлетворение от интересного занятия всегда будет более вероятным.



СОВЕТ

Проблема не в обязательной смене работы, просто, сделав работу немного интереснее, можно избежать скуки. Искусственный интеллект может эффективно в этом помочь, устранив из задач повторяющиеся действия. Однако такие примеры, как Alexa от Amazon и Home от Google, действительно предоставляют другие варианты. Чувство одиночества, которое может проникнуть в дом, на рабочее место, в автомобиль или в другое место, является мощнейшим источником скуки. Когда люди начинают чувствовать себя одинокими, начинается депрессия, и скука зачастую — только начало. Приложения, использующие интерфейс Alexa (см. <https://developer.amazon.com/>) или Actions на API Google (см. <https://developers.google.com/actions/>) для имитации общения с человеком соответствующего типа, может улучшить ощущения на рабочем месте. Однако важнее всего то, что разработка интеллектуальных интерфейсов такого типа способна помочь людям быстро выполнять различные повседневные

задачи, такие как поиск информации и взаимодействие с интеллектуальными устройствами, а не только выключение света (см. <https://www.imore.com/how-control-your-lights-amazon-echo> и https://store.google.com/product/google_home).

Как помочь работать эффективнее

У большинства людей, по крайней мере у думающих, есть некоторое представление о том, чего они хотели бы от искусственного интеллекта в плане улучшения их жизни за счет устранения задач, которые они не хотят выполнять. Недавний опрос выявил некоторые из наиболее интересных способов, которыми искусственный интеллект может это сделать: <https://blog.devolutions.net/2017/10/october-poll-results-which-tasks-in-your-job-would-you-like-to-be-automated-by-ai>. Большинство из них вполне обыденны, но один обращает на себя внимание — предлагается выяснять, когда жена чувствует себя несчастной, и посылать ей цветы. Это, скорее всего, не сработает, но сама идея интересна.

Дело в том, что люди, вероятно, предложат самые интересные идеи о том, как создать искусственный интеллект, специально предназначенный для решения именно их задач. В большинстве случаев серьезные идеи будут хороши и для других пользователей. Например, автоматическая система эксплуатационной поддержки может применяться во многих отраслях. Если некто обеспечит обобщенный интерфейс с программируемой внутренней структурой приложения для применения индивидуальных систем вопросов, то искусственный интеллект сможет сэкономить пользователям много времени и гарантировать правильные последующие действия, обеспечив грамотный опрос и запись необходимой информации.

Как искусственный интеллект борется со скукой

Скука бывает разной, и люди относятся к ней по-разному.

Бывает, скука наступает от отсутствия необходимых ресурсов, знаний или других вещей. Другой вид скуки наступает из-за незнания, что делать дальше. Искусственный интеллект может помочь с первым видом скуки, но не справится со вторым. В этом разделе рассматривается первый вид скуки. (В следующем разделе рассматривается второй.)



ЗАПОМНИ

Доступ к ресурсам всех видов помогает справиться со скукой, позволив людям творить без необходимости приобретать нужные материалы обычными способами. Вот некоторые из способов, которыми искусственный интеллект способен упростить доступ к ресурсам.

- » Поиск необходимых элементов по Интернету.
- » Автоматический заказ необходимых элементов.
- » Осуществление сенсорного сбора данных и другого мониторинга.
- » Управление данными.
- » Выполнение обыденных или повторяющихся задач.

Как искусственный интеллект не может бороться со скукой

Как уже говорилось в предыдущих главах, особенно в главах 4 и 5, искусственный интеллект не является ни творческим, ни интуитивным. Поэтому запрос к искусственному интеллекту придумать, что вам делать, вряд ли приведет к удовлетворительным результатам. Некто может запрограммировать искусственный интеллект так, чтобы он отследил десять ваших любимых занятий, а затем выбрал одно из них наугад, но результат не будет удовлетворительным, поскольку искусственный интеллект не может учесть все аспекты, включая ваше текущее настроение и увлечение. Фактически даже в самом лучшем случае искусственный интеллект не сможет общаться с вами способом, обеспечивающим удовлетворительный результат любого вида.

Искусственный интеллект не может также вас мотивировать. Вспомните, как иногда приходят друзья, чтобы вас мотивировать (или вы приходите к ним). Друг фактически полагается на комбинацию внутриличностного (сочувствие, как если бы он оказался в вашей ситуации) и межличностного знания (предложение творческих идей, чтобы получить от вас положительный эмоциональный отклик). У искусственного интеллекта нет первого вида знаний и есть лишь чрезвычайно ограниченный объем второго, как описано в главе 1. Следовательно, искусственный интеллект не сможет развеять вашу скуку, используя мотивационные методики.



СОВЕТ

Так или иначе скука — это не всегда плохо. Многие недавние исследования показали, что скука фактически порождает творческую мысль (например, см. <https://www.fastcompany.com/3042046/the-science-behind-howboredom-benefits-creative-thought>), что людям еще предстоит осмыслить. После просмотра бесчисленных статей о том, как искусственный интеллект собирается устранить рутинные работы, складывается впечатление, что деятельность, на которую он претендует зачастую скучна и не оставляет людям никакого времени на творчество. Даже сегодня люди вполне могут найти продуктивную творческую работу, если они действительно об этом подумают. Статья “7 удивительных фактов о творческом потенциале согласно науке” (7 Surprising Facts About Creativity, According To Science) (<https://www.fastcompany.com>).

com/3063626/7-surprising-factsabout-creativity-according-to-science) фактически обсуждает роль таких связанных со скукой действий, как мечтание в усилении творческого потенциала. В будущем, если люди действительно хотят достичь звезд и сделать другие удивительные вещи, творческий потенциал станет крайне важным, поэтому тот факт, что искусственный интеллект не способен развеять вашу скуку, фактически будет очень хорош.

Работа в промышленных условиях

В любых промышленных условиях, вероятно, будет небезопасно, независимо от количества вложенных в решение этой проблемы денег, времени и усилий. Можно легко найти такие статьи, как <http://www.safetyandhealthmagazine.com/articles/14054-common-workplace-safety-hazards>, описывающие семь наибольших угроз безопасности в промышленных условиях. Хотя люди и решают большинство этих проблем, скука осложняет фактическую окружающую обстановку, в которой люди работают, и служит причиной очень многих проблем. В следующих разделах описано, как автоматизация способна помочь людям жить дольше и счастливее.

Уровни автоматизации

Автоматизация в промышленных условиях намного старше, чем вы могли бы подумать. Отправной точкой автоматизации некоторые считают сборочные линии Генри Форда (см. <https://www.history.com/this-day-in-history/fords-assembly-line-starts-rolling>). Фактически автоматизация началась в 1104 году нашей эры в Венеции (см. <https://www.mouser.com/applications/factory-automation-trends/>), где 16 тысяч рабочих были способны построить весь военный корабль за один день. Американцы повторили этот подвиг с современными судами во время второй мировой войны (см. <https://www.nps.gov/nr/travel/wwiibayarea/shipbuilding.htm>). Таким образом, автоматизация существовала давно и везде.

Но чего не существовало давно и везде, так это искусственного интеллекта, который может реально помочь людям в пределах процесса автоматизации. Сегодня обычно человек-оператор сначала думает, как выполнить некую работу, и ставит задачу, а затем перепоручает ее компьютеру. Одним из недавних примеров является *автоматизация роботизированных процессов* (Robot Process Automation — RPA), позволяющая человеку научить программное обеспечение действовать в поле человека при работе с приложениями (см. <https://www.valamis.com/blog/the-tools-of-the-future-today-what-is-robotic-process-automation-artificial-intelligence-and-machine-learning>). Этот процесс

отличается от простого сценария, как при использовании VBA (Visual Basic for Applications) в Microsoft Office, RPA — это не конкретное приложение и не требует программирования. Многих удивит тот факт, что фактически есть десять уровней автоматизации, девять из которых могут полагаться на искусственный интеллект. Выбираемый вами уровень зависит от вашего конкретного случая.

1. Человек-оператор создает задачу и передает ее компьютеру для реализации.
2. Искусственный интеллект помогает человеку определить параметры задачи.
3. Искусственный интеллект устанавливает наилучшие параметры задачи, а затем позволяет человеку принять или отклонить рекомендацию.
4. Искусственный интеллект определяет параметры и на их основании определяет последовательность действий, а затем предоставляет список действий человеку для принятия или отклонения конкретных действий до их реализации.
5. Искусственный интеллект определяет параметры, последовательность действий, формирует задачу, а затем запрашивает у человека одобрение, прежде чем передать задачу компьютеру.
6. Искусственный интеллект автоматически создает задачу и передает ее в очередь задач компьютера, а человек-оператор выступает в роли посредника на случай, если выбранная задача требует доводки перед фактической реализацией.
7. Искусственный интеллект создает и реализует задачу, а затем уведомляет человека-оператора о сделанном на случай, если задачу потребуется откорректировать или аннулировать.
8. Искусственный интеллект создает и реализует задачу, уведомляя человека о сделанном, только если он спросит.
9. Искусственный интеллект создает и реализует задачу безо всякой обратной связи, если только не потребуется вмешательство человека из-за ошибки или непредвиденных результатов.
10. Искусственный интеллект не ждет приказа от человека, он сам выявляет потребность в задаче и инициализирует ее выполнение. Искусственный интеллект предоставляет обратную связь только тогда, когда требуется вмешательство человека, например при серьезной ошибке. С ошибками определенного уровня и непредвиденными результатами искусственный интеллект справляется самостоятельно.

Больше чем просто роботы

Размышляя об автоматизации промышленности, большинство людей полагает, что все делают роботы. Фактически общество сейчас находится в состоянии четвертой промышленной революции; была революция пара, массовое производство, автоматизация, а теперь — коммуникация (см. <https://www.engineering.com/ElectronicsDesign/ElectronicsDesignArticles/ArticleID/8379/New-Chips-are-Bringing-Factory-Automation-into-the-Era-of-Industry-40.aspx>). Для эффективного выполнения задачи искусственный

интеллект запрашивает информацию из источников всякого рода. Из этого следует, что чем более подробную информацию промышленная установка может получить из всякого рода источников, тем лучше ее искусственный интеллект сможет работать (если, конечно, этими данными правильно распорядиться). С учетом этого промышленные установки всех видов полагаются теперь на *индустриальный механизм связи* (Industrial Communication Engine — ICE), чтобы координировать взаимосвязь между всеми столь разными источниками, в которых нуждается искусственный интеллект.

Роботы действительно выполняют большую часть реальной работы на индустриальном производстве, однако нужны также сенсоры для оценки потенциальных рисков, таких как штормы. Но для поддержания эффективности производства координация становится все более и более важным фактором. Например, грузовики с сырьем должны поступать своевременно, а другие грузовики, отвозящие готовые товары, должны задействоваться по мере необходимости — эта задача крайне важна для эффективного использования складских помещений. Чтобы повысить эффективность работ, искусственный интеллект должен знать о техническом состоянии всего оборудования, чтобы гарантировать его наилучшее обслуживание, причем в то время, когда в этом оборудовании меньше всего потребности.

Искусственный интеллект должен также учитывать такие проблемы, как цена ресурса. Возможно, получится что-то выиграть, если использовать определенное оборудование в вечерние часы, когда цена электричества ниже.

Не автоматизацией единой

Первые примеры фабрик без людей относились к таким весьма специфическим производствам, как выпуск микросхем, для которых требовалась абсолютная чистота окружающей среды. Однако с тех пор автоматизация распространилась. Из-за опасностей для здоровья людей или очень высокой стоимости использования человека при выполнении определенных видов промышленных задач сегодня можно найти множество самых обычных фабрик, не требующих никакого человеческого вмешательства вообще (см. примеры на <https://singularityhub.com/2010/02/11/no-humans-just-robots-amazing-videos-of-the-modern-factory/>).



ЗАПОМНИ

Многие технологии до некоторой степени эффективно разрешают все связанные с производством задачи без человеческого вмешательства (см. примеры на <https://www.automationmag.com/opinion/thought-leaders/5248-top-10-industrial-automation-trends-automationdirect>). Дело в том, что в конечном счете обществу придется искать задачи и кроме рутинных фабричных работ, чтобы занять население.

Безопасность окружающей обстановки

Одной из наиболее обсуждаемых ролей искусственного интеллекта, помимо автоматизации производственных задач, является обеспечение безопасности людей различными способами. В таких статьях, как <https://futurism.com/7-reasons-you-should-embrace-not-fear-artificial-intelligence/>, описаны опасные ситуации, в которых искусственный интеллект выступает в качестве посредника, принимая на себя удар, который иначе достался бы людям. Безопасность бывает разной. Да, искусственный интеллект может сделать работу в различных обстановках более безопасной, но он также помогает сделать окружающую обстановку более здоровой, а также снизить риски, связанные с самыми обычными занятиями, включая серфинг в Интернете. В следующих разделах содержится краткий обзор путей, которыми искусственный интеллект может повышать безопасность окружающей обстановки.

Роль скуки в происшествиях

От вождения (см. <http://healthland.time.com/2011/01/04/bored-drivers-most-likely-to-have-accidents/>) до работы (см. <http://www.universaldrugstore.com/news/general-healthnews/boredom-increases-accidents-at-workplace/>) скука повышает вероятность происшествий всех видов. Фактически, когда выполнение некой задачи требует постоянного внимания, она действует усыпляюще, и результат редко оказывается хорошим. Проблема настолько серьезна и существенна, что по этой теме можно найти множество статей, таких как “Modelling human boredom at work: mathematical formulations and a probabilistic framework” (Моделирование человеческой скуки на работе: математические формулировки и вероятностная среда возникновения) (<http://www.emeraldinsight.com/doi/full/10.1108/17410381311327981>). Произойдет ли происшествие фактически (или просто создастся опасная ситуация), зависит от случая.

Представьте реальную разработку алгоритмов, позволяющих определить вероятность происшествий, происходящих из-за скуки при определенных обстоятельствах.

Как искусственный интеллект помогает избежать проблем безопасности

Никакой искусственный интеллект не сможет предотвратить происшествий из-за человеческих факторов, таких как скука. В лучшем случае, когда люди действительно решают следовать правилам, выработанным с использованием искусственного интеллекта, последний может только помочь избежать

возможных проблем. В отличие от роботов Азимова, сейчас нет никакой защиты тремя законами, работающими в любой обстановке; люди должны сами заботиться о своей безопасности. С учетом этого искусственный интеллект может оказать помощь следующими способами.

- » Предложить ротацию задач (на рабочем месте, в автомобиле или даже дома), чтобы дело осталось интересным.
- » Проконтролировать продуктивность человека и предложить лучшее время отдыха от усталости или других факторов.
- » Помочь в решении задач, в которых выгодна комбинация сообразительности человеческого интеллекта и скорости реакции искусственного.
- » Улучшить человеческие возможности по обнаружению вероятных источников опасности, чтобы они стали очевидней.
- » Устранить повторяющиеся задачи, чтобы люди меньше утомлялись и участвовали только в интересных аспектах своей работы.

Почему искусственный интеллект не может устранить проблемы безопасности

Обеспечение полной безопасности подразумевает способность предвидеть будущее. Поскольку будущее неизвестно, потенциальный риск для людей в любой момент времени также остается неизвестным, ведь неожиданные ситуации могут возникнуть когда угодно. Неожиданная ситуация — это то, о чем и не предполагали заранее разработчики некой стратегии безопасности. Люди имеют огромный опыт по части поиска новых способов попасть в затруднительное положение, частично из-за любопытства и творческого характера. Методы предотвращения опасности искусственным интеллектом следует искать в человеческой натуре, поскольку люди любознательны; мы хотим увидеть, а что будет, если что-то попробуем? Как правило, получается глупость.

Непредсказуемые ситуации — не единственная проблема для искусственного интеллекта. Даже если кому-то удалось бы найти все возможные ситуации, способные создать опасность для человека, то потребуется астрономическая вычислительная мощь, чтобы выявить и предотвратить подобное событие. Искусственный интеллект работал бы так медленно, что его ответ всегда опаздывал бы настолько, что был бы бесполезным. Следовательно, разработчикам оборудования для обеспечения безопасности фактически требуется такой искусственный интеллект, который обеспечит лишь необходимый уровень безопасности, который будет готов иметь дело с наиболее вероятными ситуациями.



Глава 7

Применение искусственного интеллекта в медицине

В ЭТОЙ ГЛАВЕ...

- » Более эффективный мониторинг пациентов
- » Помощь людям в различных ситуациях
- » Анализ потребностей пациентов
- » Решение хирургических и других медицинских задач

Медицина сложна. И не случайно для обучения врача может потребоваться 15 или более лет в зависимости от специальности (см. <http://work.chron.com/long-become-doctor-us-7921.html>). К тому времени, когда система образования упакует врача достаточным количеством информации почти до краев, большинство других людей уже трудится лет 11 (с учетом, что большинство останавливается на уровне среднего специального образования или степени бакалавра). Тем временем новые технологии, подходы и так далее — все тайно сговорилось сделать эту задачу еще сложнее. В один

прекрасный момент любой человек просто не может стать еще более опытным даже в узкой специальности. Именно поэтому незаменимому человеку нужна последовательная, логичная и беспристрастная помощь в виде искусственного интеллекта. Процесс начинается с помощи в мониторинге пациентов (как описано в первом разделе этой главы) такими способами, которые для человеческого персонала просто невозможны, поскольку количество анализов очень высоко и выполнять их требуется в определенном порядке и определенным способом. Это критически важно, и вероятность ошибки недопустима.

К счастью, сегодня у людей куда больше возможностей для самостоятельного решения многих связанных с медициной задач, чем когда-либо прежде. Например, игры позволяют пациентам выполнять некоторые из терапевтических задач самостоятельно, но все же им нужно руководство от приложения, гарантирующего выполнение упражнений способом, наилучшим образом подходящим для восстановления их здоровья. Улучшенные эндопротезы и другие медицинские средства также позволяют людям стать менее зависимыми от профессиональной помощи. Во втором разделе этой главы описано, как искусственный интеллект может помочь людям в их медицинских потребностях.

Очень трудно, если вообще возможно, починить устройство, не зная, как оно работает. Точно так же нельзя игнорировать проведение анализов при диагностике болезни. Проведение различных анализов помогает врачу выявить конкретную проблему и упростить ее решение. Сегодня врач вполне может снабдить пациента устройством мониторинга и контролировать его состояние дистанционно, а результаты направить искусственному интеллекту для их анализа и диагностики проблемы, причем все это за один визит пациента в кабинет врача (он нужен, чтобы прикрепить устройство мониторинга). Такие мониторы, как глюкометры, пациенты вполне могут купить сами, и тогда посещение кабинета врача становится вообще ненужным. Хотя в третьем разделе этой главы не затрагивается даже часть всех анализирующих устройств, вы хотя бы получите о них представление.

Конечно, требуется некоторое вмешательство, чтобы подвергнуть пациента хирургическим или другим процедурам (как описано в четвертом разделе этой главы). Роботизированные решения иногда могут выполнять такие задачи даже лучше, чем врач. В некоторых случаях помощь робота повышает эффективность врача и позволяет ему сосредоточить внимание на том, с чем может справиться только человек. Различные технологии упрощают диагностику, проводя ее быстрее и точнее. Например, искусственный интеллект может помочь врачу диагностировать рак на начальной стадии куда быстрее, чем если бы врач делал это сам.

Носимый терапевтический монитор

Медик не всегда может сказать, что происходит с пациентом, просто прослушав его сердце, проверив пульс или сделав анализ крови. Тело не всегда подает однозначные сигналы, которые позволяют медику узнать их причину сразу. Кроме того, некоторые показатели организма, такие как сахар в крови, изменяются со временем, поэтому необходим постоянный их мониторинг. Посещение кабинета врача при каждой необходимости проверить один из этих параметров — дело трудоемкое и не всегда полезное. Прежние методы определения некоторых медицинских характеристик требовали внешнего ручного вмешательства со стороны пациента, а этот процесс приводит к ошибкам, в лучшем случае. По этим и многим другим причинам искусственный интеллект может помочь контролировать статистику пациента эффективным, менее склонным к ошибкам и более единообразным способом, как описано в следующих разделах.

Ношение полезных мониторов

Все мониторы относятся к категории полезных. Фактически большинство из них не имеют никакого отношения к медицине, но все же дают положительный эффект для здоровья. Рассмотрим монитор Moov (<https://welcome.moov.cc/>), который контролирует и пульс, и движение в пространстве. Искусственный интеллект этого устройства отслеживает статистические данные и предоставляет совет о том, как лучше делать разминку. Фактически вы получаете советы о том, как лучше ставить ногу во время бега и стоит ли увеличивать шаг. Задача этого устройства — гарантировать самую эффективную разминку, которая улучшит здоровье, без риска получить травму.

Представьте, если бы такой монитор был слишком большим! Motiv (<https://mymotiv.com/>) представляет собой простое кольцо, которое контролирует то же самое, что и Моов, но меньше в объеме. Это кольцо даже следит, как вы спите, чтобы вы лучше отдыхали за ночь. Кольца действительно имеют свои преимущества и недостатки. Из статьи <https://www.wareable.com/smart-jewellery/best-smart-rings-1340> можно узнать об их проблемах больше. Достаточно интересно то, что большинство изображений на сайте ничем не напоминает фитнес-монитор; таким образом, вы получаете и стиль, и здоровье в одном пакете.

Конечно, если ваша единственная задача — контролировать свой пульс, можно обзавестись таким устройством, как Apple Watch (<https://support.apple.com/en-us/HT204666>), обеспечивающим также некий анализ с использованием искусственного интеллекта. Все эти устройства взаимодействуют со смартфоном, поэтому вы можете увязать данные с некими другими приложениями или при необходимости отослать их своему врачу.

Доверие к критически важным мониторам

Проблема с некоторыми показателями человеческого организма состоит в том, что они постоянно изменяются, поэтому их проверка лишь время от времени фактически ничего не дает. Уровень глюкозы, статистика о которой собирается при диабете, является одним из показателей, относящихся к этой категории. Чем чаще вы контролируете суточные повышения и падения уровня глюкозы, тем проще назначить лекарства и держать диабет под контролем. Такие устройства, как K'Watch (<http://www.pkvitality.com/ktrack-glucose/>), обеспечивают постоянный мониторинг, а также обладают приложением, которое человек может использовать для получения полезной информации о своем состоянии при диабете. Конечно, люди использовали регулярный мониторинг годами; просто это устройство обеспечивает тот дополнительный уровень контроля, который при наличии диабета может означать различие между незначительной неприятностью и опасной для жизни проблемой.

Постоянный мониторинг чьего-то сахара в крови или другой хронической болезни мог бы показаться чрезмерным, но у него есть практическое применение. Такие продукты, как Sentrian (<http://sentrian.com/>), позволяют использовать дистанционно собранные данные для прогноза заболевания прежде, чем оно фактически возникнет. Внося изменения в лекарства и поведение пациента прежде, чем произойдет событие, Sentrian снижает количество неизбежных госпитализаций, делая жизнь пациентов лучше и сокращая медицинские затраты.

Некоторые устройства, такие как нательный дефибриллятор Wearable Defibrillator Vest (WDV), действительно критически важны. Они непрерывно контролируют состояние сердца и при необходимости обеспечивают разряд, чтобы сердце не остановилось и работало правильно (см. <https://www.healthcentral.com/article/wearable-defibrillator-vest-pros-and-cons>). Это краткосрочное решение помогает врачу принять решение, нужна ли вживляемая версия того же устройства. У его ношения есть преимущества и недостатки, но, с другой стороны, трудно оценить постоянное наличие на теле кардиостимулятора, когда необходимо спасти свою жизнь. Конечно, основное назначение этого устройства — это мониторинг, и оно его обеспечивает. Некоторым людям на самом деле вживляемое устройство не нужно, поэтому контроль очень важен, чтобы избежать ненужной хирургической операции.

Использование носимых мониторов

Количество и разнообразие устройств для мониторинга состояния здоровья с поддержкой искусственного интеллекта на рынке сейчас колеблются (см. <https://www.mobihealthnews.com/content/31-new-digital-health->



ВНИМАНИЕ

Проблема любой медицинской технологии кроется в недостатке защиты. Вероятность взлома вживленного устройства ужасает. В статье <http://www.independent.co.uk/lifestyle/gadgets-and-tech/news/first-online-murder-will-happen-by-end-of-year-warns-us-firm-9774955.html> описано, что может случиться, если некто взломает медицинское устройство. К счастью, согласно многим источникам, никто еще не умер.

Однако представьте отказ инсулиновой помпы или вживленного дефибриллятора из-за взлома, а также вообразите вред, который может нести злоумышленник. Управление по санитарному надзору за качеством пищевых продуктов и медикаментов (Food and Drug Administration — FDA), наконец, опубликовало руководство по защите медицинских устройств, доступное по адресу <http://www.securityweek.com/fda-releases-guidance-medical-device-cybersecurity>, но в жизнь эти правила, очевидно, не претворяются. Фактически эта статья просто свидетельствует, что производители сейчас активно ищут способы защиты своих устройств.

Искусственный интеллект не несет ответственности за отсутствие защиты этих устройств, но может доказать вину злоумышленника. Дело в том, что необходимо учитывать все аспекты использования искусственного интеллекта, особенно когда дело доходит до устройств, непосредственно затрагивающих жизнь людей, таких как вживляемые медицинские устройства.

tools-showcased-ces-2017). Например, вы фактически можете купить зубную щетку с поддержкой искусственным интеллектом, которая будет контролировать ваши привычки по чистке зубов и давать советы относительно лучшей методики чистки (<http://www.businesswire.com/news/home/20170103005546/en/Kolibree-Introduces-Ara-Toothbrush-Artificial-Intelligence/>). Трудно представить, сколько сложностей пришлось преодолеть, создавая такое устройство, включая разработку схемы мониторинга состояния человеческого рта. Конечно, некоторые люди могут сказать, что улучшение чистки зубов в действительности не имеет непосредственного отношения к хорошему здоровью, но на самом деле это так (см. <https://www.mayoclinic.org/healthy-lifestyle/adult-health/indepth/dental/art-20047475>).

Как правило, при создании носимых мониторов ставится задача сделать их поменьше и менее навязчивыми. Простота — это также немаловажное требование для устройств, предназначенных для использования людьми с минимальными или вообще отсутствующими познаниями в медицине. Одним из таких

устройство является носимый кардиомонитор (ECG). Кардиомонитор в кабинете врача требует подключения пациента проводами к стационарному устройству, осуществляющему необходимые замеры. Устройство QardioCore (<https://www.prnewswire.com/news-releases/qardio-makes-a-breakthrough-in-preventative-healthcare-with-the-launch-of-qardiocore-the-first-wearable-ecg-monitor-300384471.html>) снимает кардиограмму без проводов, и любой человек с минимальными медицинскими знаниями легко может его использовать. Подобно многим другим устройствам, это полагается на ваш смартфон, чтобы сделать необходимый анализ и при необходимости подключиться к внешним источникам.



ЗАПОМНИ

Современные медицинские устройства работают хорошо, но они не портативны. Задача создания приложений с поддержкой искусственным интеллектом и специализированных устройств заключается в получении только необходимых данных, в которых фактически нуждается врач. В противном случае результатов анализов придется ждать. Даже если вы не купили умную зубную щетку или носимый кардиомонитор для контроля своего сердцебиения, тот факт, что эти устройства есть, что они малы, функциональны и удобны, уже хорош, поскольку при необходимости вы можете ими воспользоваться.

Расширение возможностей людей

Сегодня множество технологий нацелено на увеличение продолжительности здоровой человеческой жизни (сегмент жизни без существенной болезни), а не только на увеличение количества лет жизни. Они призваны дать людям больше возможностей для улучшения их здоровья различными способами. Вы легко можете найти статьи, предлагающие 30, 40 или даже 50 способов увеличения продолжительности здоровой человеческой жизни, но зачастую они сводятся к комбинации правильного питания, достаточных тренировок и хорошего отдыха. Конечно, установить, какие именно еда, упражнения и отдых лучше всего подходят именно для вас, почти невозможно. В следующих разделах описаны случаи, когда наличие устройства с поддержкой искусственного интеллекта могло бы означать различие между 60-ю годами хорошей жизни и 80-ю или даже больше. (Фактически сейчас нетрудно найти статьи о продолжительности в будущем человеческой жизни до тысячи и более лет благодаря новейшим технологиям.)

Проблема любой медицинской технологии кроется в недостатке защиты. Вероятность взлома вживленного устройства ужасает. В статье <http://www.independent.co.uk> Забота является главным назначением любого медицинского средства. Предположим, что забота осуществляется не только наилучшим образом, но и беспристрастно. Вот в чем искусственный интеллект может проявить себя в медицинской области; он гарантирует, что технические навыки останутся высокими и не будет никакого пристрастия, по крайней мере с его точки зрения.

Люди всегда будут пристрастны, поскольку они обладают внутриличностным интеллектом (как описано в главе 1). Даже самые добрые, альтруистически настроенные люди демонстрируют некоторую форму пристрастия (обычно подсознательно), создавая ситуацию, когда врач видит одно, а пациент — другое (см. раздел “Пять недостоверностей данных” в главе 2). Однако пациенты почти всегда склонны преувеличивать серьезность своей болезни, хотя, вероятно, и непреднамеренно. Использование искусственного интеллекта гарантирует беспристрастность при работе с пациентами, позволяет избежать этой проблемы. Искусственный интеллект может также помочь врачу обнаружить недостоверности (непреднамеренные или иные) в симптомах, излагаемых пациентами, повысив, таким образом, качество лечения.

В медицине время от времени возникают проблемы, поскольку чисто технических навыков зачастую недостаточно. Люди нередко жалуются на отсутствие такта у врачей и медицинского персонала. Те же самые люди, которые хотят беспристрастного лечения, так или иначе хотят и сочувствия со стороны врача (просто делающего свою работу), поскольку пристрастны теперь они. Сочувствие отличается от симпатии в корне. Люди демонстрируют *сочувствие*, когда они в состоянии чувствовать то же самое, что и пациент. Два упражнения из раздела “Программно-ориентированные решения” этой главы помогут вам понять, как некто может построить систему взглядов так, чтобы вызвать сочувствие. Искусственный интеллект никогда не сможет сочувствовать, поскольку у него нет ни чувств, ни понимания, необходимых для создания системы взглядов, а также внутриличностного интеллекта, обязательного для использования такой системы взглядов.

К сожалению, сочувствие может сыграть злую шутку с медиком относительно истинных медицинских потребностей, поскольку теперь он участвует в недостоверности точки зрения, наблюдая происходящее только с точки зрения пациента. Таким образом, врачи нередко испытывают *симпатию*, благодаря которой могут взглянуть на ситуацию извне, понять, что пациент мог бы чувствовать (а не что он чувствует), и не строить систему взглядов на основе этого.

Следовательно, врач может оказать необходимую эмоциональную поддержку, видя при этом реальную потребность в выполнении действий, которые пациенту могут вовсе и не понравиться в ближайшей перспективе. Искусственный интеллект не может решить эту задачу, поскольку ему недостает внутриличностного интеллекта и он не понимает концепции точки зрения достаточно хорошо, чтобы соответственно ее применить. k/lifestyle/gadgets-and-tech/news/first-online-murder-will-happen-by-end-of-year-warns-usfirm-9774955.html описано, что может случиться, если некто взламывает медицинское устройство. К счастью, согласно многим источникам, никто еще не умер.

Однако представьте отказ инсулиновой помпы или вживленного дефибриллятора из-за взлома, а также вообразите вред, который может нести злоумышленник. Управление по санитарному надзору за качеством пищевых продуктов и медикаментов (Food and Drug Administration — FDA), наконец, опубликовало руководство по защите медицинских устройств, доступное по адресу <http://www.securityweek.com/fda-releases-guidance-medical-device-cybersecurity>, но в жизнь эти правила, очевидно, не претворяются. Фактически эта статья просто свидетельствует, что производители сейчас активно ищут способы защиты своих устройств.

Искусственный интеллект не несет ответственности за отсутствие защиты этих устройств, но может доказать вину злоумышленника. Дело в том, что необходимо учитывать все аспекты использования искусственного интеллекта, особенно когда дело доходит до устройств, непосредственно затрагивающих жизнь людей, таких как вживляемые медицинские устройства.

Использование игр для терапии

Игровая консоль может стать мощным и интересным инструментом физической терапии. И Nintendo Wii, и Xbox 360 уже используются во многих местах для физической терапии (<https://www.webpt.com/blog/post/do-you-wii-hab-using-motion-gaming-your-therapy-clinic>). Задача этих игр в том, чтобы приучить людей двигаться определенным способом. Игра автоматически поощряет надлежащие движения пациента, получающего одновременно и терапию, и развлечение. Поскольку терапия становится развлечением, пациент выполняет ее с большей заинтересованностью и поправляется быстрее.

Конечно, одних только движений при работе с соответствующей игрой для успеха недостаточно. Фактически, играя в эти игры, кто-то сможет заработать новые болезни. Дополнение Jintronix для аппаратных средств Xbox Kinect стандартизирует использование данной игровой консоли для терапии (<https://www.forbes.com/sites/kevinanderton/2017/09/30/>

jintronix-program-uses-xbox-kinect-hardware-to-help-rehab-patients-infographic/#2cffc63a61d3), увеличивая вероятность положительного результата.

Использование экзоскелетов

Одной из самых сложных задач искусственного интеллекта является поддержка всего человеческого тела. Это случай, когда некто носит *экзоскелет* (по существу, разновидность робота). Искусственный интеллект считывает движения человека (или его намерение переместиться) и приводит экзоскелет в действие. В использовании экзоскелетов продвинулись вооруженные силы (см. <http://exoskeletonreport.com/2016/07/military-exoskeletons/>). Представьте себе возможность двигаться быстрее и носить значительно более тяжелый вес. Видео на <https://www.youtube.com/watch?v=p2W23ysgWKI> дает лишь некоторое представление об этих возможностях. Конечно, военные продолжают экспериментировать, но это пригодно и в гражданских целях. Экзоскелет, который вы в конечном счете увидите (а вы его когда-нибудь наверняка увидите), вероятно, будет создан военными.

Промышленность также интересуется технологиями экзоскелета (см. <https://www.nbcnews.com/mach/science/new-exoskeleton-does-heavy-liftingfactory-workers-ncna819291>). В настоящее время промышленные рабочие страдают от травм и массы болезней, вызванных регулярным стрессом. Кроме того, работа на заводе невероятно утомительна. Ношение экзоскелета снизит не только усталость, но и вероятность ошибки, сделав труд рабочих более эффективным. Люди, сохраняющие силы в течение всего дня, могут сделать больше и с гораздо меньшей вероятностью травмы или еще какого-либо вреда.

Экзоскелеты, используемые в промышленности сегодня, отражают свои военные начала. Их возможности и внешний вид в будущем изменятся так, чтобы больше походить на экзоскелеты из таких фильмов, как *Чужие* (Aliens) (<https://www.amazon.com/exec/obidos/ASIN/B01I0K018W/datacservip0f-20/>). Реальные примеры этой технологии (см. видео и статью <http://www.bbc.com/news/technology-26418358> в качестве примера) менее внушительны, но все равно полезны.

Самое интересное, что экзоскелеты позволяют людям делать не только то, на что они способны и сами, но и на что они не были раньше способны, и это замечательно. Например, недавно была опубликована статья об использовании экзоскелета, позволившего ходить ребенку с церебральным параличом (<https://www.smithsonianmag.com/innovation/this-robotic-exoskeleton-helps-kids-cerebral-palsy-walk-upright-180964750/>). Не все используемые в медицинских целях экзоскелеты предназначены для поддержания жизни, как бы то ни было. Например, экзоскелет может помочь пациенту после инсульта снова нормально ходить (<http://www.sciencemag.org/news/2017/07/>

watch-robotic-exoskeleton-help-stroke-patient-walk). По мере восстановления экзоскелет поддерживает человека все меньше и меньше и в конце концов становится ненужным. Некоторые пользователи даже объединили свой экзоскелет с другими устройствами, такими как Alexa от Amazon (см. <https://spectrum.ieee.org/the-human-os/biomedical/bionics/how-a-paraplegic-user-commands-this-exoskeleton-alexa-im-ready-to-walk>).



ЗАПОМНИ

Назначение экзоскелета вовсе не в том, чтобы сделать из вас Железного человека (Iron Man). Они, скорее, позволяют снизить регулярные нагрузки, ведущие к туннельному синдрому, и нагрузки после травм, а также помогают людям лучше решать те задачи, которые в настоящее время для них слишком утомительны или просто невозможны из-за физической ограниченности тела. С медицинской точки зрения использование экзоскелета — это хорошо, поскольку он сохраняет людям подвижность, а подвижность немаловажна для хорошего здоровья.

ТЕМНАЯ СТОРОНА ЭКСОСКЕЛЕТОВ

В Сети нет описания злонамеренных способов использования экзоскелетов, если не учитывать военные цели. Но ломать — не строить. Рано или поздно кто-нибудь придумает способы использования экзоскелетов во зло (как, вероятно, и любой технологии, описанной в этой главе). Предположим, например, что высокотехнологичные воры, используя экзоскелеты, украдут тяжелый объект.

Хотя эта книга призвана развеять мифы, окружающие искусственный интеллект, и ознакомить с положительными направлениями его применения, факт остается фактом, и действительно думающие люди не могут по крайней мере не задумываться о темной стороне любой технологии. Эта стратегия становится опасной, когда люди безосновательно поднимают тревогу. Да, взбесившиеся воры в экзоскелетах способны наделать дел, но просто нужно правильно обеспечить защиту от них, и, кроме того, этого еще не случалось. Этические соображения возможного применения, как позитивного, так и негативного, всегда сопровождают создание любой технологии, включая искусственный интеллект.

Повсюду в книге вы найдете различные этические и моральные соображения о позитивном применении искусственного интеллекта для помощи обществу. Безусловно, безопасность технологии важна, но имейте также в виду, что отказываться от технологии из-за ее негативного потенциала действительно неразумно.

Решение специфических задач

Когда-то потеря конечности, например, означала годы визитов к врачу, а также более короткую и менее счастливую жизнь. Однако хорошие эндопротезы и другие подобные устройства, большинство с поддержкой искусственным интеллектом, сделали этот сценарий историей для многих людей. Посмотрите, например, на танцующую пару <https://www.youtube.com/watch?v=AJOQj4NGJXA>. У женщины протез ноги. Сейчас люди могут бежать марафон или заниматься скалолазанием, даже потеряв свои настоящие ноги.

Многие люди считают термин *специальные потребности* (special needs) эквивалентным физической или умственной недостаточности или даже неполноценности. Однако почти у всех есть некая специальная потребность. В конце долгого дня кто-то с вполне нормальным зрением увеличивает размер шрифта в текстовом программном обеспечении или делает изображения крупнее. Цветовое программное обеспечение может помочь кому-то с нормальным цветным восприятием лучше различать детали, которые обычно ему не видны (по крайней мере тому, кто считает, что зрение у него нормальное). Когда люди становятся старше, они обычно нуждаются в большей помощи, чтобы услышать, увидеть, ощутить, коснуться или иначе взаимодействовать с самыми обыкновенными объектами. Аналогично помощь с такими задачами, как ходьба, может спасти кого-то от медицинского стационара и вернуть домой навсегда. Дело в том, что использование различных технологий с поддержкой искусственным интеллектом может существенно улучшить жизнь всем, о ком идет речь в следующих разделах.

Программно-ориентированные решения

Для удовлетворения специфических потребностей многие пользователи компьютеров сегодня полагаются на некий тип программно-ориентированного решения. Одним из самых известных подобных решений является программа для чтения с экрана Job Access With Speech (JAWS) (<https://www.freedomscientific.com/Products/Blindness/JAWS>), произносящая текст на дисплее, используя сложные методы. Как вы, возможно, догадались, любая полагающаяся на науку о данных и искусственный интеллект технология должна получить данные, интерпретировать их, а затем предоставить результат, как это делает программное обеспечение JAWS, являющееся хорошим примером для понимания возможностей и пределов программно-ориентированных решений. Наилучший способ убедиться в этом самом — это загрузить и установить программное обеспечение, а затем испытать его, как будто вы слепы и пытаетесь выполнить определенные действия в своей системе. (Не делайте ничего критически важного, ибо ошибки неизбежны.)

Люди, которых вы видите в Сети, особенно с большим опытом по части удивительной жизни несмотря на их специальные потребности, обычно являются специфическими людьми. Они действительно тяжело работали, чтобы добиться того, что имеют теперь. Использование устройства с поддержкой искусственным интеллектом может быть шагом в нужном направлении, но чтобы достичь цели, нужна воля сделать все, независимо от того, что именно делает это устройство и сколько часов терапии потребуется. Данная глава не пытается пролить свет на невероятный объем работы, который проделали эти удивительные люди, чтобы жить полноценно. Она скорее о технологии, оказавшей им помощь на пути к успеху. Если вы хотите увидеть нечто действительно удивительное, посмотрите на балерину <http://www.dailymail.co.uk/news/article-3653215/Schoolgirlleg-amputated-knee-foot-attached-stump-suffering-rare-bone-cancer-defies-oddscompetitive-ballet-dancer.html>. В статье описано, сколько работы потребовалось, чтобы заставить эту технологию работать.



СОВЕТ

Подобное программное обеспечение позволяет людям со специальными потребностями делать невероятные вещи. Оно способно даже помочь понять то, на что это походило бы, имей они специальную потребность. Доступно немало таких приложений, но попробуйте для примера Vischeck с <http://www.vischeck.com/vischeck/vischeckImage.php>.¹ Оно позволяет увидеть изображения так, как их видят люди с различными формами дальтонизма. Конечно, вы сразу заметите, что термин *дальтонизм* фактически неправилен; люди в этом состоянии различают цвета, но не так. Цвета просто смещаются к другому цвету, поэтому лучшим термином, вероятно, было бы *смещение цвета*.

Доверие к аппаратному усилению

Для помощи людям со специальными потребностями необходимо куда больше, чем только подходящее программное обеспечение. В разделе “Использование экзоскелетов”, приведенном выше в этой главе, обсуждались способы использования экзоскелетов для предотвращения травм, усиления естественных человеческих возможностей или удовлетворения специальных потребностей (чтобы люди, страдающие параличом нижних конечностей, могли ходить). Однако потребности людей решают и многие другие виды аппаратных

¹ Актуальность ссылки не гарантируется. — *Примеч. ред.*

средств усиления, подавляющему большинству которых для правильной работы требуется помощь искусственного интеллекта некоторого уровня.

Рассмотрите, например, систему Eye-Gaze (<http://www.eyegaze.com/>). Прежние системы полагались на шаблон, прикрепленный поверх монитора. Парализованный мог смотреть на отдельные символы, которые считывались двумя камерами (по одной на каждой стороне монитора), а затем вводились в компьютер. Вводя команды таким способом, больной мог выполнить простые задачи на компьютере.

Некоторые из первых систем, управляемых взглядом, совмещались через компьютер с роботизированным манипулятором. Роботизированный манипулятор мог делать чрезвычайно простые, но важные действия, например подать пользователю напиток или вытереть ему нос. Современные системы фактически способны соединить мозг пользователя непосредственно с роботизированным манипулятором и позволить ему выполнять такие задачи, как еда без посторонней помощи (см. <https://www.engadget.com/2017/03/29/paralyzed-man-first-to-move-his-arm-by-thinking-about-it/>).

Искусственный интеллект эндопротеза

Можно найти множество примеров применения искусственного интеллекта в эндопротезах. Да, существуют и пассивные решения эндопротеза, но более новые полагаются на динамические подходы, требующие искусственного интеллекта. Один из наиболее удивительных примеров эндопротеза с поддержкой искусственным интеллектом — это полностью динамическая стопа, созданная Хью Херром (Hugh Herr) (<https://www.smithsonianmag.com/innovation/future-robotic-legs-180953040/>). Эта стопа и лодыжка работают настолько хорошо, что фактически позволили Хью заниматься скалолазанием. Вы можете посмотреть презентацию, недавно выпущенную TED, по адресу <https://www.youtube.com/watch?v=CDsNZJTWw0w>.



ВНИМАНИЕ!

Небольшая дилемма, с которой нам, вероятно, придется встретиться когда-нибудь в будущем (к счастью, не сегодня), когда эндопротезы фактически позволят их владельцам существенно превзойти природные человеческие возможности. Например, у Ситандры в фильме *Эон Флакс* кисти вместо стоп (<https://www.awn.com/vfx-world/aeon-flux-live-action-animated-world>). Кисти — это, по существу, своего рода протез, привитый тому, кто изначально имел обычные ступни. Вопрос в том, допустим ли подобный вид протезирования, полезен ли он и даже желателен ли. В некий момент должна будет собраться группа людей и установить, где должно заканчиваться использование протезов, чтобы сохранить людей в виде

людей (если мы только решим оставаться людьми, а не развиться в некую следующую форму жизни). Сегодня вы вряд ли кого увидите с кистями вместо стоп.

Новые способы анализа

Использовать искусственный интеллект в медицинских целях можно не только для разных носимых устройств. Его возможности способны повысить потенциал специалистов и другими полезными способами. Анализ данных — это одна из областей, в которых искусственный интеллект непревзойден. Фактически роли искусственного интеллекта в современной медицине посвящены целые веб-сайты, такие как <http://medicalfuturist.com/category/blog/digitalized-care/artificial-intelligence/>.²

Казалось бы, чтобы поставить точный диагноз, специалисту достаточно сделать снимок участка потенциальной опухоли, а затем просмотреть результат. Но большинство методик получения необходимого снимка полагается на проницаемость тканей, не являющихся частью опухоли и не затеняющих результат. Кроме того, врач желает получить наилучшую информацию, просматривая снимки опухоли на ее ранних этапах.

Искусственный интеллект позволяет не только поставить диагноз и идентифицировать опухоль на ранних этапах, но и сделать это с высокой точностью и очень быстро. При многих болезнях время критически важно. Согласно статье <https://www.wired.com/2017/01/look-x-rays-moles-living-ai-coming-job/> повышение скорости существенно, а цена не столь велика при использовании этого нового подхода.

Столь же внушительной, как и возможная скорость диагностики искусственного интеллекта в этой области, является возможность его объединения с Интернетом вещей (IoT) и компиляция данных. Когда искусственный интеллект обнаруживает у пациента специфическое состояние, он может самостоятельно проверить записи пациента и вывести на экран информацию, имеющую отношение к наблюдаемому диагнозу, как демонстрируется в статье <https://www.itnonline.com/article/how-artificial-intelligence-will-change-medical-imaging>. Теперь у врача есть каждая деталь необходимой информации о пациенте, чтобы поставить диагноз и назначить конкретное лечение.

² Вероятно, имелось в виду <https://medicalfuturist.com/category/artificial-intelligence/>. — *Примеч. ред.*

Новые хирургические технологии

Сегодня роботы и искусственный интеллект уже участвуют в хирургических процедурах. Фактически некоторые хирургические операции были бы практически невозможны без использования роботов и искусственного интеллекта. Однако история применения этой технологии еще не велика. Первый хирургический робот, Arthrobot, появился в 1983 году (см. <https://www.roboticoncology.com/history-of-robotic-surgery/>). Даже в то время применение этих спасительных технологий снижало количество ошибок, улучшало результаты, ускоряло заживание и делало хирургию относительно проще. В следующих разделах рассматривается применение роботов и искусственного интеллекта в различных аспектах хирургии.

Выработка хирургических рекомендаций

К хирургическим рекомендациям можно относиться по-разному. Например, искусственный интеллект может проанализировать все данные о пациенте и предоставить хирургу рекомендации о наилучших подходах на основании записей данного конкретного пациента. Хирург может выполнить эту задачу и сам, но это отнимет больше времени и возможны ошибки, которых искусственный интеллект не сделает. Искусственный интеллект не устанет к концу дня и ничего не упустит; он последовательно просмотрит все доступные данные точно так, как и всегда.

К сожалению, даже несмотря на помощь искусственного интеллекта, неожиданности во время хирургической операции все еще случаются, и тогда в игру вступает следующий уровень рекомендаций. Согласно статье https://www.huffingtonpost.com/entry/therole-of-ai-in-surgery_us_58d40b7fe4b002482d6e6f59 у врачей теперь может быть доступ к устройству, работающему на тех же принципах, что и Alexa, Siri или Cortana (искусственный интеллект устройств, которые фактически могут быть в любом доме). Нет, устройство не будет спрашивать врача о музыке, проигрываемой во время хирургической операции, но врач может использовать устройство для поиска определенной информации без необходимости отрываться от дела. Для пациента это означает пользу от второго мнения при непредвиденном осложнении во время хирургической операции. Устройство фактически не делает ничего большего, чем уже сделано другими врачами в ходе их исследований, оно просто с готовностью дает ответ на вопрос хирурга; никакого реального мышления здесь нет.

Подготовка к хирургической операции означает также просмотр всех анализов, на которых настаивает врач. Скорость — это преимущество искусственного интеллекта перед рентгенологом. Такие системы с технологией глубокого

обучения, как Enlitic (<https://www.enlitic.com/>), способны проанализировать рентгеновский снимок за миллисекунды, в 10 тысяч раз быстрее рентгенолога. Кроме того, система на 50 процентов лучше в классификации опухолей и имеет весьма низкий процент ошибок (0 процентов против 7 у людей). Другая система той же самой категории, Arterys (<https://arterys.com/>), способна выполнить сканирование сердца за 6–10 минут, а не за час, как обычно. Пациентам даже не приходится задерживать дыхание. Удивительно, но эта система получает несколько размерностей данных (трехмерную анатомию сердца, скорость и направление кровотока) за это короткое время. Видео об Arterys можно посмотреть по адресу <https://www.youtube.com/watch?v=IcooATgPYXc>.

РАБОТА В СТРАНАХ ТРЕТЬЕГО МИРА

Люди зачастую сетуют, что ни одна из этих удивительных технологий, на которые сегодня полагаются медицинские профессионалы, не применяются в странах третьего мира. Фактически некоторые из этих технологий, такие как продукты от Bay Labs (<https://baylabs.io/>), предназначены специально для стран третьего мира. Врачи опробовали их в Африке для выявления признаков ревмокардита (Rheumatic Heart Disease — RHD) у кенийских детей. Во время визита в сентябре 2016 года врачи использовали оборудование Bay Labs для осмотра 1 200 детей за четыре дня и выявили 48 детей с RHD или врожденными пороками сердца. Без искусственного интеллекта такое оборудование не было бы возможным, поскольку иначе оно никак не могло бы быть ни достаточно компактным, ни легким, чтобы работать в этих условиях.

Помощь хирургу

Самая роботизированная хирургическая система сегодня только помогает хирургу, а не заменяет его. Первая роботизированная хирургическая система, RUMA, появилась в 1986 году. Она выполняла чрезвычайно тонкую нейрохирургическую биопсию, являющуюся лапароскопическим типом хирургической операции. Лапароскопическая хирургия минимально агрессивна, с одним или несколькими небольшими разрезами, служащими только для доступа к такому органу, как желчный пузырь, для удаления или восстановления. Первые работы не были достаточно искусны в решении таких задач.

В 2000 году хирургическая система da Vinci продемонстрировала способность выполнять лапароскопическую хирургическую операцию, используя трехмерную оптическую систему. Хирург направляет движения робота, но фактическую хирургию выполняет робот. Во время операции хирург видит

высококачественное изображение и может следить за операцией куда лучше, чем выполняя ее лично. Система da Vinci использует также меньшие отверстия, чем хирург, и риск инфекции также ниже.

Тем не менее важнейшая особенность системы da Vinci в том, что она усиливает возможности хирурга. Например, если рука хирурга немного дрогнет, хирургическая система это устранил, подобно средствам стабилизации съёмочной камеры. Система также сглаживает внешнюю вибрацию. Она также позволяет хирургу выполнять чрезвычайно малые движения, куда меньшие, чем способен человек, обеспечивая намного более высокую точность хирургической операции, чем хирург может достичь сам.



Хирургическая система da Vinci является сложным и чрезвычайно гибким устройством. Министерство здравоохранения одобрило его и для педиатрических, и для взрослых хирургических операций следующих типов.

- » Урологическая хирургия
- » Общая лапароскопическая хирургия
- » Общая не сердечно-сосудистая торакоскопическая хирургия
- » Вспомогательные торакоскопические кардиотомические процедуры

Вся эта медицинская терминология приведена здесь только для доказательства того, что хирургическая система da Vinci способна выполнять множество задач, не задействуя хирурга непосредственно. В один прекрасный момент роботы-хирурги станут более автономными, отстранив врача еще дальше от пациента во время хирургической операции. В будущем никто вообще не войдет в операционную с пациентом, чтобы снизить вероятность инфекции почти до нуля. О хирургической системе da Vinci можно прочитать больше по адресу <http://www.davincisurgery.com/da-vincisurgery/da-vinci-surgical-system/>.

Замена хирургии наблюдением

В фильме *Звездные войны* (Star Wars) роботы-хирурги регулярно чинят людей. Фактически возникает вопрос, есть ли у них врачи-люди вообще. Теоретически роботы в будущем могли бы и сами проводить некоторые типы хирургических операций, но эта возможность все еще далека. Роботам еще предстоит проделать весьма долгий путь от промышленных установок, которыми они являются сегодня. Современные роботы почти не автономны и требуют человеческого вмешательства для настройки.

Однако искусство хирургии — это достижение для роботов. Например, система STAR (Smart Tissue Autonomous Robot) превзошла по скорости хирургов

при зашивании кишечника свиньи, как описано в статье <https://spectrum.ieee.org/the-human-os/robotics/medical-robots/autonomous-robot-surgeon-bests-human-surgeons-in-world-first>. Во время хирургической операции врачи контролировали STAR, но фактически робот выполнил задачу самостоятельно, что является огромным шагом вперед в автоматизации хирургических операций. Видео о ходе хирургической операции (см. https://www.youtube.com/watch?v=vb79-_hGLkc) весьма информативно.

Автоматизация решений

Искусственный интеллект великолепен при автоматизации. Он никогда не отклоняется от инструкций, никогда не устает и не делает ошибок, если первоначальные инструкции правильны. В отличие от людей искусственный интеллект никогда не нуждается в отпуске, выходных или даже восьмичасовом рабочем дне (не у многих в медицинской профессии есть все это). Следовательно, искусственный интеллект, общающийся с пациентами с утра, будет так же работать и в обед, и вечером. Изначально у искусственного интеллекта не особенно много существенных преимуществ, если смотреть исключительно с точки зрения согласованности, точности и долговечности (случаи, когда искусственный интеллект терпит неудачу, приведены в разделе “ПРИСТРАСТИЕ, СИМПАТИЯ И СОЧУВСТВИЕ”). Ниже рассматриваются способы, которыми искусственный интеллект может помочь в автоматизации за счет улучшения доступа к таким ресурсам, как данные.

Работа с медицинскими записями

Главная помощь искусственного интеллекта в медицине заключается в работе с медицинскими записями. В прошлом для хранения всех терапевтических данных использовали записи на бумаге. У каждого пациента могла быть также специальная доска, на которой медицинский персонал ежедневно записывал текущую информацию во время пребывания пациента в стационаре. Данные пациента содержат и различные диаграммы, в которые врач может также вносить примечания. Наличие всех этих источников информации в очень разных местах существенно затрудняет контроль над состоянием пациента. Искусственный интеллект совместно с компьютерной базой данных помогает сделать эту информацию более доступной, единообразной и надежной. Такие системы, как Google Deepmind Health (<https://deepmind.com/applied/deepmind-health/working-partners/health-research-tomorrow/>), позволяют медперсоналу обрабатывать информацию пациентов так, чтобы находить в ней шаблоны, которые иначе не были бы очевидны.



ЗАПОМНИ

Врачи не всегда работают с записями таким же образом, как и все остальные. Система, подобная WatsonPaths от IBM (<http://www.research.ibm.com/cognitivecomputing/watson/watsonpaths.shtml>), помогает врачам работать с данными пациента всех видов новыми способами, чтобы принимать наилучшие диагностические решения. Видео о работе системы можно посмотреть по адресу <https://www.youtube.com/watch?v=07XPEqkHJ6U>.

Медицина — дело командное, в ней много разных людей сотрудничают друг с другом. Однако любой, кто понаблюдает процесс некоторое время, очень скоро поймет, что эти люди не общаются между собой в достаточной степени, поскольку они слишком заняты пациентами. Такие системы, как CloudMedX (<http://www.cloudmedxhealth.com/>), собирают все, что введено всеми участниками, и анализируют. Это программное обеспечение способно помочь обнаружить потенциально проблемные области, повысив вероятность хорошего терапевтического результата. Другими словами, эта система частично берет на себя разговоры, которые, вероятно, вели бы между собой различные заинтересованные лица, если бы они не были столь заняты лечением.

Предсказание будущего

Действительно удивляет программное обеспечение CareSkore (<https://www.careskore.com/>), создающее прогноз на основании медицинских записей. Оно использует алгоритмы для определения вероятности повторного попадания пациента в больницу после выписки. Выяснив причины вероятной повторной госпитализации, медицинский персонал может устранить их прежде, чем пациент покинет больницу, сделав повторную госпитализацию менее вероятной. Согласно этой стратегии Zephyr Health (<https://zephyrhealth.com/>) помогает врачам подбирать подходящие методы лечения и достигать положительного результата с наибольшей вероятностью, что также снижает для пациента риск повторной госпитализации. Видео на <https://www.youtube.com/watch?v=9y930hioWjw> поможет узнать больше о Zephyr Health.

В некотором отношении именно ваша генетика определяет то, что будет с вами в будущем. Следовательно, знание генетики увеличивает шансы понять свои преимущества и недостатки, что поможет вам жить лучше. Система Deep Genomics (<https://www.deepgenomics.com/>) выясняет, как генетические мутации влияют на человека. Мутации не всегда приводят к отрицательному результату; некоторые мутации делают людей лучше, поэтому знание о мутациях может быть положительной практикой. Более подробная информация по этой теме содержится в видео на <https://www.youtube.com/watch?v=hVibPJyf-xg>.

Повышение безопасности процедур

Чтобы принимать правильные решения, врачи нуждаются в большом количестве данных. Однако при повсеместном распространении данных врачи зачастую не могут проанализировать эти разнородные данные достаточно быстро и принимают несовершенные решения. Чтобы сделать процедуры более безопасными, врач должен иметь доступ не только к данным, но также к неким средствам их организации и анализа способом, отражающим специальность врача. Одним из таких средств является Oncora Medical (<https://oncoramedical.com/>), оно собирает и организует медицинские записи онкологов-радиологов. В результате им удастся установить точный объем радиации только для правильного места, чтобы получить необходимый результат при минимальном нежелательном побочном эффекте.

Врачам также бывает трудно получить необходимую информацию потому, что используемые ими машины обычно велики и дороги. Изобретатель Джонатан Ротберг (Jonathan Rothberg) решил заменить все это системой Butterfly Network (<https://www.butterflynetwork.com/#News>). Вообразите устройство на базе iPhone, способное сделать МРТ и ультразвуковое сканирование. Изображение на веб-сайте просто удивительно.

Создание лучших медикаментов

Сегодня все жалуются на цену лекарств. Да, лекарства способны на удивительные вещи, но иногда стоят так дорого, что некоторые закладывают дома, чтобы их купить. Частично проблема кроется в продолжительности проверки. Анализ тканей для выявления действия нового препарата может занять до года. К счастью, такие системы, как 3Scan (<http://www.3scan.com/>), способны сократить время того же самого анализа тканей всего до одного дня.

Конечно, было бы просто прекрасно, если бы производящая препараты компания имела лучшее представление о вероятном действии лекарства до того, как инвестирует деньги в его исследование. Система Atomwise (<http://www.atomwise.com/>) использует огромную базу данных молекулярных структур, чтобы выполнять исследования, при которых молекулы будут отвечать поставленным требованиям. В 2015 году исследователи использовали систему Atomwise для разработки лекарства, сделавшего эболу не более чем очередной инфекционной болезнью. Анализ, который занял бы у людей месяцы или даже годы исследований, система Atomwise провела всего за один день. Вообразите этот сценарий во время потенциально глобальной эпидемии. Если Atomwise способна выполнить анализ, необходимый для выявления вируса или неинфекционной бактерии за один день, возможная эпидемия может быть предотвращена.

Фармацевтические компании производят огромные количества лекарств. Причина столь внушительного разнообразия продукции, помимо доходов, в том, что каждый человек немного отличается от других. Препарат, хорошо помогающий и не дающий побочных эффектов одним людям, может вовсе не помогать и даже вредить другим. Система Turbine (<http://turbine.ai/>) позволяет фармацевтическим компаниям моделировать препараты так, чтобы они с высокой вероятностью подходили для конкретных людей. Сейчас Turbine специализируется на лечении рака, однако несложно предположить, что тот же самый подход применим во многих других областях.



СОВЕТ

Лекарства бывают разными. Некоторые полагают, что они бывают в виде пилюль и уколов; на самом деле диапазон значительно шире, вплоть до микробиомов. Фактически количество микробов в теле человека примерно равно одной десятой части количества его собственных клеток, и большинство этих микробов важны для его жизни; без них он умер бы очень быстро. Система Whole Biome (<https://www.wholebiome.com/>) использует множество методов, чтобы наладить работу микробиомов так, чтобы вы не обязательно нуждались в пилюле или уколе, чтобы нечто исправить. Посмотрите видео на <https://www.youtube.com/watch?v=t1Y2AckssyI>, чтобы узнать больше.

Некоторым компаниям все еще предстоит реализовать свой потенциал, и в конечном счете они, вероятно, так и сделают. Одна из таких компаний, Recursion Pharmaceuticals (<https://www.recursionpharma.com/>), использует автоматизацию для разработки новых способов применения известных лекарств, биологически активных и фармацевтических препаратов, которые ранее недооценили, для решения новых проблем. Компания имела некоторый успех в лечении редких генетических заболеваний и задалась целью победить сто болезней за следующие десять лет (безусловно, задача чрезвычайно сложная).

Объединение роботов и медицинских специалистов

В общество начинают интегрироваться полуавтономные роботы с ограниченными возможностями. В Японии они уже используются некоторое время (см. <https://www.japantimes.co.jp/news/2017/05/18/national/science-health/japans-nursing-facilities-using-humanoid-robots-improve-lives-safety-elderly/>). В Америке используются такие роботы, как Rudy (см. <https://www.roboticsbusinessreview.com/rbr/>

rudy_assistive_robot_helps_elderly_age_in_place/). В большинстве случаев эти роботы способны выполнять очень простые задачи, такие как напоминания о приеме лекарств и простые игры, безо всяких возможностей в смысле вмешательства. Однако при необходимости врач или другой медицинский специалист может дистанционно связаться с роботом и выполнить с его помощью более сложные задачи. Применение подобных средств позволяет оказать человеку мгновенную помощь, и это не так уж и дорого.



ЗАПОМНИ

Сейчас подобные роботы находятся еще в младенчестве, но со временем они усовершенствуются. Хотя эти роботы являются вспомогательным инструментом медицинского персонала и никак не могут заменить врача или медсестру при выполнении многих специальных задач, они действительно обеспечивают постоянное наблюдение за пациентом наряду с эффектом присутствия. Кроме того, роботы могут выполнять простейшие человеческие задачи (такие, как распределение пилюль, напоминания и помощь при ходьбе), причем весьма хорошо.



Глава 8

Искусственный интеллект в человеческом общении

В ЭТОЙ ГЛАВЕ...

- » Общение новыми способами
- » Обмен идеями
- » Использование средств массовой информации
- » Улучшение сенсорного восприятия людей

Люди общаются друг с другом бесчисленным количеством способов. Фактически многие даже не догадываются, сколько есть способов коммуникации. Когда люди обычно думают о коммуникации, они подразумевают литературу или разговор. Однако общение может иметь и множество других форм, включая визуальный контакт, интонацию голоса и даже запах (см. <https://www.smithsonianmag.com/sciencenature/the-truth-about-pheromones-100363955/>). Примером компьютерной версии улучшения человеческих возможностей является электронный нос, использующий для решения

своей задачи комбинацию электроники, биохимии и искусственного интеллекта, который нашел широкое применение в промышленности и научных исследованиях (см. <https://www.ncbi.nlm.nih.gov/pmc/articles/PMC3274163/>). Данная глава сосредоточена в основном на обычной коммуникации, однако затрагивает и жестикуляцию. Вы получите лучшее представление о том, как искусственный интеллект может улучшить коммуникацию между людьми, причем куда дешевле электронного носа.

Искусственный интеллект может также улучшить способ обмена идеями между людьми. В некоторых случаях он предоставляет совершенно новые методы коммуникации, как правило, едва уловимо (а иногда и вполне отчетливо) улучшая существующие способы обмена идеями. Обмениваясь идеями, люди создают новые технологии на основании уже существующих, а также узнают о технологиях, необходимых для повышения уровня их знаний. Идеи абстрактны, что зачастую особенно затрудняет обмен ими, поэтому искусственный интеллект способен перебросить необходимый мост между людьми.

Когда некто хочет поделиться своими знаниями, он обычно полагается на литературу. В некоторых случаях можно повысить свою коммуникабельность, используя графику. Но только некоторые люди могут использовать обе эти формы массовой информации, чтобы получить новые знания; большинству людей требуется куда больше; именно поэтому такие сетевые источники, как YouTube (<https://www.youtube.com/>), стали настолько популярными. Используя искусственный интеллект, можно существенно увеличить мощь уже существующих средств массовой информации, и эта глава рассказывает, как именно.

Заключительный раздел этой главы поможет понять, как искусственный интеллект способен обеспечить почти сверхчеловеческое сенсорное восприятие. В конце концов, вы, возможно, тоже хотите иметь электронный нос; он действительно дает существенные преимущества в обнаружении запахов, которые могут быть значительно менее ароматными, чем у людей. Предположим, удастся достичь того же уровня обоняния, что и у собаки (она использует сто миллионов рецепторов запаха против одного миллиона у человека). Достижение этой цели можно двумя путями: используя мониторы, общающиеся с человеком косвенно, и непосредственно воздействуя на сенсоры человека.

Новые способы коммуникации

Люди использовали разговорную речь для общения задолго до появления письменности. Единственная проблема разговорной речи в том, что для общения обе стороны должны находиться в одном и том же месте в одно и то же время. Следовательно, письменное общение лучше во многих отношениях, по-

скольку оно не зависит от времени и не требует от общающихся сторон встречи. Три основных метода невербальной коммуникации между людьми полагаются на следующее.

- » **Алфавит.** Абстракция компонентов человеческих слов или символов.
- » **Язык.** Свод слов или символов, позволяющий создавать предложения или передавать идеи в письменной форме.
- » **Жестикуляция.** Усиление языка контекстом.

Первые два метода — прямые абстракции разговорных слов. Их не всегда просто реализовать, но люди делали это на протяжении тысяч лет. Реализовать жестикуляцию труднее всего, поскольку придется создавать абстракцию физического процесса. Литература способна передавать жестикуляцию, используя специфическую терминологию, такую как описано в статье <https://writers-write.co.za/cheat-sheets-for-writing-body-language/>. Но когда письменных слов недостаточно, люди усиливают их такими символами, как эмодзи (смайлики) и эмодзи (см. <https://www.britannica.com/story/whats-the-difference-between-emoji-and-emoticons>). Ниже эти проблемы обсуждаются более подробно.

Создание новых алфавитов

Вначале обсудим два новых алфавита, появившихся в век компьютеров: эмодзи (<http://cool-smileys.com/text-emoticons>) и эмодзи (<https://emojipedia.org/>). На сайтах с изображениями этих двух графических алфавитов их содержатся сотни. По большей части люди легко интерпретируют эти графические алфавиты, поскольку они напоминают выражения лица, но у компьютерных приложений нет человеческого восприятия эмоций, поэтому искусственный интеллект им зачастую требуется только для того, чтобы выяснить, какую эмоцию человек пытается передать этим небольшим изображением. К счастью, для них есть стандартизированные списки, такие как таблица эмодзи Unicode по адресу <https://unicode.org/emoji/charts/full-emoji-list.html>. Конечно, стандартный список не поможет с переводом. Статья <https://www.geek.com/tech/ai-trained-on-emoji-can-detect-socialmedia-sarcasm-1711313/> предоставляет больше подробностей о том, как можно обучать искусственный интеллект интерпретировать эмодзи и реагировать на них (а также на дополнительные эмодзи). С этим процессом в действии можно ознакомиться по адресу <https://deepemoji.mit.edu/>.

Эмодзи — это устаревшая технология, и многие люди пытаются забыть их из всех сил (но, вероятно, безуспешно). Эмодзи является новой и весьма

захватывающей технологией, удостоившейся мультфильма (см. <https://www.amazon.com/exec/obidos/ASIN/B0746ZZR71/datacservip0f-20/>). Вы можете положиться на искусственный интеллект Google даже для превращения в эмодзи собственного селфи (см. <https://www.fastcodesign.com/90124964/exclusive-new-google-tooluses-ai-to-create-custom-emoji-of-you-from-a-selfie>). На случай, если вы не хотите просеивать все 2 666 официальных эмодзи с поддержкой Unicode (или 564 септильонов эмодзи, создаваемых Allo, от Google (<https://allo.google.com/>)), то можете положиться на систему Dango (<https://play.google.com/store/apps/details?id=co.dango.emoji.gif&hl=en>), способную предложить вам соответствующую эмодзи (см. <https://www.technologyreview.com/s/601758/this-app-knows-just-the-right-emoji-for-any-occasion/>).



ЗАПОМНИ

Люди создавали новые алфавиты для решения конкретных задач испокон времен. Эмотиконы и эмодзи представляют собой лишь два из многих алфавитов, созданных людьми для Интернета и искусственного интеллекта. Фактически искусственный интеллект может потребоваться только для того, чтобы разобраться в них всех.

Автоматизация перевода

У мира всегда была проблема с отсутствием общего языка. Да, английский язык стал более или менее универсальным, но он все еще не полностью универсален. Перевод на другие языки дорог, громоздок и склонен к ошибкам, поэтому переводчики необходимы во многих ситуациях, что тоже не всегда является наилучшим решением. Тем, кто испытывает затруднения при общении на других языках, могут помочь такие приложения, как Google Translate (рис. 8.1).

Обратите внимание: программа Google Translate на рис. 8.1 предлагает автоматически определить язык. И что интересно, в большинстве случаев это средство работает чрезвычайно хорошо, частично благодаря системе GNMT (Google Neural Machine Translation — нейронный машинный перевод Google). Фактически она просматривает все предложение, чтобы перевести его осмысленно, что куда лучше перевода на основании отдельных фраз или слов (подробности — на <http://www.wired.co.uk/article/google-ai-language-create>).



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Но что еще удивительнее, система GNMT способна осуществлять перевод даже тогда, когда у нее нет подходящего словаря. В этом случае используется искусственный язык *интерлингва* (см. <https://en.oxforddictionaries.com/definition/interlingua>). Интерлингва не является универсальным языком, это скорее универсальный мост между языками. Скажем, у системы GNMT нет китайско-испанского

словаря, но есть китайско-английский и англо-испанский. Построив трехмерную сеть, представляющую эти три языка (интерлингва), система GNMT способна осуществлять китайско-испанский перевод. К сожалению, она не сможет переводить с китайского на марсианский, поскольку нет марсианского словаря ни для одного из естественных языков. Однако, чтобы система GNMT работала, ее словарную базу должны были создать люди.

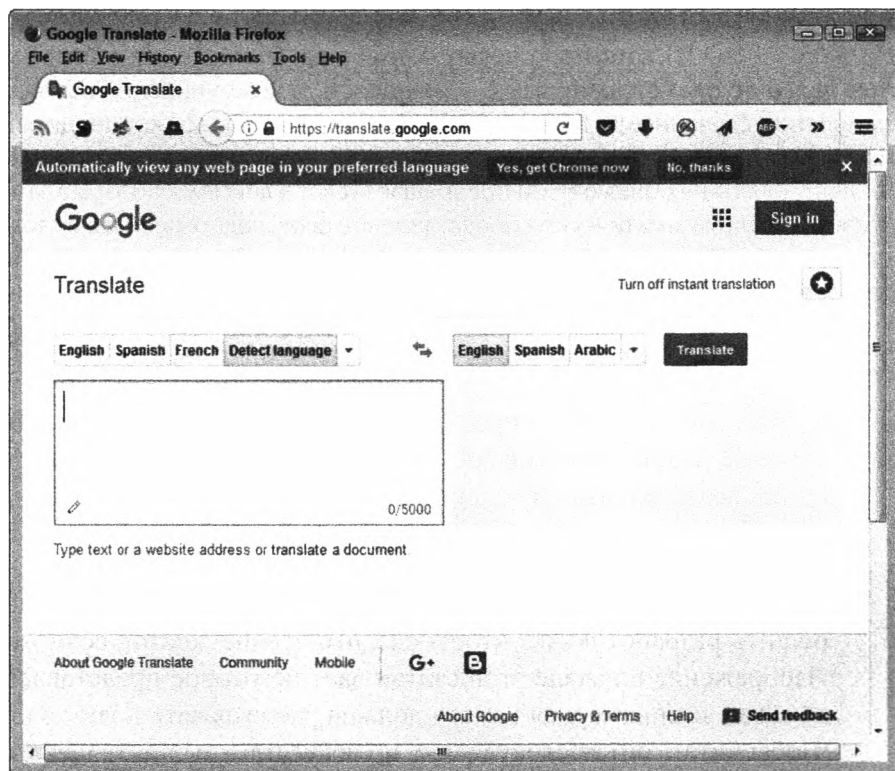


Рис. 8.1. Переводчик Google как пример искусственного интеллекта, выполняющего важную повседневную задачу

Включение жестикуляции

Существенной частью общения между людьми являются мимика и жесты, а потому использование эмодзи важно. Однако люди все чаще используют камеры для видеосвязи и других форм связи, не подразумевающих письменности. В данном случае компьютер может слышать вводимые человеком данные, разделять их на лексемы, представляющие человеческую речь, а затем обрабатывать эти лексемы так, чтобы выполнить запрос подобно таким системам, как Alexa или Google Home.

После чтения таких статей, как <https://www.fastcodesign.com/90132632/ai-is-inventing-its-own-perfect-languageshould-we-let-it>, может сложиться впечатление, что некоторые типы искусственного интеллекта так или иначе могут создавать новые непонятные людям языки, а затем использовать их для общения. Концепция удивительно похожа на сценарий фильма *Терминатор* (<https://www.amazon.com/exec/obidos/ASIN/B00938UVC2/dataservip0f-20/>), хотя это, конечно, фантастика. Тем не менее, если читать статью далее, то окажется, что язык — не новость, а коммуникации выглядят на удивление случайными. В предыдущей главе упоминалась концепция понимания со стороны искусственного интеллекта. Фактически искусственный интеллект ничего не понимает; он превращает текст в лексемы, которые затем анализирует как математическое представление слов, содержащихся в таблице подстановок. В разделе “Аргумент китайской комнаты” главы 5 обсуждается использование таблиц подстановок для уверения наблюдателя в том, что кто-то понимает некий язык, когда в действительности это не так.



ЗАПОМНИ

К сожалению, простое преобразование произнесенных слов в лексемы не решит всей задачи, поскольку остается еще невербальное общение. В данном случае искусственный интеллект должен быть способен читать жесты непосредственно. В статье <https://www.cmu.edu/news/stories/archives/2017/july/computer-reads-body-language.html> обсуждаются некоторые из проблем, которые предстоит решить разработчикам, чтобы сделать чтение жестов возможным. Изображение в начале этой статьи дает некоторое представление о том, как компьютерная камера должна распознавать позы человека, чтобы читать его жестикуляцию, но искусственный интеллект зачастую требует, чтобы ввод осуществлялся сразу с нескольких камер, чтобы восполнить пробелы, когда одна часть тела человека закрыта другой с точки зрения первой камеры. Чтение жестикуляции подразумевает интерпретацию таких человеческих характеристик.

- » Поза
- » Движение головы
- » Выражение лица
- » Визуальный контакт
- » Жесты

Конечно, есть и другие характеристики, но даже если искусственный интеллект сможет распознать эти пять областей, ему еще предстоит пройти длинный путь для поддержки правильной интерпретации жестикуляции. Кроме жестикуляции, текущие реализации искусственного интеллекта учитывают тембр голоса, что приближает чрезвычайно сложный искусственный интеллект к тому, что человеческий мозг делает, по-видимому, без усилий.



Как только искусственный интеллект сможет читать жесты, он должен будет также иметь средства для их воспроизведения при общении с людьми. В то время как чтение жестов находится еще в младенчестве, роботизированное или графическое представление жестикуляции еще даже не разрабатывается. Статья <https://spectrum.ieee.org/video/robotics/roboticssoftware/robots-learn-to-speak-body-language> указывает, что роботы в настоящее время могут интерпретировать жесты, а затем реагировать соответственно в некоторых случаях. Современные роботы не способны хорошо имитировать выражение лица, поэтому согласно статье <http://theconversation.com/realistic-robot-faces-arentenough-we-need-emotion-to-put-us-at-ease-with-androids-43372> они в лучшем случае должны реализовать позу, движение головы и жесты. Пока результат внушительным назвать нельзя.

Обмен идеями

У искусственного интеллекта нет идей, потому что у него отсутствуют и внутриличностный интеллект, и способность что-либо понимать. Однако искусственный интеллект вполне может позволить людям обмениваться идеями способом, куда лучшим, чем сумма его составляющих. Как правило, искусственный интеллект обменом не занимается. Обмен осуществляют люди, вовлеченные в некое занятие, полагаясь на искусственный интеллект для улучшения процесса коммуникации. Следующие разделы знакомят с дополнительными деталями этого процесса.

Установление связи

Один человек может обмениваться идеями с другим человеком, только когда эти два человека знают друг о друге. Проблема в том, что многие эксперты в конкретной области фактически не знают друг друга, по крайней мере знают недостаточно хорошо, чтобы общаться. Искусственный интеллект способен выполнить исследование на основании предоставленного человеком

направления мысли, а затем связать его с другими людьми со схожим (или подобным) направлением.

Одним из способов создания таких связей являются сайты социальных сетей, такие как LinkedIn (<https://www.linkedin.com/>), идея которого в том, чтобы объединять людей на основании многих критериев. Социальная сеть становится средством, которое искусственный интеллект глубоко внутри системы LinkedIn предлагает для потенциальных связей. В конечном счете цель этих связей, с точки зрения пользователя, заключается в том, чтобы обеспечить доступ к новым человеческим ресурсам, установить деловые контакты, обеспечить продажу и выполнять другие задачи, используя различные связи, установить которые позволяет LinkedIn.

Улучшение коммуникаций

Для успешного обмена идеями два человека должны общаться хорошо. Единственная проблема в том, что люди иногда не общаются хорошо, а иногда они не общаются вообще. Проблема кроется в переводе не только слов, но и идей. Социальные и личные пристрастия людей могут препятствовать их общению, поскольку идея одной группы может быть вообще непонятна для другой. Например, законы в одной стране могут заставить кого-то думать одним способом, а законы в другой стране могут заставить другого человека думать совершенно иначе.

Теоретически искусственный интеллект может помочь общению между несоизмеримыми группами многочисленными способами. Конечно, перевод (если он точен) является одним из этих методов. Но искусственный интеллект мог бы предварительно маркировать материалы как культурно приемлемые и напротив. Используя классификацию, искусственный интеллект может также предложить такие средства, как альтернативная графика, чтобы помочь общению способом, который понравится обеим сторонам.

Определение тенденций

Идеи нередко возникают на базе тенденций. Но чтобы показать, как работает идея, другие стороны по обмену идеями должны также видеть эти тенденции, и передача информации этого вида связана с известными трудностями. Искусственный интеллект может выполнять анализ данных на разных уровнях и представлять графический вывод. Искусственный интеллект способен проанализировать данные большим количеством способов и куда быстрее, чем человек, чтобы рассказываемая данными история была именно такой, как требуется. Данные останутся теми же, но их представление и интерпретация будут нужными.

Практика показывает, что информацию в виде графики люди воспринимают лучше, чем в виде таблиц, и графический вывод, определенно, делает тенденции понятнее. Как описано на сайте <http://sphweb.bumc.bu.edu/otlt/mph-modules/bs/datapresentation/DataPresentation2.html>, табличные данные лучше использовать для представления только специфической информации; для демонстрации тенденций всегда лучше подходит графика. Применение управляемого искусственным интеллектом приложения позволит оформить графические данные правильно и согласно конкретным требованиям. Не все люди воспринимают графику одинаково, поэтому выбор наилучшего типа графиков для своей аудитории очень важен.

Использование средств массовой информации

Большинство людей учится, используя несколько органов чувств и несколько подходов. Подход, хорошо зарекомендовавший себя при обучении одних, может совсем не подходить другим. Следовательно, чем больше путей передачи идей и концепций, тем более вероятно, что другие люди поймут то, что человек пытается сообщить. Обычно средства массовой информации используют текст, звук, графику и анимацию, но некоторые из них идут дальше.

СРЕДСТВА МАССОВОЙ ИНФОРМАЦИИ И СПЕЦИАЛЬНЫЕ ПОТРЕБНОСТИ

У большинства людей есть некие специальные потребности, и их учет средствами массовой информации важен для этих людей. Главная цель средств массовой информации — выражать мысли как можно более многочисленными способами, чтобы почти все могли понять представляемые идеи и концепции. Даже при успешном использовании средств массовой информации в целом отдельные идеи могут теряться, когда для их представления используют только один метод. Например, при только устном оповещении хорошее представление об идее получают почти наверняка только люди с хорошим слухом. Из людей, обладающих хорошим слухом, также не все осознают идею, поскольку может быть слишком шумно или они просто плохо воспринимают на слух. Очень важно использовать по возможности больше методов выражения каждой мысли, если хотите оповестить о них как можно больше людей.

Искусственный интеллект может помочь средствам массовой информации множеством способов. Один из самых важных заключается в создании содержимого или разработке авторских материалов. Искусственный интеллект можно найти в приложениях, помогающих во всем, от создания текстов до публикации в средствах массовой информации. Например, цветокоррекция изображений с использованием искусственного интеллекта может обеспечить необходимые эффекты и помочь визуализировать нужное быстрее, чем попытка сделать это самостоятельно одной цветовой комбинацией.

После применения средств массовой информации для публикации идеи в нескольких формах узнавшие об идее люди должны обработать информацию. Здесь искусственный интеллект применяется снова, теперь уже используя нейронные сети для обработки информации различными способами. Категоризация средств массовой информации является немаловажным элементом современных технологий. Но в будущем искусственный интеллект, возможно, обеспечит трехмерное воссоздание сцен на основании двухмерных изображений. Вообразите полицию, способную воссоздать виртуальную сцену преступления с каждой деталью, воспроизведенной абсолютно точно.

Под новыми формами средств массовой информации люди обычно понимают нечто фантастически новое. Вспомните, например, газету из фильмов о Гарри Поттере с динамическими фотографиями. Большинство элементов подобных технологий фактически доступно уже сегодня, но проблема одна — рынок. Чтобы технология стала успешной, у нее должен быть рынок, т.е. шанс на самоокупаемость.

Улучшение сенсорного восприятия человека

В чем искусственный интеллект действительно великолепен, так это в улучшении человеческого общения за счет усиления людей одним из двух способов: позволив им использовать свои органы чувств для работы с более сложными данными или существенно усилив органы чувств. В следующих разделах обсуждаются оба подхода к усилению человеческого восприятия, а следовательно, и к улучшению коммуникабельности.

Смещение спектра данных

При сборе различных видов информации люди, как правило, используют технологии, фильтрующие или сдвигающие спектр данных по цвету, звуку или запаху. Человек все еще использует свои естественные возможности, но технология немного изменяет исходные данные так, чтобы эта естественная возможность с ними работала. Одним из наиболее популярных примеров смещения

спектра является астрономия, в которой фильтрация и смещение частот источников света позволяют людям видеть такие астрономические объекты, как туманности, которые невооруженным человеческим глазом не видны. Это существенно улучшает наше понимание Вселенной.

Однако смещение и фильтрация цветов, звуков или запахов вручную может потребовать большого количества времени, а результаты могут разочаровать, даже если обработка выполнена квалифицированно. Вот где искусственный интеллект может сыграть свою роль. Он может опробовать различные комбинации куда быстрее, чем человек, и куда быстрее найти потенциально полезные комбинации, поскольку решает задачу единообразным способом.



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Кстати, наиболее интригующие методики исследования мира полностью отличны от ожидаемых большинством людей. Что если бы вы могли чувствовать запах цвета или видеть звук? Практика *синестезии* (synesthesia) (http://www.science20.com/news_releases/synaesthesia_smelling_a_sound_or_hearing_a_color), когда один орган чувств интерпретирует ввод от другого, хорошо известна. Люди используют искусственный интеллект для получения такого результата, как тот, который описан по адресу <http://journals.plos.org/ploscompbiol/article?id=10.1371/journal.pcbi.1004959>. Тем не менее использование этой технологии интересно; она создает условия, в которых люди фактически могут использовать синестезию как другое средство восприятия мира (см. <https://www.fastcompany.com/3024927/this-app-aids-yourdecision-making-by-mimicking-its-creators-synesthesia>). На случай, если вы хотите увидеть, как использовать синестезию, посмотрите приложение ChoiceMap на странице <https://choicemap.co/>.¹

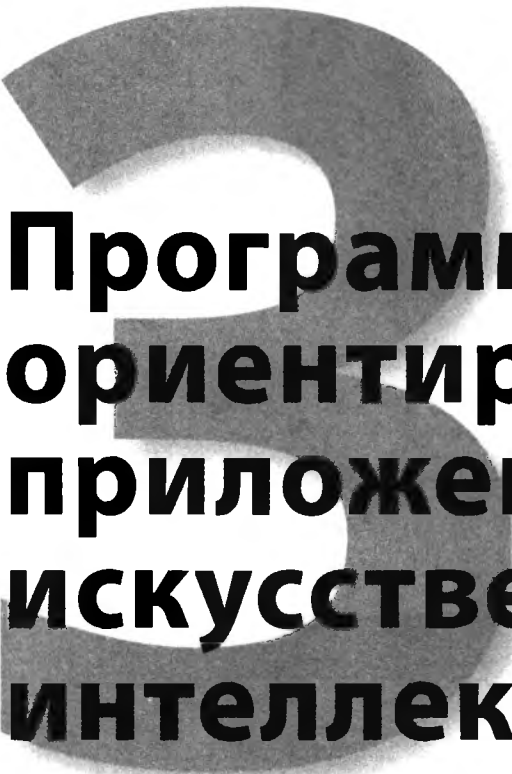
Усиление человеческих органов чувств

Альтернативой использованию внешнего приложения для сдвига спектра данных, позволяющего сделать их доступными для использования людьми, является усиление их органов чувств. Устройство усиления, внешнее или внедренное, позволит человеку непосредственно обрабатывать сенсорный ввод новым способом. Многие люди рассматривают эти новые возможности как попытку создания киборгов (см. страницу <https://www.theatlantic.com/technology/archive/2017/10/cyborgfuture-artificial-intelligence/543882/>). Идея стара: инструменты позволяют сделать людей более способными к выполнению разнообразных задач. В этом сценарии люди получают две формы усиления: физическую и интеллектуальную.

¹ Актуальность ссылки не гарантируется. — *Примеч. ред.*

Физическое усиление человеческих органов чувств уже имеет место разными способами, и, как можно убедиться, благодаря различным видам дополнений люди становятся более восприимчивыми. Например, современные очки ночного видения позволяют людям видеть ночью, а несколько более совершенные модели обеспечивают даже цветное зрение, контролируемое специально для этого разработанным процессором. В будущем усиление (или замена) глаз может позволить людям видеть в любой части спектра и контролироваться мыслью, чтобы люди видели только ту часть спектра, которая нужна для выполнения конкретной задачи.

Усиление интеллекта (Intelligence Augmentation — IA) требует более радикальных мер, но обещает намного большее увеличение возможностей. В отличие от искусственного интеллекта, в центре усиления интеллекта находится человек. Человек предоставляет творческий потенциал и намерение, отсутствующие в настоящее время у искусственного интеллекта. О различии между искусственным интеллектом и усилением интеллекта можно прочитать на странице <https://www.financialsense.com/contributors/guild/artificialintelligence-vs-intelligence-augmentation-debate>.



Программно- ориентированные приложения искусственного интеллекта

В ЭТОЙ ЧАСТИ...

- » Анализ данных**
- » Отношения между искусственным интеллектом и машинным обучением**
- » Отношения между искусственным интеллектом и глубоким обучением**



Глава 9

Анализ данных для искусственного интеллекта

В ЭТОЙ ГЛАВЕ...

- » Как работает анализ данных
- » Эффективный анализ данных с использованием машинного обучения
- » На что способно машинное обучение
- » Виды алгоритмов машинного обучения

Накопление данных не является современным явлением; люди накапливали данные на протяжении многих столетий. Независимо от текстовой или числовой формы информации люди всегда ценили данные, описывающие окружающий мир, и продолжают использовать их для прогресса цивилизации. Данные значимы сами по себе. Используя их, человечество может изучать критически важную информацию и передавать ее потомкам (нет необходимости повторно изобретать колесо).

Люди лишь недавно узнали, что данные могут содержать куда больше информации, чем кажется на первый взгляд. Если данные находятся в соответствующей числовой форме, к ним можно применить специальные методики,

разработанные математиками и статистиками. Эти методики анализа данных позволяют извлечь из них куда больше информации. Кроме того, даже простой анализ данных позволяет извлечь из них осмысленную информацию, а подвергнув данные более совершенному анализу, с использованием алгоритмов машинного обучения можно предсказывать будущее, классифицировать информацию и эффективно принимать решения.

Анализ данных и машинное обучение позволяют перейти на следующий уровень использования данных, теперь — для разработки более умного искусственного интеллекта. Эта глава посвящена анализу данных. Она демонстрирует применение данных в качестве инструмента обучения при решении сложных проблем искусственного интеллекта, таких как правильная рекомендация товара клиенту, понимание разговорного языка и языкового перевода, автоматизация вождения автомобиля и многие другие.

Анализ данных

Наше время называют веком информации не просто потому, что сейчас накоплено богатое разнообразие данных, но и потому, что общество достигло определенной зрелости в анализе данных и извлечении из них информации. Такие компании, как Alphabet (Google), Amazon, Apple, Facebook и Microsoft (пять самых дорогих компаний в мире), построили свой бизнес на данных. Они не просто собирают и хранят данные, полученные в результате цифровых процессов; они знают, как, используя точный и сложный анализ данных, сделать их такими же ценными, как нефть. Компания Google, например, собирает данные не только из веба вообще, но, между прочим, и из собственного поискового механизма.

Вы, возможно, уже встречали в новостях, журналах или на конференциях расхожую фразу “Данные — это новая нефть”. Она подразумевает, что данные могут сделать компанию богатой, но для этого придется тяжело и эффективно работать. Хотя эту концепцию использовали многие и сделали ее невероятно успешной, именно британский математик Клайв Хамбли (Clive Humby) впервые приравнял данные к нефти на основании своей практики с данными о потребителях в розничном секторе. Хамбли известен тем, что был среди основателей британской торговой компании Dunhumby; его идеи также легли в основу программы дисконтных карт Tesco. В 2006 году Хамбли также подчеркнул, что данные — это не просто деньги, которые падают с неба; чтобы сделать их полезными, требуются усилия. Подобно тому, как нельзя непосредственно использовать неочищенную нефть (ее следует превратить в ходе химических процессов в бензин, пластмассы или другие химикаты), данные также следует существенно переработать, чтобы они приобрели значимость.

Самые простые преобразования данных — это *анализ данных* (data analysis); вы можете считать его простым химическим преобразованием, которым нефть очищается на заводе прежде, чем стать ценным топливом или пластмассой. Используя подходящий анализ данных, вы можете заложить фундамент для более сложных аналитических процессов. Анализ данных, в зависимости от контекста, сводится к большому количеству возможных операций, иногда специфичных для конкретной отрасли или задачи. Все эти преобразования можно отнести к четырем основным категориям, концептуально отличающимся происходящим во время анализа.

- » **Преобразование.** Изменяет внешний вид данных. Термин *преобразование* применим к разным процессам, хотя данные, как правило, помещает в упорядоченные ряды и столбцы — *матричный формат* (или *двумерный файл*). Например, вы не можете эффективно обработать данные о товарах, купленных в супермаркете, прежде чем поместите каждого клиента в отдельный ряд и добавите купленные товары в один столбец в пределах этого ряда в виде числовых элементов, содержащих значения количества или платы. Чтобы сделать набор данных подходящим для алгоритма, могут также потребоваться специальные числовые преобразования, такие как *масштабирование*, вычисление *среднего*, минимального и максимального значений.
- » **Чистка.** Исправление дефектных данных. В зависимости от средств сбора данных могут возникнуть различные проблемы с отсутствием информации, выбросами значений из диапазона или просто неправильными значениями. Например, данные из супермаркета могут содержать ошибки, если у товаров неправильные ценники. Некоторые данные могут быть *подложными*, т.е. созданными специально, чтобы исказить заключение. Например, у товара могут быть поддельные отзывы в Интернете, которые изменят его ранг. Чистка помогает удалить подложные случаи из данных и сделать заключение объективней.
- » **Проверка.** Проверка данных. Анализ данных — это по большей части человеческая работа, хотя программное обеспечение играет в ней важную роль. Люди могут легко распознавать шаблоны и выявлять странные элементы данных. Поэтому анализ данных подразумевает множество статистических выкладок и графических представлений, таких как у Health InfoScape от MIT Senseable City Lab и General Electric (<http://senseable.mit.edu/healthinfoscape/>), позволяющих сразу схватить информативное содержимое. Например, на основании обработанных данных из 72 миллионов медицинских записей можно увидеть, как взаимосвязаны болезни.

» **Моделирование.** Выявление отношений между элементами в данных. Для решения этой задачи необходимы такие статистические инструменты, как корреляции, t-проверки, линейная регрессия и многие другие, позволяющие определить, действительно ли одно значение не зависит от другого или они взаимосвязаны. Например, анализируя продажи супермаркета, вы можете прийти к мнению, что люди, покупающие подгузники, имеют тенденцию покупать и пиво. Статистический анализ считает эти два товара взаимосвязанными, поскольку они многократно обнаруживаются в одних и тех же корзинах. (Это исследование — настоящая легенда аналитики данных; см. рассказ об этой истории в статье *Forbes* по адресу <https://www.forbes.com/global/1998/0406/0101102s1.html>).

Магии в анализе данных нет. Вы осуществляете преобразования, очистку, проверку и моделирование, используя суммирование и умножение массивов на основании матричного исчисления (представляющего собой не более чем длинные последовательности суммирований и умножений, которым людей обычно учат в школе). Арсенал анализа данных включает и такие статистические инструменты, как поиск среднего и дисперсии, описывающие распределение данных, и такие сложные инструментальные средства, как корреляция и линейный регрессионный анализ, показывающие, можно ли связать между собой некие события (такие, как покупка подгузников и пива) на основании доказательств. Более подробная информация об этих методиках обработки данных приведена в книгах *Machine Learning For Dummies* и *Python for Data Science For Dummies* Джона Пола Мюллера и Луки Массарона, практически представляющих собой краткий обзор и объяснение каждой из них.



ЗАПОМНИ

Анализ данных существенно сложнее в случае их больших объемов. Для этого требуются специальные инструментальные средства, такие как Hadoop (<http://hadoop.apache.org/>) и Apache Spark (<https://spark.apache.org/>). Эти два программных инструмента применяются для работы с большими массивами данных. Несмотря на такие передовые инструменты, как эти, все еще остается вопрос пота: до 80 процентов данных приходится готовить вручную. Вызывает интерес интервью с Моникой Рогати (Monica Rogati) (см. <https://www.nytimes.com/2014/08/18/technology/for-big-data-scientists-hurdle-to-insights-is-janitor-work.html>) — экспертом и советником в области искусственного интеллекта многих компаний, обсуждающей эту проблему более подробно.

Почему анализ столь важен

Анализ данных важен для искусственного интеллекта. Фактически никакой современный искусственный интеллект невозможен без визуализации, очистки, преобразования и моделирования данных прежде, чем передовые алгоритмы вступят в игру и поднимут значимость информации на куда более высокий уровень, чем прежде.

В начале, когда искусственный интеллект состоял просто из алгоритмических решений и экспертных систем, ученые и эксперты тщательно готовили передаваемые ему данные. Поэтому, если некто хотел, например, чтобы алгоритм сортировал информацию, эксперт по данным помещал данные в *списки* (упорядоченную последовательность элементов данных) или другие структуры данных, которые могли содержать информацию и позволять манипулировать ею желательным образом. Затем эксперты по данным собирали и организовывали данные так, чтобы их содержимое и форма были точно такими, как ожидалось, согласно конкретной цели, для которой они были созданы. Манипулирование известными данными в специфической форме налагало серьезные ограничения, поскольку обработка данных требовала много времени и энергии; а следовательно, алгоритмы получали меньше информации, чем доступно сегодня.

Сегодня внимание сместилось с создания данных на их подготовку для анализа. Дело в том, что различные источники уже производят данные в таких больших количествах, что в них уже можно найти все нужное без необходимости создавать данные для задачи специально. Представьте, например, что искусственный интеллект должен контролировать дверцу для домашнего животного в двери жилого дома, чтобы впускать вашего кота или собаку, но не других животных. Современные алгоритмы искусственного интеллекта обучаются на основании специфических для задачи данных, а значит, предстоит обработка больших количеств изображений с примерами собак, котов и других животных. Вероятнее всего, такой огромный набор изображений поступит из Интернета, возможно, с социальных сайтов или поисковиков изображений. Ранее выполнение подобной задачи означало, что алгоритмы используют лишь несколько изображений, чтобы получить исходные данные о форме, размере и отличительных характеристиках животных. Недостаток данных означал возможность выполнения весьма ограниченных задач. Фактически нет никаких примеров того, что искусственный интеллект мог контролировать дверь для животных, используя классические алгоритмы или экспертные системы.

Анализ данных приходит на помощь современным алгоритмам, предоставляя информацию об изображениях, полученных из Интернета. Анализ данных позволяет искусственному интеллекту отобрать изображения по размеру, разнообразию, количеству цветов, слов в их подписях и т.д. Это этап проверки данных, и в данном случае он необходим для их очистки и преобразования.

Например, анализ данных может помочь определить фотографию животного, ошибочно помеченную как кот¹ (вы не хотите запутать свой искусственный интеллект), и преобразовать изображения так, чтобы использовать одинаковый формат цвета (например, оттенки серого) и одинаковый размер.

Пересмотр значения данных

Взрывообразное повышение доступности данных для цифровых устройств (как обсуждалось в главе 2) придало им новые аспекты значимости в дополнение к первоначальным возможностям для обучения и передачи знаний. Изобилие предоставляемых на анализ данных формирует новые функции, отличные от только информативных.

- » Данные описывают мир лучше, когда предоставляется широкое разнообразие фактов при более высокой подробности по каждому факту. Их стало так много, что учитывается практически каждый аспект действительности. Вы можете использовать их для представления взаимосвязи между даже, казалось бы, никак не связанными вещами и фактами.
- » Данные демонстрируют, как факты связаны с событиями. Вы можете вывести общие правила и узнать, как изменится мир с учетом неких конкретных предположений. Когда люди действуют определенным образом, данные обеспечивают также определенную способность прогнозирования.

В некотором отношении данные дают нам суперсилу. Крис Андерсон (Chris Anderson), предыдущий главный редактор *Wired*, обсуждает, как большие объемы данных могут помочь в научных исследованиях даже вне пределов научного метода (см. статью <https://www.wired.com/2008/06/pb-theory/>). Автор приводит примеры достижений Google в бизнес-секторах рекламы и перевода, в которых компания Google достигла выдающегося положения не за счет использования специфических моделей или теорий, а скорее, за счет применения алгоритмов для обучения на основе данных.

Как и в рекламе, научные данные (из физики или биологии) могут обеспечить новшество, которое позволит ученым подойти к проблеме без гипотез, а рассматривать варианты найденные в больших объемах данных с использованием поисковых алгоритмов. Галилео Галилей полагался на научный метод для создания современной физики и астрономии (см. <https://www.biography.com/people/galileo-9305220>). Первые достижения полагались на наблюдения и контролируемые эксперименты, позволяющие выяснить причины

¹ См. рис. 11.3. — Примеч. ред.

происходящего. Возможность открытий при использовании одних только данных является главным прорывом в способе нашего познания мира.

В прошлом ученые делали бесчисленные наблюдения и с помощью дедукции описывали физику Вселенной. Этот ручной процесс позволил открыть основные законы мира, в котором мы живем. Анализ данных, когда наблюдения выражены как вводы и выводы, позволяет определять, как нечто работает, и выяснять, по каким примерно правилам или законам нашего мира, без проведения экспериментов вручную и дедукции. Теперь процесс протекает быстрее и с большим процентом автоматического выполнения.

ОТКРЫТИЕ БОЛЕЕ УМНОГО ИСКУССТВЕННОГО ИНТЕЛЛЕКТА ЗАВИСИТ ОТ ДАННЫХ

Наличие больших количеств данных не сделает искусственный интеллект возможным. Некоторые люди сказали бы, что искусственный интеллект — это результат выполнения сложных математических алгоритмов, и это, конечно, так. Такие действия, как зрение и понимание речи, требуют алгоритмов, которые нелегко объяснить в непрофессиональных терминах, и миллионов вычислений в секунду. (Аппаратные средства здесь также имеют значение.)

И все же искусственный интеллект — это даже больше, чем алгоритмы. Д-р Александер Уисснер-Гросс (Alexander Wissner-Gross), американский исследователь, предприниматель и действительный член Гарвардского института прикладной информатики, продемонстрировал свою способность проникать в суть в недавнем интервью Edge (<https://www.edge.org/response-detail/26587>). Он размышляет, почему технологии искусственного интеллекта потребовалось так много времени для взлета. В интервью Уисснер-Гросс заключает, что это, возможно, был вопрос качества и доступности данных, а не алгоритмов.

Уисснер-Гросс рассматривает эволюцию большинства революционных достижений искусственного интеллекта за последние годы, демонстрируя, как данные и алгоритмы способствовали успеху каждого технического прорыва, и подчеркивает, что каждый из них не зависел от времени достижения новой вехи. Уисснер-Гросс демонстрирует, что данные всегда относительно новы и обновляемы, в отличие от алгоритмов, которые скорее полагаются на консолидацию прежних технологий.

В заключение Уисснер-Гросс упоминает, что в среднем алгоритмы на 15 лет старше данных. Он указывает, что данные способствуют успеху искусственного интеллекта, и ставит перед читателем вопрос, что было бы, если бы доступным ныне алгоритмам были предоставлены лучшие данные с точки зрения качества и количества.

Машинное обучение

Апофеозом анализа данных является машинное обучение. Вы можете успешно применить машинное обучение только после того, как анализ данных предоставляет правильные исходные данные. Но только машинное обучение способно ассоциировать наборы выходных и входных данных, а также эффективно выявить использованные при этом правила. Анализ данных концентрируется на понимании и манипулировании данными таким способом, чтобы они могли стать более полезными и способными обеспечить проникновение в суть вещей, тогда как машинное обучение строго сосредоточено на том, чтобы получить исходные данные и, сделав свою работу, выработать внутреннее представление о сути вещей, которое можно использовать практически. Машинное обучение позволяет людям решать такие задачи, как предсказание будущего, классификация осмысленным способом и выработка наиболее рационального решения в данном контексте.



ЗАПОМНИ

Главная идея, лежащая в основе машинного обучения, — можно представить действительность, используя математические функции, которые алгоритму заранее неизвестны, но которые он может предположить на основании наблюдения некоторых данных. Вы можете выразить действительность и всю ее комплексную сложность в терминах неизвестных математических функций, которые алгоритмы машинного обучения способны выявить и сделать доступными. Эта концепция лежит в основе идей всех видов алгоритмов машинного обучения.

Обучение при машинном обучении является просто математической функцией, и она заканчивается ассоциацией определенных входных данных с определенными выходными. Она не имеет отношения к пониманию того, что алгоритм изучил (анализ данных вырабатывает понимание до некоторой степени), поэтому процесс обучения зачастую называют *тренировкой*, поскольку алгоритм учится находить правильный ответ (вывод) на каждый предоставленный вопрос (ввод). Более подробная информация по этому вопросу приведена в книге *Machine Learning For Dummies* Джона Пола Мюллера и Луки Массарона (издательство Wiley).

Несмотря на отсутствие осмысленного понимания, поскольку это просто математический процесс, машинное обучение может оказаться весьма полезным. Оно дает приложению с искусственным интеллектом силу наибольшей рациональности в текущем контексте, когда обучение происходит с использованием правильных данных. В следующих разделах работа машинного

обучения описана более подробно, включая ожидаемые преимущества и пределы его применения в приложениях.

Как работает машинное обучение

Многие люди привыкли к идее, что приложения начинают работу с получения данных на входе, а затем предоставляют некий результат. Например, программист мог бы создать функцию `Add()`, которая получает как ввод два значения, например 1 и 2, а затем возвращает результат, 3. Выводом этого процесса является значение. В прошлом написание программы означало понимание того, как функция должна манипулировать данными, чтобы получить заданный результат при определенных входных данных. Машинное обучение осуществляет переворот в этом процессе. В данном случае вы знаете входные данные, это 1 и 2. Вы также знаете, что желаемый результат — 3. Однако вы не знаете, какую функцию применить, чтобы получить желаемый результат. Обучение предоставляет алгоритм со всякого рода примерами входных данных и ожидаемыми результатами для этих входных данных. Затем алгоритм использует заданный ввод, чтобы создать функцию. Другими словами, обучение — это процесс, в ходе которого обучаемый алгоритм сопоставляет с данными гибкую функцию. Выводом обычно является вероятность определенного класса или числового значения.

Чтобы дать общее представление о происходящем в ходе учебного процесса, вообразите ребенка, учащегося отличать деревья от других объектов. Прежде чем ребенок сможет сделать это самостоятельно, учитель показывает ему изображения деревьев, отображающие все факты, отличающие дерево от других объектов. К таким фактам могут относиться материал дерева (древесина), его части (ствол, ветви, листья и корни) и расположение (в почве). Ребенок вырабатывает представление о том, как выглядит дерево, в отличие от изображений других объектов, таких как предметы мебели, которые тоже состоят из древесины, но не имеют других характеристик, схожих с характеристиками дерева.

Классификатор машинного обучения работает точно так. Он формирует свои когнитивные способности, создавая математические формулировки, включающие все заданные средства, способом, который определяет функцию, способную отличить один класс от другого. Предположим, что существует некая математическая формулировка, *целевая функция* (*target function*), способная выразить характеристики дерева. В таком случае классификатор машинного обучения может искать свое представление как ее реплику или приближение (функция иная, но работает похоже). Способность выразить такую математическую формулировку является возможностью представления классификатора.

С математической точки зрения вы можете выразить процесс представления в машинном обучении, используя сопоставление эквивалентного термина. Сопоставление происходит, когда вы обнаруживаете конструкцию функции, наблюдая ее вывод. Успешное сопоставление в машинном обучении подобно ребенку, усваивающему идею объекта. Ребенок понимает абстрактные правила, следующие из фактов реального мира, поэтому когда ребенок видит дерево, например, он сразу его узнает.

Такое представление (абстрактные правила, проистекающие из реальных фактов) возможно потому, что у алгоритма обучения есть множество внутренних параметров (состоящих из векторов и матрицы значений), эквивалентных памяти алгоритма для идей, которые лучше всего подходят для ассоциации средств с классами ответа. Размерности и тип внутренних параметров разграничивают вид целевых функций, которые алгоритм может изучать. Чтобы выяснить скрытую целевую функцию, механизм оптимизации алгоритма во время обучения изменяет значения параметров, начиная с их исходных значений.

Во время оптимизации алгоритм ищет возможные варианты комбинаций параметров, чтобы найти ту, при которой возможно правильное сопоставление средств и классов при обучении. Этот процесс вычисляет множество возможных целевых функций — потенциальных кандидатов из числа тех, которые может предположить обучающий алгоритм. Набор всех потенциальных функций, которые смог обнаружить обучающий алгоритм, является *пространством гипотез* (hypothesis space). Вы можете вызвать результирующий классификатор с его параметрами для набора гипотез в ходе машинного обучения, чтобы сказать, что алгоритм установил параметры для репликации целевой функции и теперь готов определить правильные классификации (факт, демонстрируемый позже).

Пространство гипотез должно содержать все варианты параметра всех алгоритмов машинного обучения, которые вы хотите попробовать сопоставить с неизвестной функцией при решении проблемы классификации. У различных алгоритмов могут быть разные пространства гипотез. Действительно имеет значение то, что пространство гипотез содержит целевую функцию (или ее подобие, поскольку, в конце концов, необходимо нечто работоспособное).

Можете считать эту фазу временем, когда ребенок экспериментирует со многими разными творческими идеями, накапливая знания и опыт (аналогия получению средств), чтобы получить представление о дереве. Естественно, на этой фазе задействованы родители, и они предоставляют корректные исходные данные об окружающей среде. В машинном обучении кто-то должен предоставить правильные алгоритмы обучения, некие не учебные параметры (гиперпараметры), выбрать набор примеров для изучения, а также выбрать сопутствующие примерам средства. Подобно тому, как ребенок не всегда может

научиться различать, что правильно и что неправильно, если он остается в изоляции, алгоритмы машинного обучения нуждаются в людях, чтобы учиться успешно.

Преимущества машинного обучения

Сегодня вы найдете искусственный интеллект и машинное обучение используемыми в очень многих приложениях. Единственная проблема — технология работает настолько хорошо, что вы даже не знаете, что она существует. Фактически вы могли бы быть немало удивлены, обнаружив, что многие устройства в вашем доме уже используют обе технологии. Обе технологии, определенно, присутствуют в вашем автомобиле и на рабочем месте. Фактически случаи использования и искусственного интеллекта, и машинного обучения исчисляются миллионами, и все без очевидных опасностей, даже когда их характер весьма критичен. Вот только некоторые из способов, которыми искусственный интеллект может использоваться.

- » **Обнаружение мошенничества.** Вы получаете из банка запрос, платили ли вы своей кредитной карточкой за определенную покупку. Банк не любопытен; он просто обеспокоен тем фактом, что некто мог сделать покупку, используя вашу карточку. Искусственный интеллект банка обнаружил незнакомый шаблон расходов и предупредил кого следует.
- » **Планирование ресурсов.** Многим организациям необходимо эффективное планирование использования ресурсов. Например, больнице, вероятно, придется решать, куда поместить пациента, исходя из его потребностей, доступности квалифицированного персонала и ожидаемого периода пребывания пациента в больнице.
- » **Комплексный анализ.** Люди нередко нуждаются в комплексном анализе, когда приходится учитывать слишком много факторов. Например, один и тот же набор симптомов может свидетельствовать о нескольких проблемах. Чтобы спасти пациенту жизнь, врачу или другому специалисту может понадобиться помощь в своевременной постановке диагноза.
- » **Автоматизация.** Любая форма автоматизации может извлечь пользу из применения искусственного интеллекта для реакции на непредвиденные изменения или события. Проблема некоторых типов автоматизации сегодня заключается в том, что непредвиденное событие, такое как нахождение объекта в неподобающем месте, может фактически привести к ее отказу. Добавление искусственного интеллекта к автоматизации может позволить справиться с непредвиденным событием и продолжить работу, как будто ничего не случилось.

- » **Клиентская служба.** На другом конце линии клиентской службы, в которую вы сегодня звонили, вполне могло и не быть человека. Автоматизация ныне достаточно хороша, чтобы, используя различные сценарии и ресурсы, справиться с подавляющим большинством ваших вопросов. При корректных интонациях голоса (также предоставляемых искусственным интеллектом) вы не сможете даже с уверенностью сказать, что говорили с компьютером.
- » **Системы безопасности.** Большинство систем безопасности современных машин различных видов полагается на искусственный интеллект, чтобы перехватить управление транспортным средством в критической ситуации. Например, многие автоматические тормозные системы полагаются на искусственный интеллект, чтобы остановить автомобиль, исходя из всех исходных данных, которые может предоставить транспортное средство, например направления заноса.
- » **Эффективность машин.** Искусственный интеллект может помочь контролировать машину так, чтобы добиться максимальной эффективности. Искусственный интеллект контролирует использование ресурсов, чтобы система не превышала скорости и для других задач. Каждая унция мощности будет использована точно так, как необходимо, чтобы оказать желаемую услугу.

Этот список не затрагивает даже части возможностей. Вы можете найти искусственный интеллект используемым и многими другими способами. Однако полезно также узнать об использовании машинного обучения и вне областей, традиционных для искусственного интеллекта. Вот несколько случаев использования машинного обучения, которое могли бы и не быть связаны с искусственным интеллектом.

- » **Управление доступом.** Во многих случаях доступ либо разрешают, либо запрещают. Смарт-карта сотрудника предоставляет доступ к ресурсу точно так же, как и ключ, многие столетия назад. Некоторые замки действительно позволяют задавать время и даты, когда их можно открывать, но такой контроль — не ответ на каждую потребность. Используя машинное обучение, можно определить, может ли сотрудник получить доступ к данному ресурсу, исходя из его роли и потребности. Например, сотрудник может получить доступ к учебной комнате, если обучение является его ролью.
- » **Защита животных.** Океан может казаться достаточно большим, чтобы морские животные и суда могли без проблем в нем сосуществовать. К сожалению, суда ежегодно ранят множество животных. Алгоритм машинного обучения может позволить судам избегать животных, изучив издаваемые ими звуки и другие характеристики животных и судов.

- » **Предсказание времени ожидания.** Большинству людей не нравится ждать, когда они понятия не имеют, как долго придется это делать. Машинное обучение позволяет приложению определить время ожидания на основании уровня персонала, их загрузки, сложности решаемых ими проблем, доступности ресурсов и т.д.

Будешь полезным — станешь обыденным

Хотя в фильмах и утверждается, что искусственный интеллект приведет к огромному подъему, и вы действительно иногда видите невероятные случаи его использования в реальности, применяется он, как правило, для решения обыденных и даже скучных задач. К примеру, недавно компания Verizon использовала язык R для машинного обучения, чтобы проанализировать данные о нарушении правил безопасности и автоматизировать составление ежегодных отчетов о них (<https://www.computerworld.com/article/3001832/dataanalytics/how-verizon-analyzes-security-breach-data-with-r.html>). Подобный анализ — не столь грандиозное свершение по сравнению с использованием других видов искусственного интеллекта, но, используя язык R, компания Verizon экономит деньги, выполняя анализ, и получает лучшие результаты.

Пределы машинного обучения

Для анализа огромных наборов данных машинное обучение полагается на алгоритмы. В настоящее время машинное обучение не способно обеспечить такой вид искусственного интеллекта, как в кино. Даже наилучшие алгоритмы не могут думать, чувствовать, обладать любой формой самосознания или иметь свободу выбора. На что машинное обучение способно, так это выполнять прогнозирующую аналитику намного быстрее, чем любой человек. В результате машинное обучение может помочь людям работать эффективнее. Да, искусственный интеллект в текущем состоянии способен выполнить выдающийся анализ, но смысл его результатов все еще должны осознавать люди, они же принимают необходимые моральные и этические решения на его основании. По существу машинное обучение обеспечивает лишь часть обучения искусственного интеллекта, и эта часть ничуть не готова создать искусственный интеллект того вида, который вы видите в фильмах.

Основная причина несоответствия между обучением и интеллектом — в человеческом предположении, что простой способности машины справляться со своей работой (обучение) уже достаточно для сознания (интеллект). Это предположение ничем не подтверждено для машинного обучения. То же самое происходит, когда люди полагают, что компьютер преднамеренно создает для них проблемы. Компьютер не имеет эмоций и поэтому действует только на основании предоставленных данных и инструкций для их обработки. Истинный

искусственный интеллект получится только тогда, когда компьютеры, наконец, смогут подражать следующей сложной комбинации, используемой в природе.

- » **Генетика.** Медленное обучение из поколения в поколение.
- » **Обучение.** Быстрое обучение на базе организованных источников.
- » **Исследование.** Спонтанное обучение на базе средств массовой информации и общения между собой.

Кроме того факта, что машинное обучение состоит из математических функций, оптимизированных для определенной цели, пределы машинного обучения обуславливают и другие недостатки. Необходимо учесть три важных предела.

- » **Представление.** Представление некоторых проблем с использованием математических функций не всегда просто, особенно для таких комплексных проблем, как имитация работы человеческого мозга. В настоящее время машинное обучение может решать отдельные специфические задачи, подразумевающие ответы на такие простые вопросы, как “Что это такое?”, “Сколько стоит?” и “Что будет дальше?”
- » **Переобучение.** Алгоритму машинного обучения может казаться, что он изучает то, о чем вы просили, но фактически это не так. Их внутренние функции по большей части только запоминают данные, но не учатся на них. *Переобучение* (overfitting) происходит, когда ваш алгоритм учится на ваших данных слишком много и достигает момента создания функций и правил, которых в действительности не существует.
- » **Нехватка эффективного обобщения из-за ограниченных данных.** Алгоритм изучает то, что вы ему даете. Если снабдить алгоритм плохими или недостоверными данными, он поведет себя неожиданным образом.

Что касается представления, отдельный обучаемый алгоритм может узнать много разных вещей, но не каждый алгоритм подходит для определенных задач. Некоторые алгоритмы являются достаточно общими, они могут играть в шахматы, распознавать лица в Facebook и диагностировать рак у пациентов. Алгоритм ограничивает поступающие данные, и ожидаемым результатом этих данных в любом случае будет функция, но функция, специфическая для задач такого вида, для которого предназначен алгоритм.

Тайна машинного обучения — в обобщении. Однако в обобщении кроются проблемы переобучения и смещенных данных. Задача в том, чтобы обобщить функцию вывода так, чтобы эти проблемы не повлияли на данные учебных примеров. Рассмотрим, к примеру, фильтр спама. Скажем, что ваш словарь

содержит 100 000 слов (небольшой словарь). Учебный набор данных, ограниченный 4000 или 5000 словосочетаний, должен создать обобщенную функцию, способную затем найти спам в $2^{100\,000}$ комбинациях, которые функция будет встречать при работе с фактическими данными. В таких условиях алгоритму будет казаться, что он изучил правила языка, но в действительности это не так. Алгоритм может правильно реагировать на ситуации, подобные использованым при обучении, но в совершенно новых ситуациях окажется некомпетентным. Или могут неожиданно проявиться пристрастия из-за вида использованных при обучении данных.

Например, компания Microsoft обучала свой искусственный интеллект, Тау, общаться с людьми через Twitter, а также учиться на их ответах. К сожалению, общение пошло неправильно, поскольку пользователи научили Тау нецензурной речи, поставив под вопрос совершенство любого искусственного интеллекта на базе технологий машинного обучения. (Частично об этой истории можно узнать по адресу <https://www.theverge.com/2016/3/24/11297050/taymicrosoft-chatbot-racist>). Проблема была в том, что алгоритму машинного обучения были плохо поданы данные, без фильтрации (Microsoft не использовала соответствующий анализ данных для их чистки и правильной балансировки), что привело в результате к переобучению. Переобучение привело к выбору неправильного набора функций для общего представления мира способом, который должен был бы избежать нетолерантности, такой как нецензурная речь. Другой обучаемый беседе с людьми искусственный интеллект, заслуженный Mitsuku (<http://www.mitsuku.com/>)², не подвержен таким рискам, как Тау, поскольку его обучение строго контролируется как анализом данных, так и человеком.

Обучение на основе данных

В машинном обучении все вращается вокруг алгоритмов. Алгоритм — это процедура или формула решения задачи. Предметная область включает и вид необходимого алгоритма, но базовая предпосылка всегда одинакова: решить своего рода задачу, такую как вождение автомобиля или игра в домино. В первом случае задача сложна и многогранна, но сводится в конечном счете к одному — доставить пассажира из одного пункта в другой, не разбив автомобиль. Аналогично задача при игре в домино заключается в победе.

Обучение бывает разным в зависимости от алгоритма и его целей. Алгоритмы машинного обучения можно отнести к трем основным группам на основании их задач.

² Переадресация на <https://www.pandorabots.com/mitsuku/>. — *Примеч. ред.*

- » Контролируемое обучение
- » Неконтролируемое обучение
- » Обучение с подкреплением

В следующих разделах более подробно описаны виды алгоритмов машинного обучения.

Контролируемое обучение

Контролируемое обучение (supervised learning) происходит, когда алгоритм учится на основе данных примеров и относящихся к цели ответов, которые могут состоять из числовых значений или строковых меток, таких как классы или тэги, чтобы впоследствии спрогнозировать правильный ответ, когда появятся новые примеры. Контролируемый подход подобен обучению человека под присмотром учителя. Учитель обеспечивает ученику хорошие примеры для запоминания, а ученик усваивает общие правила из этих конкретных примеров.

Необходимо различать *задачи регрессии* (regression problem), целью которых является числовое значение, и *задачи классификации* (classification problem), целью которых является такая качественная переменная, как класс или тэг. Задачей регрессии может быть определение средней цены домов в окрестности города Бостон, а примером задачи классификации может быть поиск различий между цветами ириса на основании размеров их лепестков и чашелистиков. Вот некоторые из примеров контролируемого обучения с важными применениями в искусственном интеллекте, описанные по их вводу и выводу данных, а также реальному применению для решения.

Ввод данных (X)	Вывод данных (Y)	Реальное применение
История покупок клиентов	Список товаров, которые клиенты никогда не покупали	Система рекомендаций
Образы	Список коробок, помеченных неким названием	Обнаружение и распознавание образов
Английский текст в виде вопросов	Английский текст в виде ответов	Chatbot, приложение, способное беседовать
Английский текст	Немецкий текст	Машинный перевод
Аудио	Текстовая транскрипция	Распознавание речи
Образ, полученный сенсором	Поворот, торможение или ускорение	Планирование поведения при автономном вождении

Неконтролируемое обучение

Неконтролируемое обучение (unsupervised learning) происходит, когда алгоритм учится на основе простых примеров без какого-либо ассоциированного ответа, позволяя алгоритму самостоятельно определять шаблоны данных. Этот тип алгоритма стремится реструктурировать данные в нечто еще, такое как новая возможность, способная представить класс, или новая последовательность некоррелированных значений. Полученные данные весьма полезны людям со способностью проникновения в суть смысла первоначальных данных, а также как новые полезные входные данные для алгоритмов контролируемого машинного обучения.

Неконтролируемое обучение напоминает методы, используемые людьми для выявления определенных объектов или событий того же самого класса, с учетом наблюдаемой степени сходства между объектами. Некоторые системы рекомендаций, которые можно найти на сайтах интернет-магазинов, построены на базе обучения этого типа. Алгоритм автоматизированных интернет-магазинов создает свои рекомендации на основании того, что вы покупали в прошлом. Рекомендации создаются на основании оценки того, какую группу клиентов вы напоминаете больше всего, а затем вычисляет ваши вероятные предпочтения на основании предпочтений этой группы.

Обучение с подкреплением

Обучение с подкреплением (reinforcement learning) происходит, когда алгоритму предоставляют примеры без меток, как при неконтролируемом обучении. Но вы можете сопроводить примеры положительной или отрицательной оценкой в зависимости от предлагаемого алгоритмом решения.

Обучение с подкреплением связано с приложениями, алгоритм которых должен принимать решения (поэтому направление весьма перспективно, а не только интересно как неконтролируемое обучение), и эти решения имеют последствия. В мире людей это точно то же, что и обучение методом проб и ошибок. Ошибки помогают учиться, поскольку убыток от них (потеря денег, времени, сожаление, боль и т.д.) наглядно демонстрирует, что определенный образ действия менее вероятен для успеха, чем другие. Интересен пример обучения с подкреплением, когда компьютеры учатся играть в видеоигры.

В данном случае приложение предоставляет алгоритму примеры определенных ситуаций, когда игрок прячется в лабиринте от врага. Приложение позволяет алгоритму узнать результат его действий, и обучение происходит в результате попыток избежать обнаружения опасным преследователем и выжить. Вы можете увидеть, как созданная Google DeepMind программа обучения с подкреплением играет в видеоигры старого Атари по адресу <https://www>.

[youtube.com/watch?v=V1eYniJ0Rnk](https://www.youtube.com/watch?v=V1eYniJ0Rnk). Просматривая видео, обратите внимание на то, что сначала программа неуклюжа и неумела, но постоянно совершенствуется и учится, пока не становится чемпионом. В поучительном видео от TEDx Talks на <https://www.youtube.com/watch?v=mqma6GpM7vM> Райя Хадзелл (Raia Hadsell), главный исследователь группы Deep Learning компании TEDx Talks, описывает сильные и слабые стороны процесса.



Глава 10

Машинное обучение в искусственном интеллекте

В ЭТОЙ ГЛАВЕ...

- » Применение инструментальных средств различных научных школ при обучении на основе данных
- » Чем вероятность полезна искусственному интеллекту
- » Прогнозирование с использованием наивного байесовского классификатора и байесовских сетей
- » Разделение данных на ветви и листья деревьями решений

Обучение было важной частью искусственного интеллекта с самого начала, поскольку искусственный интеллект может подражать интеллекту человека. Достижение такого уровня мимикрии (*mimicry*), который фактически напоминает обучение, займет много времени и потребует множества подходов. Сегодня машинное обучение может демонстрировать квазичеловеческий (*quasi-human*) уровень обучения при решении специфических задач, таких как классификация (*classification*) изображений или обработка звука, но он стремится достичь подобного уровня обучения и для многих других задач.

Машинное обучение автоматизируется не полностью. Вы не можете указать компьютеру прочитать книгу и ожидать, что он что-нибудь поймет. Автоматизация

подразумевает, что компьютеры могут знать, как себя программировать самим, чтобы выполнить некие задачи, а не ожидать программирования от людей. В настоящее время автоматизация требует больших объемов данных, выбранных человеком, а также анализа этих данных и обучения (также под управлением человека). Это как взять ребенка за руку, когда он делает свои первые шаги. Кроме того, у машинного обучения есть пределы, продиктованные обучением на основе данных.

У каждого семейства алгоритмов есть собственные способы решения задач, им и посвящена данная глава. Задача заключается в том, чтобы понять, как искусственный интеллект принимает решения и делает прогнозы. Подобно тому, как за занавесом в книге *Волшебник страны Оз* был обнаружен человек, в этой главе вы раскроете и механизмы, и механику лежащую в основе искусственного интеллекта. Но вы все же поразитесь, узнав о достижениях, которые может обеспечить машинное обучение.

Способы обучения

Подобно тому как существуют разные способы обучения людей, ученые, подходя к проблеме обучения искусственного интеллекта, также следовали разными путями. Каждый верил в свой рецепт подражания интеллекту. До сих пор ни одна из моделей не доказала своего превосходства над другими. Теорема *no free lunch* (бесплатных завтраков не бывает) гласит, что ничего не бывает бесплатно, за каждое преимущество приходится платить в полной мере. Каждое из направлений имеет доказательства эффективности решения специфических задач. Поскольку алгоритмы эквивалентны в абстрактном (см. приведенный ниже раздел “БЕСПЛАТНЫХ ЗАВТРАКОВ НЕ БЫВАЕТ”), ни один из алгоритмов не превосходит другой, если речь не идет о некой конкретной практической задаче. В следующих разделах содержится дополнительная информация по использованию методов обучения.

Пять основных подходов к обучению искусственного интеллекта

Алгоритм — это своего рода контейнер. Он похож на коробку для хранения метода решения задач определенного вида. При обработке алгоритм проводит данные через серию четко определенных состояний. Состояния не обязаны быть определены, тем не менее они определяются. Задача заключается в том, чтобы получить результат, решающий задачу. В некоторых случаях алгоритм получает входные данные, помогающие определить вывод, но основное внимание всегда сосредоточено на выводе.

Так в математическом фольклоре (mathematical folklore) обычно упоминают теорему Дэвида Вулперта (David Wolpert) и Уильяма Макриди (William Macready), утверждающую, что два любых алгоритма оптимизации эквивалентны, если их эффективность остается средней для всех возможных задач. По существу, независимо от используемого алгоритма оптимизации, не будет никакого преимущества от его применения для всех возможных задач. Чтобы получить преимущество, алгоритм следует использовать для тех задач, для которых он подходит лучше всего. Статья Ё-Ши Хо (Yo-Chi Ho) и Дэвида Л. Пепина (David L. Pepyne) по адресу https://www.researchgate.net/publication/3934675_Simple_explanation_of_the_no_free_lunch_theorem_of_optimization содержит доступное, но строгое доказательство этой теоремы. Для большего количества деталей имеет также смысл ознакомиться с обсуждением по адресу <http://www.no-free-lunch.org/>.

Алгоритмы должны выражать переходы между состояниями, используя четкий и формальный язык, понятный компьютеру. При обработке данных и решении задачи алгоритм определяет, детализирует и выполняет функцию. Функция является всегда специфической для вида задачи, решаемой алгоритмом.

Как описано в разделе “Не обмануться в ожиданиях от искусственного интеллекта” главы I, у каждой из этих пяти научных школ есть своя методика и стратегия решения задач, на основании которых созданы уникальные алгоритмы. Объединение этих алгоритмов должно в конечном счете привести к созданию верховного алгоритма, который будет в состоянии решить любую задачу. В следующих разделах представлен краткий обзор пяти основных типов алгоритмов.

Символьное рассуждение

Одна из первых научных школ, символисты, полагала, что знание может быть получено при работе с символами (знаками, имеющими определенный смысл или означающими событие) и выводе правил из них. При формировании достаточно сложных систем правил можно достичь логической дедукции результата, который вы хотели узнать. Таким образом, символисты сформировали свои алгоритмы так, чтобы выводить правила из данных. При символьном рассуждении *дедукция* расширяет область человеческого знания, в то время как *индукция* повышает уровень человеческого знания. Индукция обычно открывает новые области для исследования, а дедукция исследует эти области¹.

¹ Согласно авторам. — *Примеч. ред.*

Взаимосвязи, моделируемые нейронами мозга

Коннекционисты, вероятно, являются самой известной из пяти научных школ. Эта научная школа стремится воспроизвести функции мозга, используя кремний вместо нейронов. По существу, каждый из нейронов (созданный как алгоритм, моделирующий дубликат из реального мира) решает одну малую часть задачи, а используя множество нейронов параллельно, можно решить задачу в целом.

Используя обратное распространение ошибки, можно попытаться определить условия, изменяя *весовые коэффициенты* (*weight*) (насколько конкретный ввод влияет на результат) и *пристрастия* (*bias*) (какие средства выбираются), при которых ошибки удаляются из сетей, построенных подобно сети человеческих нейронов. Цель заключается в том, чтобы продолжать изменять весовые коэффициенты и пристрастия до тех пор, пока фактический вывод не станет соответствовать задаче. В настоящее время искусственный нейрон вырабатывает и передает свое решение следующему нейрону в серии. Решение, созданное только одним нейроном, является лишь частью целого решения. Каждый нейрон передает информацию следующему нейрону в серии, пока группа нейронов не сформирует окончательный вывод. Такой метод доказал свою высочайшую эффективность в решении задач, подобных решаемым человеком, таких как распознавание объектов, письменности, разговорной речи и общения с людьми.

Эволюционные алгоритмы проверки вариантов

Для решения задач эволюционисты полагаются на принципы развития. Другими словами, эта стратегия основана на выживании сильнейшего (удаляются любые решения, которые не соответствуют желаемому результату). Функция выживания определяет жизнеспособность каждой функции в решении задачи. Используя древовидную структуру, метод решения ищет наилучшее решение на основании функции вывода. Для построения функции следующего уровня выбирается победитель текущего уровня развития. Идея в том, что каждый следующий уровень будет ближе к решению задачи, и если решение еще не полное, то просто необходим другой уровень. Для решения проблемы эта научная школа полностью полагается на рекурсии и поддерживающие ее языки. Интересным результатом этой стратегии стали развивающиеся алгоритмы: одно поколение алгоритмов фактически создает следующее поколение.

Байесовский вывод

Группа ученых, байесианцы, полагает, что ключевым аспектом при наблюдении и обучении является неопределенность и что, скорее всего, имеет место непрерывная модификация предыдущих представлений, которые становились все более точными. Это восприятие вынудило байесианцев применять

статистические методы, в частности — следствия теоремы Байеса, позволяющей вычислять вероятности при определенных условиях (например, посмотрите карту с определенным значением *начального числа* псевдослучайной последовательности, нарисованным на доске после трех других карт для того же самого начального числа).

Системы, учащие на аналогиях

Для распознавания шаблонов в данных аналогисты используют многоядерные машины. Распознав шаблон одного набора входных данных и сравнив его с шаблоном известного вывода, можно создать решение задачи. Цель заключается в том, чтобы, используя сходство, определить наилучшее решение задачи. Это тот вид рассуждения, при котором полагают, что если в прошлом данное конкретное решение сработало при неких обстоятельствах, то при подобных же стечениях обстоятельств оно также должно сработать. Одним из наиболее известных результатов этой научной школы являются системы рекомендаций. Например, когда вы покупаете товар на Amazon, система рекомендаций предлагает другой связанный с этим товар, которые вы также могли бы захотеть купить.

Окончательная цель машинного обучения состоит в том, чтобы объединить технологии и стратегии этих пяти научных школ и выработать единый алгоритм (верховный алгоритм), способный что-нибудь изучить. Конечно, до достижения этой цели еще далеко. Но несмотря на это такие ученые, как Педро Домингос (Pedro Domingos) (<http://homes.cs.washington.edu/~pedrod/>), работают над данной задачей уже сейчас.

Три наиболее перспективных подхода к обучению искусственного интеллекта

В следующих разделах этой главы рассматриваются подробности базовых алгоритмов, выбранных байесианцами, символистами и коннекционистами. Эти научные школы представляют настоящую и будущую грани обучения на основе данных, поскольку любой прогресс в искусственном интеллекте, подобном человеческому, ожидается именно от них, по крайней мере пока не произойдет новый прорыв с новыми, более невероятными и мощными алгоритмами обучения. Конечно, сцена машинного обучения значительно шире этих трех алгоритмов, но основное внимание в данной главе уделяется именно этим трем научным школам из-за их роли в текущем положении искусственного интеллекта. Вот краткое описание рассматриваемых подходов.

- » **Наивный байесовский классификатор.** Этот алгоритм может быть куда точнее врача при диагностировании некоторых болезней. Кроме того, один и тот же алгоритм может и обнаруживать спам, и

прогнозировать впечатление от текста. Он также широко используется в Интернете для упрощения обработки больших объемов данных.

- » **Байесовская сеть (форма графа).** Этот граф дает представление о сложности мира в терминах вероятности.
- » **Дерево решений.** Тип алгоритма, представляющий символистическое решение лучше всего. Дерево решений имеет длинную историю и демонстрирует, как искусственный интеллект способен принимать решения, поскольку он напоминает серию вложенных решений, которые можно представить как дерево (отсюда и название).

Следующая глава, “Усиление искусственного интеллекта глубоким обучением”, знакомит с нейронными сетями — образцовым типом алгоритма, предложенного коннекционистами, и реальным механизмом ренессанса искусственного интеллекта. Сначала в главе 11 обсуждается работа нейронной сети, а затем объясняется, что такое глубокое обучение и почему оно настолько эффективно.



ЗАПОМНИ

Во всех этих разделах обсуждаются типы алгоритмов. Далее алгоритмы каждого типа делятся на подкатегории. Например, деревья решений подразделяются на деревья регрессий, деревья классификации, бустинговые деревья (boosted tree), бустинговые объединенные деревья (bootstrap aggregated tree) и лес случайностей (rotation forest). Можно даже выделить подтипы и подразделы. Классификатор леса случайностей — это своего рода бустинговое объединенное дерево, и далее в том же духе. Закончив с уровнями, вы увидите, что количество реальных алгоритмов исчисляется тысячами. Короче говоря, эта книга дает лишь краткий обзор более сложной темы, для рассмотрения которой во всех подробностях может потребоваться много томов. Главное — получить представление о типах алгоритмов и не утонуть в подробностях.

Ожидание следующего прорыва

В 1980-х годах, когда миром искусственного интеллекта правили экспертные системы, большинство ученых и практиков считали машинное обучение второстепенной ветвью искусственного интеллекта, сосредоточенной на обучении тому, как лучше соответствовать простой окружающей обстановке (представленной данными) с использованием оптимизации. Сегодня у машинного обучения верховенство в искусственном интеллекте, оно даже превзошло экспертные системы во многих приложениях и исследовательских разработках, а приложения искусственного интеллекта оно подняло на такую высоту, которую ученые ранее считали невозможной по точности и эффективности. Нейронные сети, решение, предложенное коннекционистами, также добивались прогресса

за последние несколько лет, ставшего возможным благодаря сочетанию увеличения возможностей аппаратных средств, более подходящих данных и усилий таких ученых, как Джеффри Хинтон (Geoffrey Hinton), Ян Лекун (Yann LeCun), Йошуа Бенгио (Yoshua Bengio) и многих других.

Возможности алгоритмов нейронной сети (недавно заклеенные сторонниками глубокого обучения за чрезмерную сложность) ежедневно расширяются. Новости пестрят сообщениями о новых достижениях в областях распознавания звука, образов и видео, языковом переводе и даже чтении по губам. (Хотя глубокому обучению и недостает эффективности HAL9000, оно уже приближается к эффективности человека; см. статью <https://www.theverge.com/2016/11/7/13551210/ai-deep-learning-lipreading-accuracy-oxford>). Усовершенствования — это результат хорошего финансирования большими и малыми компаниями, позволяющего привлечь исследователей и создать лучшее программное обеспечение, такое как TensorFlow от Google (<https://www.tensorflow.org/>) и CNTK (Computational Network Toolkit) от Microsoft (<https://blogs.microsoft.com/ai/microsoft-releases-cntk-its-open-source-deep-learning-toolkit-on-github/>), предоставляющее доступ к технологии и ученым, и практикам.

В ближайшем будущем ожидаются еще более сенсационные новости в области искусственного интеллекта. Конечно, исследователи всегда могут снова наткнуться на стену, как случилось в предыдущие зимы искусственного интеллекта. Никто не может знать, достигнет ли искусственный интеллект человеческого уровня при уже существующей технологии или кто-то откроет верховный алгоритм, поскольку Педро Домингос прогнозирует (см. <https://www.youtube.com/watch?v=qIZ5PXLVZfo>), что он решит все проблемы искусственного интеллекта. Однако машинное обучение — это, конечно, не причуда и не миф; оно существует и будет существовать в улучшенной форме или в форме неких новых алгоритмов.

Поиск правды в вероятностях

На некоторых веб-сайтах утверждается, что статистика и машинное обучение — это две совершенно разные технологии. Например, когда читаешь блог “Statistics vs. Machine Learning, fight!” (<http://brenocon.com/blog/2008/12/statistics-vs-machine-learning-fight/>), складывается впечатление, что эти две технологии не только различны, но и совершенно враждебны одна другой. Хотя статистика демонстрирует скорее теоретический подход к задачам, а машинное обучение просто основано на данных, у статистики и машинного обучения есть много общего. Кроме того, статистика представляет одну из пяти научных школ, делающих машинное обучение возможным.

Статистика часто использует вероятности (как способ выразить неопределенность событий реальности), а следовательно, она присуща и машинному обучению, и искусственному интеллекту (в куда большей степени, чем чистой статистике). Не все задачи похожи на игру в шахматы или Го, позволяющие предпринимать большое, но ограниченное количество действий. Если вы хотите узнать, как будет перемещаться робот в коридоре, переполненном людьми, или создать беспилотный автомобиль, успешно участвующий в дорожном движении, вам стоит учесть, что у некоторых планов (таких, как перемещение из точки А в точку В) не всегда будет единственный результат, что возможно множество результатов, каждый со своей вероятностью. В некотором смысле вероятность поддерживает системы искусственного интеллекта в их рассуждении, принятии решений и выработке того, что кажется наилучшим, самым рациональным выбором, несмотря на неопределенность. Неопределенность может возникать по различным причинам, и искусственный интеллект должен быть осведомлен об уровне неопределенности, чтобы эффективно использовать вероятности.

1. Некоторые ситуации нельзя прогнозировать с уверенностью, поскольку они случайны по своей природе. Подобные ситуации изначально являются стохастическими. Например, при игре в карты вы не можете знать, какие карты окажутся на руках после сдачи.
2. Даже если ситуация не случайна, не факт, что ни один из ее аспектов (неполное наблюдение) не создаст неопределенности по мере развития событий. Например, робот, попавший в коридор с людьми, не может знать намеченный путь каждого человека (он не может прочесть их мысли), но может сделать предположение об этом на основании частичного наблюдения за их поведением. Как и с любым предположением, у робота есть шанс оказаться как правым, так и неправым.
3. Ограниченность записывающих данные аппаратных средств (сенсоров) и приближения при их обработке обуславливают некоторую неопределенность результатов, полученных на их основании. Измерение нередко подвержено ошибкам из-за используемых инструментальных средств и способа измерения. Кроме того, люди зачастую подвержены когнитивным пристрастиям и легко становятся жертвой иллюзий или предвзятости. Точно так же искусственный интеллект ограничен качеством полученных данных. Приближения и ошибки ввода привносят неопределенность в каждый алгоритм.

На что способны вероятности

Вероятность указывает, насколько ожидаемо некое событие, и выражает это как число. Например, подбросив монету, вы не можете знать заранее, как она упадет, орлом или решкой, но вы можете предсказать вероятность обоих

результатов. Вероятность события измеряется в диапазоне от 0 (отсутствие вероятности события) до 1 (полная уверенность в событии). Промежуточные значения, такие как 0,25; 0,5 и 0,75; указывают, что событие произойдет при определенном количестве попыток. Если умножить вероятность на целое число, представляющее количество попыток, вы получите оценку того, как часто событие должно происходить в среднем, если осуществить все попытки. Например, если у вас есть событие, происходящее с вероятностью $p = 0,25$ и вы сделаете 100 попыток, то событие произойдет, вероятно, $0,25 * 100 = 25$ раз.

При выборе случайной карты из колоды масть выпадает с вероятностью $p = 0,25$. Французские игральные карты являются классическим примером для объяснения вероятностей. Колода содержит 52 карты, равномерно разделенных на четыре масти: трефы и пики являются черными, а бубны и черви — красными. Так, если вы хотите определить вероятность выбора туза, то следует предположить, что есть четыре туза разных мастей. Ответ в терминах вероятности таков: $p = 4/52 = 0,077$.

Вероятности находятся в пределах от 0 до 1; никакая вероятность не может превышать эти пределы. Вы определяете вероятности опытным путем из наблюдений. Вы просто помните, как часто происходит некое событие относительно всех остальных интересующих вас событий. Скажем, например, что вы хотите вычислить вероятность мошенничества при банковских операциях или как часто люди болеют некой болезнью в некой стране. Наблюдения позволяют оценить вероятность события, достаточно подсчитать количество искомых событий и поделить на количество всех событий.

Вы можете вычислить количество случаев мошенничества или заболеваний, используя данные из записей (обычно их берут из баз данных), а затем поделить их на общее количество событий или доступных наблюдений. Поэтому вы делите количество мошенничеств на общее количество транзакций за год или подсчитываете количество случаев заболевания за год относительно всего населения некой области. В результате получится число в пределах от 0 до 1, которое вы можете использовать как свою базовую вероятность для определенного события при неких обстоятельствах.

Подсчет всех случаев события не всегда возможен, поэтому необходима выборка. При выборке, сделанной на основании определенных ожиданий вероятности, вы можете наблюдать малую часть большего набора событий или объектов, и все же быть в состоянии предусмотреть правильные вероятности для события, а также точно осуществить количественные или качественные измерения, связанные с набором объектов. Например, если вы хотите отследить продажи автомобилей в США за прошлый месяц, вы не обязаны отслеживать каждую продажу в стране. Используя выборку продаж всего для нескольких автомобильных дилеров в стране, вы можете определить такие количественные

показатели, как средняя цена проданного автомобиля, или такие качественные показатели, как модель автомобиля, продаваемая чаще всех.

Учет предыдущих знаний

Вероятность имеет смысл в терминах времени и пространства, но на нее влияют и некоторые другие условия при измерении. Контекст важен. Оценивая вероятность события, вы можете полагать (иногда ошибочно), что можете применить вычисленную вероятность к каждой возможной ситуации. Для выражения этой веры есть термин, *априорная вероятность* (a priori probability), означающий общую вероятность события.

Например, когда вы бросаете монету, если она настоящая, априорная вероятность орла составляет приблизительно 50 процентов (если вы принимаете также существование крошечной вероятности падения монеты на ребро). Независимо от того, сколько бы раз вы ни бросали монету, перед каждым следующим броском вероятность орла все равно будет примерно 50 процентов. Но в некоторых других ситуациях, если вы изменяете контекст, априорная вероятность окажется неверной, поскольку нечто может вмешаться и изменить ее. В данном случае вы можете выразить эту веру как *апостериорную вероятность* (a posteriori probability), являющуюся априорной вероятностью после вмешательства чего-то, меняющего счет.

Например, априорная вероятность того, что человек является женщиной, составляет примерно 50 процентов. Но вероятность может решительно отличаться, если вы рассматриваете конкретные возрастные диапазоны, поскольку женщины обычно живут дольше и после определенного возраста возрастная группа содержит больше женщин, чем мужчин. Вот другой связанный с полом пример: в настоящее время женщины в среднем превосходят численностью мужчин в главных университетах (см. примеры этого явления по адресам <https://www.theguardian.com/education/datablog/2013/jan/29/how-many-men-and-women-are-studying-at-my-university> и <https://www.ucdavis.edu/news/gender-gap-more-female-studentsmales-attending-universities/>). Поэтому в данных двух контекстах апостериорная вероятность отличается от ожидаемой априорной. На распределение по половому признаку вполне может влиять и характер национальной культуры, он вполне может создать иную апостериорную вероятность. Следующие разделы помогут лучше понять пользу вероятности.

Условная вероятность и наивный байесовский классификатор

Вы можете считать такие связанные с полом случаи, как упомянутые в предыдущем разделе, *условной вероятностью* (conditional probability) и выражать

ее как $p(y|x)$, где вероятность события y выше, если событие x уже произошло. Условная вероятность — это очень мощный инструмент для машинного обучения и искусственного интеллекта. Фактически, если априорная вероятность может измениться неким образом из-за определенных обстоятельств, то, зная эти обстоятельства, можно увеличить возможность правильного предсказания события при наблюдении примеров, а это именно то, для чего предназначается машинное обучение. Как, например, уже упоминалось, случайный человек с общей вероятностью примерно в 50 процентов является мужчиной или женщиной. Но что если добавить такой фактор, как длина волос? Вы можете оценить вероятность наличия длинных волос как 35 процентов у населения в целом; но если рассматривать только женскую часть населения, вероятность повышается до 60 процентов. Если процент настолько высок у женского населения вопреки априорной вероятности, то такой алгоритм машинного обучения, как *наивный байесовский классификатор* (Naïve Bayes), потребует ввода, указывающего длину волос человека.

Фактически алгоритм наивного байесовского классификатора использует в своих интересах факт увеличения шанса правильного прогноза при знании обстоятельств, сопутствующих прогнозу. Все началось с преподобного Томаса Байеса (Reverend Bayes) и его революционной теоремы вероятностей. Как уже упоминалось в книге, в его честь названа одна из научных школ машинного обучения. Для решения задач байесианцы используют различные статистические методы на основании наблюдения вероятностей желательного результата в правильном контексте до и после наблюдения самого результата. На основании этих наблюдений они решают задачу восхода солнца (оценка вероятности, что солнце завтра взойдет) регулярными наблюдениями и непрерывной модификацией их оценки вероятности восхода солнца, повышающейся пропорционально количеству засвидетельствованных длин последовательностей рассветов. Вы можете почитать о рассуждении байесианцев относительно новорожденного ребенка, наблюдающего восход солнца, на сайте *Economist* по адресу <http://www.economist.com/node/382968>.

У аналитиков данных есть большие ожидания в области разработки усовершенствованных алгоритмов на основании вероятности Байеса. Журнал *Technology Review* Массачусеттского технологического института упоминает машинное обучение Байеса как новую технологию, способную изменить мир (<http://www2.technologyreview.com/news/401775/10-emergingtechnologies-that-will-change-the/>). И все же основы теоремы Байеса не столь сложны (хотя они могут быть немного интуитивно непонятны, если вы обычно мыслите, как и большинство людей, только априорными вероятностями, не учитывая имеющийся опыт).

Теорема Байеса

Томас Байес был не только священником в пресвитерианской церкви, но также статистиком и философом, сформулировавшим свою теорему в первой половине XVIII века. При его жизни эта теорема никогда не публиковалась, но впоследствии она сделала революцию в теории вероятности, введя понятие условной вероятности, упомянутой в предыдущем разделе. Благодаря теореме Байеса прогноз вероятности того, что человек является мужчиной или женщиной, становится более простым, если есть доказательство, что у человека длинные волосы. Вот формула, используемая Томасом Байесом:

$$P(B|E) = P(E|B) * P(B) / P(E)$$



ЗАПОМНИ

Преподобный Байес не разрабатывал наивный байесовский классификатор; он только сформулировал теорему. По правде говоря, никакого точного определения этого алгоритма нет. Впервые он появился в книге 1973 года безо всякой ссылки на его создателя и оставался незамеченным более десятилетия, до 1990 года², когда исследователи обратили внимание на то, с какой невероятной точностью он выполняет прогнозы, если предоставить ему достаточно много точных данных. Применение этой формулы к предыдущему примеру в качестве ввода может дать лучшее представление об этой малопонятной в противном случае формуле.

- » **$P(B|E)$.** Вероятность веры (B) при данном наборе доказательств (E) (апостериорная вероятность). Считайте *веру* (belief) альтернативным выражением гипотезы. В данном случае гипотеза — это то, что человек — женщина, а доказательство — длинные волосы. Знание вероятности веры способно помочь прогнозировать пол человека с некой уверенностью.
- » **$P(E|B)$.** Вероятность наличия длинных волос, когда человек — женщина. Этот термин относится к вероятности доказательства в подгруппе, которая сама является условной вероятностью. В данном случае примерно 60 процентов, что преобразуется в значение 0,6 в формуле (априорная вероятность).
- » **$P(B)$.** Общая вероятность того, что человек — женщина, т.е. априорная вероятность веры. В данном случае вероятность составляет 50 процентов, или значение 0,5 (вероятность).
- » **$P(E)$.** Общая вероятность наличия длинных волос. Это другая априорная вероятность, на сей раз связанная с наблюдаемым доказательством. Это 35-процентная вероятность, которая в формуле является значением 0,35 (доказательство).

² См. лучше https://ru.wikipedia.org/wiki/Теорема_Байеса. — *Примеч. ред.*

Если решить приведенную выше задачу, используя ее значения и формулу Байеса, получится результат $0,6 * 0,5 / 0,35 = 0,857$. Это высокий процент от вероятности, позволяющий утверждать, что при данных обстоятельствах человек, вероятно, является женщиной.

Другим наиболее популярным примером, способным открыть некоторым глаза и обычно приводимым в учебниках и научных журналах, является пример положительного медицинского анализа. Он весьма интересен для лучшего понятия того, как априорные и апостериорные вероятности могут действительно существенно изменяться в зависимости от обстоятельств.

Скажем, вы обеспокоены тем, что могли заболеть очень редкой болезнью, встречающейся у 1 процента населения. Вы сдаете анализ, и результат положительный. Медицинские анализы никогда не бывают совершенно точными, и в лаборатории вам говорят, что, когда дело плохо, анализ положителен в 99 процентах случаев, тогда как если вы здоровы, анализ будет отрицательным в 99 процентах случаев. Теперь, уже имея представление, вы сразу полагаете, что дело плохо, с учетом весьма высокого процента положительного результата, когда человек действительно болен (99 процентов). Однако действительность совершенно иная. В данном случае теорема Байеса применима следующим образом:

» $0,99 = P(E|B)$

» $0,01 = P(B)$

» $0,01 * 0,99 + 0,99 * 0,01 = 0,0198 = P(E)$

Вычисление дает $0,01 * 0,99 / 0,0198 = 0,5$, что соответствует только 50-процентной вероятности заболевания. В конце концов, ваши шансы не быть больным куда больше, чем вы ожидали. Возникает вполне резонный вопрос, как это так? Факт в том, что количество людей, видящих положительный результат анализа, таков.

- » **Действительно больные и получающие правильный результат анализа.** Эта группа действительно позитивна, и она составляет 99 процентов больных от 1 процента всего населения.
- » **Здоровые, получившие неправильный результат анализа.** Эта группа составляет 1 процент от 99 процентов людей, получивших позитивный результат, даже при том что они не больны. И снова, это умножение 99 процентов и 1 процента. Эта группа соответствует ошибочно позитивным.

Если взглянуть на проблему с этой точки зрения, становится очевидным, почему. Ограничив контекст людьми, получившими положительный результат

анализа, можно сказать, что есть вероятность оказаться в группе истинно положительных, но с той же вероятностью можно оказаться и в группе ошибочно положительных.

Представление реального мира как графа

Теорема Байеса может помочь рассчитать, с какой вероятностью произойдет некое событие в определенном контексте, на основании общей вероятности самого факта и исследованных доказательств совместно с вероятностью доказательств данного факта. Иногда одна часть доказательства уменьшает сомнения и обеспечивает достаточную уверенность в прогнозе, чтобы нечто гарантировать. Как истинный детектив, чтобы достигнуть уверенности, вы должны собрать в своем расследовании больше доказательств и заставить отдельные части сложиться вместе. Обратите внимание, что факта наличия у человека длинных волос недостаточно, чтобы утверждать, что человек является женщиной или мужчиной. Но дополнительные данные о росте и весе вполне могут повысить уверенность.

Алгоритм наивного байесовского классификатора позволяет упорядочить все собранные доказательства и достичь более обоснованного прогноза с более высокой вероятностью правильности. Полученное доказательство, рассматриваемое отдельно, не может предохранить вас от риска неправильного предсказания, но суммирование всех доказательств позволяет достичь более категорического результата. Следующий пример демонстрирует, как работает наивный байесовский классификатор. Это известная задача, но она демонстрирует виды возможностей, которых можно ожидать от искусственного интеллекта. Набор данных взят из статьи Джона Росса Квинлана (John Ross Quinlan) “Induction of Decision Trees” (<https://dl.acm.org/citation.cfm?id=637969>). Квинлан — программист, участвовавший в разработке другого фундаментального алгоритма машинного обучения, деревьев решений, но его пример хорошо применим для алгоритмов обучения любых видов. Задача требует, чтобы искусственный интеллект предложил наилучшие погодные условия для игры в теннис. Квинлан описал следующий набор возможностей.

- » **Метеоусловия.** Солнечно, пасмурно или дождь.
- » **Температура.** Холодно, умеренно или жарко.
- » **Влажность.** Высокая или нормальная.
- » **Ветер.** Есть или нет.

В следующей таблице содержатся записи базы данных, используемые для примера.

Метеоусловия	Температура	Влажность	Ветер	Можно ли играть
Солнечно	Жарко	Высокая	Нет	Нельзя
Солнечно	Жарко	Высокая	Есть	Нельзя
Пасмурно	Жарко	Высокая	Нет	Можно
Дождь	Умеренно	Высокая	Нет	Можно
Дождь	Холодно	Нормальная	Нет	Можно
Дождь	Холодно	Нормальная	Есть	Нельзя
Пасмурно	Холодно	Нормальная	Есть	Можно
Солнечно	Умеренно	Высокая	Нет	Нельзя
Солнечно	Холодно	Нормальная	Нет	Можно
Дождь	Умеренно	Нормальная	Нет	Можно
Солнечно	Умеренно	Нормальная	Есть	Можно
Пасмурно	Умеренно	Высокая	Есть	Можно
Пасмурно	Жарко	Нормальная	Нет	Можно
Дождь	Умеренно	Высокая	Есть	Нельзя

Возможность играть в теннис зависит от четырех аргументов, представленных на Рис. 10.1.



Рис. 10.1. Модель наивного байесовского классификатора позволяет отследить доказательства до правильного результата

Результатом этого примера обучения искусственного интеллекта является решение о том, играть ли в теннис при данных погодных условиях (доказательство). Только метеоусловий (солнечно, пасмурно или дождливо) не будет достаточно, поскольку температура и влажность также могут быть слишком

высокими или ветер может быть слишком сильным. Эти аргументы представляют реальные условия, у которых есть несколько причин или причины которых взаимосвязаны. Алгоритм наивного байесовского классификатора прекрасен, когда справедливо предполагается наличие нескольких причин.

Алгоритм вычисляет балл на основании вероятности принятия конкретного решения, умноженной на вероятность доказательства, связанного с этим решением. Например, чтобы принять решение, играть ли в теннис, когда погода солнечная, но ветер силен, алгоритм вычисляет балл для положительного ответа, умножая общую вероятность игры (9 сыгранных игр из 14 начатых) на вероятность того, что день был солнечным (2 из 9 сыгранных игр) и наличия ветра при игре в теннис (3 из 9 сыгранных игр). Те же самые правила применимы и в негативном случае (у которого другие вероятности для отказа от игры при данных конкретных условиях).

вероятность игры: $9/14 * 2/9 * 3/9 = 0,05$

вероятность отказа от игры: $5/14 * 3/5 * 3/5 = 0,13$

Поскольку балл для вероятности выше, алгоритм решает, что при таких условиях лучше не играть. Эта вероятность вычисляется суммированием двух баллов и делением обоих баллов на их сумму:

вероятность игры: $0,05 / (0,05 + 0,13) = 0,278$

вероятность отказа от игры: $0,13 / (0,05 + 0,13) = 0,722$

Используя *байесовскую сеть*, состоящую из графов, представляющих взаимосвязь событий, наивный байесовский классификатор можно дополнять и далее, чтобы представлять куда более сложные отношения, чем набор факторов, которые и так намекают на вероятный результат. У байесовских графов есть узлы, представляющие события, и дуги, представляющие связь данного события с другими, а также таблицы условных вероятностей, описывающих эти связи в терминах вероятности. На Рис. 10.2 представлен известный пример байесовской сети, взятый из академической статьи “Local computations with probabilities on graphical structures and their application to expert systems” Стефана Л. Лауритцена (Steffen L. Lauritzen) и Дэвида Спигерелхолтера (David J. Spiegelhalter) в журнале *Journal of the Royal Statistical Society* за 1988 год (см. <https://www.jstor.org/stable/2345762>).

Представлена сеть *Азия* (Asia). Она демонстрирует возможные состояния пациента и их причины. Например, если у пациента одышка, это может быть симптомом туберкулеза, рака легких или бронхита. Знание, курит ли пациент, был ли в Азии и имеет ли аномалии на рентгеновских снимках (придающее уверенность определенным частям доказательства априорно, на языке Байеса), позволяет вычислить реальные (апостериорные) вероятности наличия в графе любой из патологий.

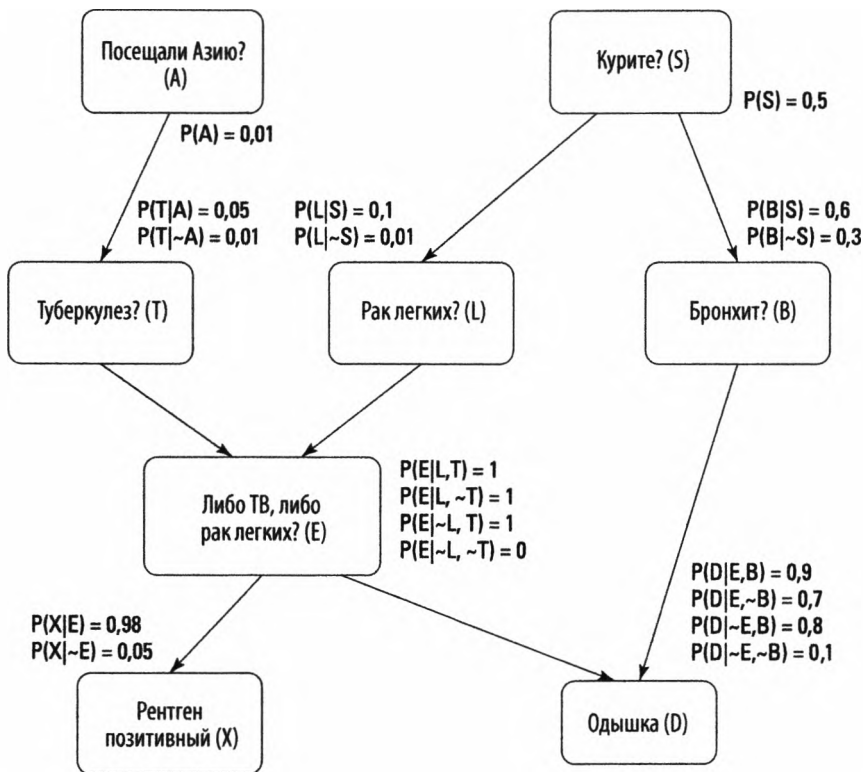


Рис. 10.2. Байесовская сеть способна предоставить медицинское решение

В основе байесовских сетей, несмотря на их интуитивную понятность, лежит сложный математический механизм, они куда мощнее простого алгоритма наивного байесовского классификатора, поскольку подражают реальности как последовательности причин и следствий на основании вероятностей. Байесовские сети настолько эффективны, что вы можете использовать их для моделирования любой ситуации. Они применяются для множества целей, таких как медицинская диагностика, обработка неполных данных, поступающих от нескольких сенсоров, экономическое моделирование, а также контроль таких сложных систем, как автомобиль. Например, поскольку вождение в напряженной дорожной обстановке подразумевает сложные ситуации с участием многих транспортных средств, консорциум Analysis of Massive Data Streams (AMIDST) в сотрудничестве с автомобилестроительной компанией Daimler разработал байесовскую сеть, способную распознавать маневры транспортных средств и повышать безопасность.

Рост деревьев, способных классифицировать

Дерево решений (decision tree) — это другой тип ключевого алгоритма машинного обучения, влияющий на реализацию искусственного интеллекта и обучения. Алгоритмы дерева решений не новы, их история действительно длинна. Первый такой алгоритм относится к 1970-м годам (со многими последующими вариантами). Когда рассматриваешь первые эксперименты и оригинальное исследование по использованию деревьев решений даже в прошлом, они весьма впечатляют. В качестве базового алгоритма символистов деревья решений давно популярны, поскольку интуитивно понятны. Он просто преобразует вывод в правила, а потому сделать вывод понятным людям довольно легко. Кроме того, деревья решений чрезвычайно удобны. Все эти характеристики делают их эффективными и легко реализуемыми при создании моделей, требующих сложных входных преобразований матрицы данных или чрезвычайно точной настройки гиперпараметров.



ЗАПОМНИ

Символизм (symbolism) — это подход к созданию искусственного интеллекта на базе логических операторов и широкого применения дедукции. *Дедукция* (deduction) развивает знание от того, что уже известно, а *индукция* (induction) формулирует общие правила, начинающиеся с доказательства.

Предсказание результатов при разделении данных

Если есть группа показателей и вы хотите описать ее, используя одно число, то вы используете *среднее арифметическое значение* (сумму всех показателей, деленную на их количество). Точно так же, если у вас есть группа классов или качеств (например, набор данных, содержащий записи о множестве пород собак или типов товаров), чтобы представить все классы группы, вы можете использовать класс, встречающийся в группе наиболее часто, т.е. моду. *Мода* (mode) — это такой же статистический показатель, как и среднее значение, но он содержит конкретное значение (показатель или класс), встречающееся чаще всего. И среднее значение, и мода стремятся выразить число или класс, предоставляющий наиболее вероятный следующий элемент группы, поскольку его предположение дает наименьшее количество ошибок. В некотором смысле это предсказание на основании изучения существующих данных. Деревья решений манипулируют средними значениями и модами как прогнозами при разделении набора данных на меньшие наборы, средние значения или моды которых являются наилучшими возможными прогнозами для текущей проблемы.



СОВЕТ

Разделение проблемы на меньшие задачи для облегчения поиска решения является весьма популярной стратегией многих алгоритмов наподобие *разделяй и властвуй*. Это как с враждебными странами: если удастся натравить их одна на другую, можно завоевать обе с минимумом усилий.

Используя как отправную точку выборку наблюдений, алгоритм выводит правила, создавшие выходные классы (или числовые значения при решении задачи регрессией) за счет разделения исходной матрицы на все меньшие и меньшие части, пока процесс не выработает правило для остановки. Такое воссоздание от частных к общим правилам типично для человеческой обратной дедукции, проистекающей из логики и философии.



ЗАПОМНИ

В контексте машинного обучения такое инверсное рассуждение достигается за счет поиска среди всех возможных путей разделения обучения на выборки и решения, жадным способом, использовать ли разделение, максимизирующее статистические показатели в результирующем разделении. Алгоритм считается *жадным* (greedy), если он всегда решает максимизировать результат на текущем этапе процесса оптимизации, независимо от того, что может случиться на следующих этапах. В результате жадный алгоритм не может достичь глобальной оптимизации.

Разделение осуществляется для приведения в действие простого принципа: каждое разделение начальных данных должно упростить предсказание результата, характеризующееся другим, более благоприятным распределением классов (или значений), чем исходная выборка. Алгоритм создает разделение, разделяя данные. Разделение данных он определяет первой оценкой возможностей. Затем он вычисляет значения и возможности, способные обеспечить максимальное улучшение некоего статистического показателя, т.е. показателя, играющего роль функции стоимости в дереве решений.

Решение о разделении в дереве решений принимается на основании множества статистических показателей. При этом разделение всегда должно изменять к лучшему исходную выборку или другое возможное разделение, чтобы сделать прогноз более достоверным. Среди показателей наиболее популярными являются индекс Джини (gini impurity), прирост информации (information gain) и снижение дисперсии (variance reduction) (для регрессионных задач). Эти показатели работают практически одинаково, поэтому данная глава сосредоточивается на приросте информации, поскольку это наиболее интуитивно понятный показатель, вполне позволяющий продемонстрировать, как дерево решений способно повысить достоверность прогноза (или снизить риск)

самым простым способом для некоего разделения. Росс Квинлан создал алгоритм дерева решений на основании прироста информации (ID3) в 1970-х годах, но он все еще весьма популярен благодаря своей недавно улучшенной версии C4.5. Прирост информации полагается на формулу для информативной энтропии (найдена американским математиком и инженером Клодом Шенноном (Claude Shannon), известным как отец теории информации), обобщенная формулировка которой описывает ожидаемое значение от информации, содержащейся в сообщении:

$$\text{Энтропия Шеннона } E = - \sum (p(i) \times \log_2(p(i)))$$

В формуле все классы учтены по одному, вы суммируете результат умножения каждого из них. При умножении каждый класс должен дать $p(i)$ — вероятность для данного класса (выраженная в диапазоне от 0 до 1) и $\log_2$ — логарифм по основанию 2. Начнем со случая, когда необходимо классифицировать два класса, имеющих одинаковую вероятность (распределение 50/50). Максимально возможная энтропия составит Энтропия = $-0,5 \times \log_2(0,5) - 0,5 \times \log_2(0,5) = 1,0$. Но когда алгоритм дерева решений обнаруживает возможность разделить набор данных на два, где распределение двух классов составит 40/60, средняя информативная энтропия уменьшается: Энтропия = $-0,4 \times \log_2(0,4) - 0,6 \times \log_2(0,6) = 0,97$

Обратите внимание на сумму энтропии для всех классов. При распределении 40/60 сумма оказывается меньше теоретического максимума 1 (снижение энтропии). Считайте энтропию мерой беспорядочности данных: чем меньше беспорядка, тем больше порядка и тем проще предположить правильный класс. После первого разделения алгоритм пытается разделить полученные части далее, используя ту же самую логику снижения энтропии. Это последовательно разделяет все получаемые фрагменты данных до тех пор, пока дальнейшее деление не окажется невозможным, поскольку подвыборка сведется к одиночному примеру или будет удовлетворено правило остановки.

Правило остановки (stopping rule) ограничивает развитие дерева. Эти правила работают с учетом трех аспектов разделения: размер исходного раздела, размер результирующего раздела и прирост информации, достигнутый разделением. Правила остановки важны, поскольку алгоритмы дерева решений аппроксимируют большое количество функций; но искажения и ошибки в данных могут легко повлиять на этот алгоритм. Следовательно, в зависимости от выборки неустойчивость и дисперсия результирующих оценок влияют на прогнозы дерева решений.

Принятие решений на основании деревьев

В качестве примера использования дерева решений в данном разделе используется тот же самый набор данных Росса Квинлана, который обсуждался ранее, в разделе “Представление реального мира как графа”. Используя этот набор данных, можно продемонстрировать и описать алгоритм ID3 — специальный вид дерева решений, опубликованный в статье “Induction of Decision Trees”, упоминавшейся ранее в этой главе. Набор данных очень прост, он состоит всего из 14 наблюдений погодных условий, результаты которых указывают, имеет ли смысл играть в теннис.

Пример содержит четыре возможности: метеоусловия, температура, влажность и ветер, выражаемые с использованием качественных классов вместо показателей (температуру, влажность и силу ветра можно выразить и в цифровой форме), чтобы в более понятной форме объяснить влияние погоды на результат. После обработки этих возможностей алгоритмом набор данных можно представить в виде древовидной структуры, показанной на Рис. 10.3. Рисунок демонстрирует, что вы можете просмотреть и прочесть набор правил разделения набора данных для создания частей, прогнозы для которых проще при поиске класса, встречающегося чаще (в данном случае результатом является игра в теннис).

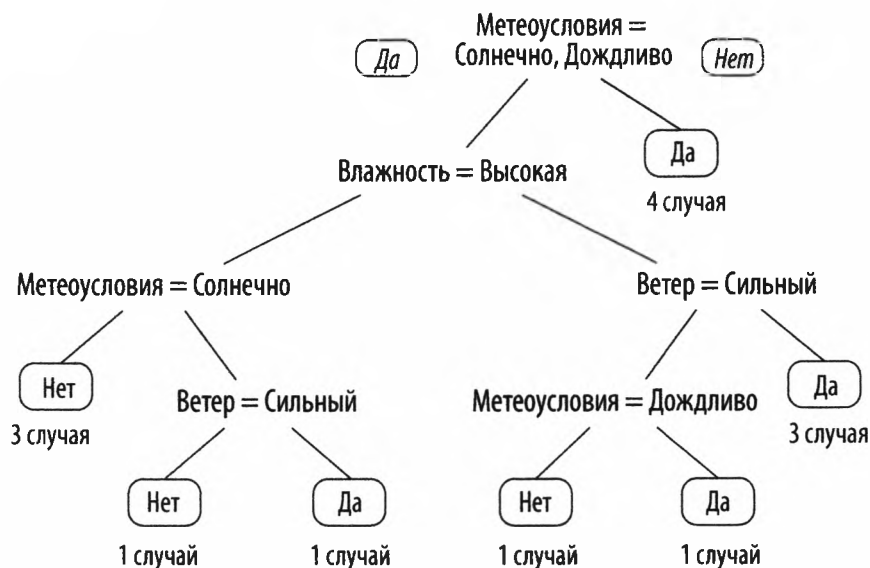


Рис. 10.3. Визуализация дерева решений, созданного для данных примера игры в теннис

Чтение узлов дерева начинается с верховного узла, соответствующего исходным учебным данным; затем читайте правила. Обратите внимание, что у каждого узла есть два последствия: левая ветвь означает, что правило выше истинно (надпись “Да” в прямоугольнике), а правая означает, что оно ложно (надпись “Нет” в прямоугольнике).

Справа от первого правила находится важное конечное правило (конечный лист), объявляющее о положительном результате (Да) который вы можете прочитать как `play tennis=True`. Согласно этому узлу, когда погода солнечная (Солнечно) или дождливая (Дождливо), играть можно. (Числа ниже конечного листа отмечают четыре случая, подтверждающих это правило и нуль отказов.) Обратите внимание, что правило можно было бы понять лучше, если бы вывод просто гласил, что при пасмурной погоде игра возможна. Зачастую правила дерева решений не являются непосредственно пригодными для использования, их следует предварительно интерпретировать. Однако они однозначно понятны (и намного лучше, чем коэффициент вектора значений).

Слева дерево продолжается другими правилами, связанными с влажностью. Далее слева, когда влажность высока и погода солнечна, большинство конечных листьев содержит отказ, кроме тех случаев, когда ветер не силен. Когда вы исследуете ветви справа, вы видите, что дерево прогнозирует возможность игры всегда, когда ветер не сильный или когда ветер сильный, но не идет дождь.

Отсечение переросших деревьев

Хотя набор данных для игры в теннис из предыдущего раздела иллюстрирует подробности дерева решений, у него практически нет проблем, поскольку он предлагает набор детерминированных действий (т.е. у него нет никаких противоречивых инструкций). При обучении на реальных данных таких однозначных правил обычно не бывает, поэтому имеют место неоднозначности и вероятности желаемого результата.

У деревьев решений бывает и больше вариаций, чем пристрастия при оценках. Чтобы пример меньше зависел от данных, установим, что минимальное разделение должно задействовать по крайней мере пять примеров; это также сократит дерево. Отсечение применяется, когда дерево становится слишком большим.

Начиная с листьев, пример сокращает ветви дерева, обеспечивая небольшое преимущество за счет сокращения прироста информации. Вначале дереву разрешают разворачивать ветви при небольших преимуществах, поскольку они могут открыть более интересные ветви и листья. Проход от листьев к

корню и сохранение только тех ветвей, которые имеют значение для прогноза, ограничивает дисперсию модели, делая результирующие правила более однозначными.



СОВЕТ

Отсечение в дереве решений напоминает коллективное обсуждение. Сначала код создает все возможные ветвления дерева (как с идеями при мозговом штурме). Затем, когда обсуждение заканчивается, остается только то, что действительно может работать.



Глава 11

Усиление искусственного интеллекта глубоким обучением

В ЭТОЙ ГЛАВЕ...

- » Сначала был ограниченный перцептрон
- » Получение стандартных блоков нейронной сети и обратное распространение ошибки
- » Восприятие и распознавание объектов в образах с использованием сверточности
- » Последовательности и их получение с использованием RNN
- » Признаки творчества со стороны искусственного интеллекта благодаря GAN

Газеты, бизнес-журналы, социальные сети и не технические веб-сайты — все говорят одно и то же: искусственный интеллект — интересный материал, и он изменит мир благодаря глубокому обучению. Искусственный интеллект — это намного более широкое поле, чем машинное обучение, а глубокое обучение — лишь небольшая часть машинного обучения.

Важно распознавать обман, которым обычно соблазняют инвесторов, и представлять себе реальные возможности этой технологии, в чем, собственно, и заключается задача этой главы. В статье <https://blogs.nvidia.com/blog/2016/07/29/whats-difference-artificial-intelligence-machine-learning-deep-learning-ai/> сравниваются роли трех методов манипулирования данными (искусственный интеллект, машинное обучение и глубокое обучение), рассматриваемыми в этой главе.

Глава поможет понять, что такое глубокое обучение с практической и технической точек зрения, чего оно может достичь в ближайшем времени благодаря его возможностям и ограничениям. Глава начинается с истории и основ нейронных сетей. Затем она знакомит с ультрасовременными результатами сверточных нейронных сетей, рекуррентных нейронных сетей (обе — контролируемого обучения) и генеративно-состязательных сетей (своего рода неконтролируемое обучение).

Формирование нейронных сетей, подобных человеческому мозгу

В следующих разделах вы познакомитесь с семейством алгоритмов обучения, черпающих вдохновение из работы человеческого мозга. Речь идет о нейронных сетях — базовом алгоритме научной школы коннекционистов, наилучшим образом подражающем нейронам внутри человеческого мозга, но в меньшем масштабе.



ЗАПОМНИ

Коннекционизм (connectionism) — способ машинного обучения на базе неврологии, а также хороший пример биологически подобных сетей.

Знакомство с нейронами

Мозг человека содержит миллионы нейронов, представляющих собой специализированные клетки, способные получать, обрабатывать и передавать электрические и химические сигналы. Каждый нейрон обладает ядром с нитями, осуществляющими обмен данными: *дендриты* (dendrite) получают сигналы от других нейронов, а единственная нить вывода, *аксон* (axon), завершается синапсами, предназначенными для внешних коммуникаций. Нейроны соединяются и передают информацию между собой, используя химические вещества, но в самом нейроне информация обрабатывается электрически. Больше о нейронной структуре можно прочитать по адресу <https://www.dummies.com/>

В главе 10 обсуждаются байесовские сети и демонстрируется пример того, как подобные сети способны давать врачу советы при диагностике. Для этого байесовская сеть требует хорошо подготовленных заранее данных о вероятности. Глубокое обучение позволяет перекинуть мост между возможностью алгоритмов вырабатывать наилучшие возможные решения, используя все необходимые данные, и теми данными, которые фактически доступны, а они редко бывают в формате, понятном алгоритмам машинного обучения. Фотографии, образы, звуковые записи, данные из веба (особенно из социальных сетей) и записи компаний — все они требуют, чтобы анализ данных предварительно преобразовал их в подходящий формат.

Будущий алгоритм глубокого обучения вполне может существенно помочь врачам, применяя обширные медицинские знания (из всех доступных источников, включая книги, официальные документы и последние медицинские исследования) к информации о пациенте. Информация о пациенте, в свою очередь, может состоять из предыдущих диагнозов, выписанных ранее лекарств и даже из публикаций в социальных сетях (чтобы врач не выпрашивал, был ли пациент, например, в Азии, искусственный интеллект просмотрит его фотографии в Instagram или Facebook). Это может показаться фантастикой, но создание такой системы почти возможно уже сегодня. Например, искусственный интеллект глубокого обучения уже способен диагностировать пневмонию по рентгеновскому снимку даже лучше практикующего рентгенолога благодаря стэнфордской группе машинного обучения Stanford Machine Learning Group (<https://stanfordmlgroup.github.io/projects/chexnet/>).

Глубокое обучение уже присутствует во многих приложениях. Его можно найти в социальных сетях, при автоматической классификации образов и содержимого; в поисковых механизмах, при получении запросов; в сетевой рекламе, при выявлении возможных потребителей; в мобильных телефонах и цифровых помощниках, для распознавания речи или перевода; в беспилотных автомобилях, для наблюдения окружающей обстановки; а также в игре AlphaGo против чемпиона. Менее широко известны такие области применения глубокого обучения, как робототехника и прогнозирование землетрясений. Вы могли бы также найти полезными такие приложения, как TinEye (<https://tineye.com/>): вы предоставляете TinEye образ, а она сама находит его в Интернете.

[education/science/biology/whats-the-basic-structure-of-nerves/](https://www.khanacademy/science/biology/whats-the-basic-structure-of-nerves/) или в книге *Neuroscience For Dummies* Фрэнка Амтора (Frank Amthor).

Обратное проектирование обработки сигналов мозгом помогает коннекционистам создавать нейронные сети по аналогии с биологическими, и для их

компонентов используются такие неврологические термины, как “нейроны”, “активизация” и “соединения”, равно как и для математических операций. Если проверить математические формулировки, то компоненты нейронных сетей напоминают не более чем последовательности суммирований и умножений. Но все же эти алгоритмы очень эффективны при решении сложных задач, таких как распознавание образов и звуков, а также машинный перевод. Используя специализированные аппаратные средства, они способны вычислять результат очень быстро.

Сначала был чудесный перцептрон

Базовый алгоритм нейронной сети — это *нейрон* (neuron), он же *модуль* (unit). Множество нейронов, упорядоченных во взаимосвязанной структуре, составляют нейронную сеть, в которой каждый нейрон связан с входами и выходами других нейронов. Таким образом, нейрон может получать исходные данные от одних нейронов и передавать результаты другим в зависимости от своего расположения в нейронной сети.

Несколько десятилетий назад Фрэнк Розенблатт (Frank Rosenblatt) из Корнеллской лаборатории авиационной техники создал первый экземпляр нейрона этого вида — *перцептрон* (perceptron). Он разработал его в 1957 году при финансировании научно-исследовательской лаборатории военно-морского флота США (Naval Research Laboratory — NRL). Розенблатт был психологом, а также пионером в области искусственного интеллекта. Будучи профессионалом в когнитивистике, он предложил создать компьютер, способный учиться методом проб и ошибок, как человек.

Перцептрон был просто интеллектуальным способом выявления черты, разделяющей простое пространство, образованное двумя координатами исходных данных (в данном случае — размер и уровень приручения животного), как показано на рис. 11.1, чтобы различить два класса (собаки и кошки в данном примере). Формулировка перцептрона дает черту в декартовом пространстве, в котором сущности причисляются к группам более или менее отчетливо. Этот подход подобен подходу, используемому наивным байесовским классификатором, описанным в главе 10, который для классификации суммирует условные вероятности, умноженные на общую.

Перцептрон не оправдал в полной мере ожиданий ни его создателя, ни инвесторов. Он продемонстрировал весьма ограниченные способности даже в своей специализации — распознавании изображений. Общее разочарование запустило первую зиму искусственного интеллекта и стало причиной отказа от коннекционизма до 1980-х годов. Несмотря на потерю финансирования некоторые исследования все же продолжались (больше о происходившем в те времена можно узнать из статьи доктора Нильса Ж. Нильссона (Dr. Nils J. Nilsson),

ныне уволенного, но прежде профессора по искусственному интеллекту в Стэнфорде: <https://www.singularityweblog.com/ai-is-so-hot-weve-forgotten-all-about-the-ai-winter/>).

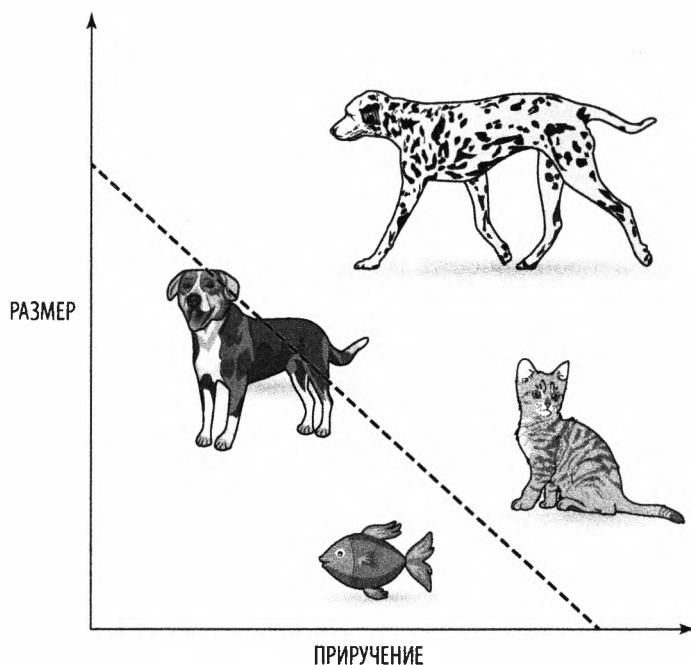


Рис. 11.1. Пример перцептрона в простых и сложных задачах классификации

Впоследствии эксперты пытались создать более усовершенствованный перцептрон, и им это удалось. Нейроны в нейронной сети — это дальнейшее развитие перцептрона: их много, они соединены между собой и подражают человеческим нейронам, когда активизируются определенным стимулом. Наблюдая функции человеческого мозга, ученые обратили внимание на то, что нейроны получают сигналы, но не всегда издают собственный сигнал. Подача сигнала зависит от объема полученного сигнала. Когда нейрон получает достаточно много стимулов, он производит ответ, а в противном случае молчит. Подобным образом действуют и алгоритмические нейроны: после получения данных они их накапливают и используют функцию активизации, чтобы вычислить результат. Если полученный ввод достигает определенного порогового значения, нейрон преобразует их и передает, а в противном случае — просто бездействует.



Для выработки результата нейронные сети используют специальные функции — *функции активизации* (activation function). Это ключевой компонент нейронной сети, поскольку они позволяют сети решать сложные задачи. Они похожи на двери, позволяя передавать сигнал

или не передавать. Они не просто позволяют передачу сигнала, они преобразуют его полезным способом. Глубокое обучение, например, невозможно без эффективных функций активизации, таких как *блок линейной ректификации* (Rectified Linear Unit — ReLU), а следовательно, функции активизации — это важный аспект истории.

Подражание учащемуся мозгу

В нейронной сети сначала следует рассмотреть архитектуру — расположение компонентов нейронной сети. В следующих разделах обсуждаются архитектурные соображения нейронной сети.

Простые нейронные сети

В отличие от других алгоритмов, обладающих фиксированным конвейером получения и обработки данных, нейронные сети требуют решения об информационных потоках при фиксированном количестве модулей (нейронов) и их распределении по уровням, чтобы получилась *архитектура нейронной сети* (neural network architecture), подобная представленной на рис. 11.2.

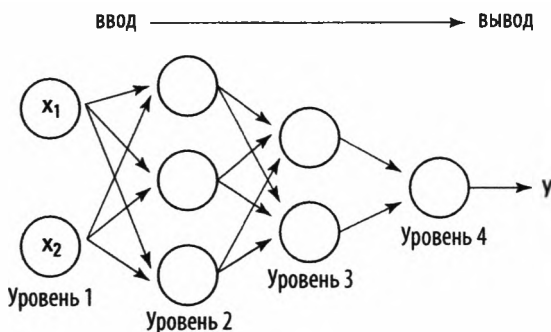


Рис. 11.2. Архитектура нейронной сети от ввода до вывода

На рисунке представлена простая архитектура нейронной сети. Обратите внимание, что уровни фильтруют и обрабатывают информацию прогрессивным способом. Это *ввод с прямой подачей* (feed-forward input) поскольку данные в сеть подаются с одного направления. Модули одного уровня соединяются исключительно с модулями следующего уровня (информация течет слева направо). Нет никаких соединений между модулями на одном и том же уровне или на уровне дальше следующего. Кроме того, информация продвигается только слева направо. Обработанные данные никогда не возвращаются к нейронам предыдущего уровня.

Нейронная сеть похожа на многослойный фильтр для воды: вы льете воду сверху, она фильтруется и вытекает вниз. Вода не может подняться вверх; она движется только вниз и никогда в стороны. Точно так же нейронные сети вынуждают данные течь через сеть и объединяться между собой так, как диктует архитектура сети. При использовании соответствующей архитектуры смешения данных нейронная сеть создает новые составные элементы на каждом уровне и позволяет достичь наилучших прогнозов. К сожалению, нет никакого способа определить наилучшую архитектуру иначе как опытным путем, опробуя различные решения и проверяя, помогают ли выходные данные лучше прогнозировать ваши цели после прохода сквозь сеть.



СОВЕТ

Иногда концепции становятся понятнее после непосредственной проверки на практике. Компания Google предоставляет игровую модель нейронной сети (<http://playground.tensorflow.org>), позволяющую на практике проверить работу нейронной сети интуитивно понятным способом, добавляя или удаляя уровни и изменяя виды активизации.

Раскроем секреты коэффициентов

Нейронные сети имеют несколько уровней, каждый с собственным коэффициентом. *Коэффициенты* (weight) определяют силу связей между нейронами в сети. Когда коэффициент связи между двумя уровнями низок, это значит, что сеть формирует дамп значений, передаваемых между ними, а также что следование этим маршрутом, вероятно, не будет влиять на окончательный прогноз. Аналогично большое положительное или отрицательное значение коэффициента существенно влияет на получаемое следующим уровнем значение, а значит, оно существенно повлияет на прогноз. Этот подход аналогичен используемому клетками мозга, которые работают не самостоятельно, а совместно с другими клетками. По мере накопления опыта человеком связи между его нейронами слабеют или усиливаются, активизируя или деактивируя определенные фрагменты сети клеток мозга, меняя способ обработки информации или предпринимаемое действие (например, реакция на опасность, если обработка поступивших информационных сигналов определила ситуацию как опасную).

На каждом последующем уровне модулей нейронной сети значения, полученные от предыдущих средств, обрабатываются, как ленточный конвейер. По мере передачи по сети данные достигают каждого модуля как некое суммарное значение, созданное из предыдущих значений на предыдущем уровне с учетом коэффициента связи текущего уровня. Когда полученные от других нейронов данные превышают определенное пороговое значение, функция активизации увеличивает значение, хранимое в модуле; в противном случае она

гасит сигнал, снижая его. После обработки функцией активизации результат готов для передачи по связи на следующий уровень. Эти этапы повторяются на каждом уровне, пока значение достигнет конца, это и будет результат.

Коэффициенты связей позволяют смешивать и компоновать входные данные новым способом, создавая новые средства за счет смешения обработанных входных данных творческим способом благодаря функции активизации и коэффициентам. Кроме того, активизация осуществляет обработку нелинейно, обеспечивая рекомбинацию входных данных, получаемых связью. Оба эти компонента нейронной сети позволяют алгоритму изучить сложные целевые функции, представляющие отношения между входными средствами и результатом.

Роль обратного распространения ошибки

В процессе обучения в человеческом мозге формируются и модифицируются синапсы между нейронами на основании стимулов, получаемых при накоплении эмпирического опыта. Нейронные сети реплицируют этот процесс как математическую формулировку — *обратное распространение ошибки* (backpropagation). Вот как эта архитектура взаимосвязанных вычислительных модулей может решать задачи: модули получают пример и, если не делают правильного предположения, отслеживают проблему в обратном порядке по системе существующих коэффициентов, используя обратное распространение ошибки, а затем устраняют ее, изменяя некоторые значения. Этот процесс может потребовать множества итераций, прежде чем нейронная сеть обучится. Итерации в нейронной сети называют *эпохами*, и это хорошее название, поскольку нейронной сети могут понадобиться дни или недели обучения для решения сложных задач.



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Математика обратного распространения ошибки довольно сложна, она требует знания таких концепций, как производные. Более подробная информация по этой теме приведена в книге *Machine Learning For Dummies* Джона Пола Мюллера и Луки Массарона. Концептуально обратное распространение ошибок достаточно интуитивно понятно, поскольку оно напоминает то, что делают люди, когда решают задачи методом проб и ошибок, последовательно приближаясь к решению.

С момента появления алгоритма обратного распространения в 1970-х годах разработчики многократно совершенствовали его и в настоящее время обсуждают, не стоит ли его переделать. (С мнением Джеффри Хинтона (Geoffrey Hinton), одного из соавторов метода, можно ознакомиться по адресу <https://>

Обратное распространение ошибки лежит в основе нынешнего Ренессанса искусственного интеллекта. В прошлом, каждое усовершенствование процесса обучения нейронной сети вело к новым приложениям и возобновлению интереса к методу. Кроме того, текущая революция глубокого обучения, приводящая к возобновлению работ над нейронными сетями (остановленных в начале 1990-х годов), следовала из ключевых достижений в способе, которым нейронные сети учатся на своих ошибках.

Знакомство с глубоким обучением

После обратного распространения ошибок следующее усовершенствование нейронных сетей привело к глубокому обучению. Исследования продолжались, несмотря на зиму искусственного интеллекта, и нейронные сети преодолели технические проблемы, такие как исчезающий градиент, ограничивающие размерность нейронных сетей. Для решения определенных задач разработчики нуждались в больших нейронных сетях, настолько больших, что в 1980-х годах они были невообразимы. Кроме того, исследователи использовали в своих интересах последние достижения в разработке процессоров CPU и GPU (графические процессоры больше известны своим применением в играх).



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Исчезающий градиент (vanishing gradient) — это когда вы пытаетесь передать сигнал через нейронную сеть, но он быстро затухает почти до нуля и не может больше проходить через функции активизации. Это происходит потому, что в нейронной сети происходит поэтапное умножение. Каждое умножение на значение ниже нуля быстро уменьшает значение сигнала, а функции активизации нуждаются в достаточно больших значениях, чтобы позволить дальнейшую передачу сигнала. Чем дальше уровень нейрона от ввода, тем выше вероятность, что он не получит дополнений, поскольку сигнал слишком слаб и функции активизации не пропускают его. В результате сеть прекращает обучаться вообще или учится, но невероятно медленно.

Новое решение позволило избежать проблемы исчезающего градиента и многих других технических проблем, позволив создавать большие *глубокие сети* (deep network) в отличие от более простых *поверхностных сетей* (shallow network) прошлого. Глубокие сети появились благодаря исследованиям ученых из Университета Торонто в Канаде, таких как Джеффри Хинтон (<https://>

www.utoronto.ca/news/artificialintelligence-u-t), настаивавших на продолжении работ над нейронными сетями, даже когда казалось, что они — уже прошлое машинного обучения.

Процессор GPU — это мощнейший блок для матричных и векторных вычислений, необходимых для обратного распространения ошибки. Эти технологии сделали обучение нейронных сетей вполне осуществимым за более короткое время и доступным для большего количества людей. Исследование также открыло мир новых приложений. Нейронные сети могут учиться на огромных объемах данных и использовать в своих интересах большие данные (образы, текст, транзакции и данные социальных сетей), создавая, таким образом, такие модели, которые непрерывно становятся лучше, в зависимости от путей передаваемых им данных.

Новые тенденции определяют крупные игроки, такие как Google, Facebook, Microsoft и IBM, и с 2012 года они начали приобретать компании и нанимать экспертов (сейчас Хинтон работает в Google; Лекун, создатель сверточных нейронных сетей, возглавляет исследования искусственного интеллекта в Facebook) в новых областях глубокого обучения. Проект Google Brain реализовали Эндрю Ын и Джефф Дин (Jeff Dean), соединившие 16 тысяч компьютеров, чтобы создать сеть глубокого обучения с более чем миллиардом коэффициентов, обеспечив, таким образом, неконтролируемое обучение из видео на YouTube. Компьютерная сеть могла даже различать котов безо всякого человеческого вмешательства (об этом можно прочитать в статье от *Wired* на странице <https://www.wired.com/2012/06/google-x-neural-network/>).

Различия в глубоком обучении

Глубокое обучение может показаться просто большой нейронной сетью, выполняющейся на большем количестве компьютеров. Другими словами, это только математика и прорыв вычислительных технологий, который сделал такие большие сети возможными. Однако глубокое обучение — это, несомненно, качественное изменение по сравнению с поверхностными нейронными сетями. Это куда больше, чем просто смещение парадигмы. Глубокое обучение сместило парадигму машинного обучения со средств создания (средств, упрощающих обучение и выполняющих предварительный анализ используемых данных) к средствам обучения (автоматически создаваемым сложным средствам на основании фактических средств). Такой аспект не мог быть определен на меньших сетях, но становится очевидным при использовании множества уровней нейронной сети и больших объемов данных.

Рассматривая глубокое обучение, можно с удивлением обнаружить множество старых технологий, но еще удивительнее то, что все работает как никогда

При нынешнем положении дел у людей сложилось нереальное представление о том, как глубокое обучение может помочь обществу в целом. Вы видите, что приложение глубокого обучения побеждает кого-то в шахматах и думаете, что если они способны на эту действительно удивительную вещь, то на какие еще удивительные вещи они способны? Проблема в том, что в глубоком обучении не очень хорошо разбираются даже его сторонники. В технических изданиях о глубоком обучении авторы нередко описывают организованные в сеть уровни обработки туманно, без каких-либо объяснений фактически происходящего в каждой из этих коробок. Главное — помнить, что даже глубокое обучение фактически ничего не понимает. Оно использует огромные массивы примеров, чтобы получить статистически обоснованное соответствие шаблону, и использует математические принципы. Когда искусственный интеллект выигрывает игру, подразумевающую действия в лабиринте, понятия о лабиринте он не имеет; он просто знает, что определенные входные данные при обработке некоторыми способами создают определенные выходные данные, ведущие к победе.

В отличие от людей, глубокое обучение должно полагаться на огромные количества примеров, чтобы выявить конкретные отношения между вводом и выводом. Если вы скажете ребенку, что все люди от одного возраста до другого (уже не ребенок, но еще не взрослый) являются подростками, то ребенок будет в состоянии распознать любого соответствующего этой категории с высоким процентом точности, даже если это совершенно чужой человек. Для решения той же задачи глубокое обучение требовало бы специального обучения, и его будет довольно просто ввести в заблуждение, поскольку отсутствовавшие в его практике примеры ему непонятны.

Люди способны также создавать иерархии знаний безо всякого обучения. Например, без большого усилия понятно, что собаки и коты — это животные. Кроме того, зная, что собаки и коты — это животные, человек может легко сделать скачок и увидеть других животных как животных даже без специфического обучения. Глубокое обучение требовало бы отдельного обучения по каждому объекту, который является животным. Короче говоря, глубокое обучение не может перенести то, что оно знает, на другие ситуации, как люди.

Даже при этих ограничениях глубокое обучение — это удивительный, но не единственный инструмент в арсенале искусственного интеллекта. Наилучший способ применения этой технологии — поиск неочевидных людям шаблонов. *Шаблоны (pattern)* — это основной элемент обнаружения чего-то нового. Например, люди довольно давно борются с раком. Выявление глубоким обучением неочевидных людям шаблонов может стать серьезным шагом к решению с куда меньшими усилиями, что очень нужно людям.

прежде. Поскольку исследователи, наконец, сумели заставить сотрудничать некоторые простые, но очень хорошие решения, появилась возможность автоматически фильтровать, обрабатывать и преобразовывать большие данные. Например, такие новые активизации, как ReLU, не так уж и новы; они были известны еще со времен перцептрона. Кроме того, возможности по распознаванию образов, сделавшие глубокое обучение столь популярным, тоже не новы. Первоначально глубокое обучение получило хороший толчок благодаря *сверточным нейронным сетям* (Convolutional Neural Network — CNN), открытым в 1980-х годах французским ученым Яном Лекуном (Yann LeCun) (чья личная домашняя страница находится по адресу <http://yann.lecun.com/>). Такие сети и сейчас дают удивительные результаты, поскольку используют много уровней нейронов и много данных. То же самое относится к технологии, позволяющей машине понимать человеческую речь или переводить с одного языка на другой; этим технологиям уже десятилетия, но исследователи вернулись к ним в новой парадигме глубокого обучения.

Конечно, часть различий относится также к данным (подробности — далее), улучшению применения GPU и работе с компьютерными сетями. Вместе с *параллелизмом* (много собранных в кластеры компьютеров, работающих параллельно), GPU позволяют создавать большие сети и успешно обучать их на большом количестве данных. Фактически процессоры GPU выполняют определенные операции примерно в 70 раз быстрее любого процессора CPU, позволяя сократить время обучения нейронных сетей с недель до дней или даже часов.



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Более подробная информация о том, насколько процессоры GPU способны ускорить машинное обучение с использованием нейронной сети, приведена в технической статье по адресу <https://icml.cc/2009/papers/218.pdf>.

Поиск еще более умных решений

Глубокое обучение влияет на эффективность искусственного интеллекта в решении таких задач, как распознавание образов, машинный перевод и распознавание речи, которыми классический искусственный интеллект и машинное обучение занимались изначально. Кроме того, оно представляет новые, более выгодные решения.

- » Непрерывное обучение с использованием *дистанционного обучения*.
- » Многократное использование решений с использованием *переноса обучения*.

- » Демократизация искусственного интеллекта с помощью среды с открытым исходным кодом.
- » Простые и понятные решения с использованием *сквозного обучения*.

Дистанционное обучение

Нейронные сети существенно гибче других машинных алгоритмов обучения, они способны не прекращать обучение, даже уже осуществляя прогнозирование и классификацию. Эта возможность следует из алгоритмов оптимизации, позволяющих нейронным сетям учиться не только на небольших регулярных выборках примеров (*групповое обучение* (batch learning)), но даже на отдельных примерах (*дистанционное обучение* (online learning)). Сети глубокого обучения могут вырабатывать свои знания шаг за шагом и быть восприимчивыми к новой информации (как ум ребенка, который всегда открыт для новых стимулов и нового опыта). Приложение глубокого обучения на веб-сайте социальной сети, например, может обучаться на образах кота. Когда люди публикуют фотографии котиков, приложение распознает их и соответствующим образом отмечает. Когда люди начинают публиковать в социальной сети фотографии собак, нейронная сеть не обязана обучаться заново; она может продолжить обучение и на образах собак. Эта возможность особенно полезна при копировании столь разнообразных данных Интернета. Сеть глубокого обучения может быть открытой для нового и адаптировать свои коэффициенты так, чтобы приспособиться к изменениям.

Перенос обучения

Гибкость удобна, ведь когда сеть закончит свое обучение, вы сможете использовать ее многократно и для других целей, отличных от первоначальной. Сети, различающие объекты и правильно их классифицирующие, требуют долгого времени обучения и больших вычислительных мощностей. Расширение возможностей сети новым видам образов, которые не были частью первоначального обучения, означает перенос знаний на эту новую задачу (*перенос обучения* (transfer learning)).

Например, вы вполне можете перестроить сеть, способную различать собак и кошек, так, чтобы она отличала блюда с макаронами от блюд с сыром. Большинство уровней сети вы используете как есть (вы замораживаете их), а завершающие уровни (уровни вывода), вы видоизменяете, осуществляя *точную настройку* (finetuning). В скором времени и с меньшим количеством примеров сеть применит то, что она изучила при различении собак и котиков, к макаронам и сыру. Получится даже лучше, чем нейронная сеть, обученная распознавать только макароны и сыр.

Перенос обучения — это нечто новое для большинства машинных алгоритмов обучения и открывающее рынок, возможный для передачи знаний из одного приложения в другое, от одной компании другой. Компания Google уже фактически делает это, позволяя совместно использовать свое огромное хранилище данных публичным сетям, построенным на этих данных (см. подробности на <https://techcrunch.com/2017/06/16/object-detection-api/>). Это этап демократизации глубокого обучения, когда доступ к его потенциальным возможностям разрешается всем.

Среда с открытым исходным кодом

Сегодня сети доступны всем, включая доступ к инструментальным средствам создания сетей глубокого обучения. Это вопрос не только публичного разглашения научных трудов, объясняющих принципы работы глубокого обучения, но и вопрос программирования. В первые дни глубокого обучения вы должны были строить каждую сеть с самого начала, как приложение, разработанное на таком языке, как C++, что ограничивало круг посвященных несколькими хорошо подготовленными специалистами. Возможности современных языков (таких, например, как Python; см. <http://www.python.org>) куда лучше благодаря многочисленным реализациям глубокого обучения средами с открытым исходным кодом, такими как TensorFlow от Google (<https://www.tensorflow.org/>) или PyTorch от Facebook (<http://pytorch.org/>). Используя простые команды, эти среды выполнения позволяют повторить последние достижения в области глубокого обучения.



ЗАПОМНИ

Когда есть свет, есть и тени. Для работы нейронные сети нуждаются в огромных объемах данных, а данные доступны не для всех, поскольку большие организации не стремятся разглашать их публично. Перенос обучения может снизить влияние нехватки данных, но только частично, поскольку определенные приложения действительно требуют фактических данных. Следовательно, демократизация искусственного интеллекта ограничивается. Кроме того, системы глубокого обучения настолько сложны, что их выводы приходится объяснять, и иногда это довольно трудно (с учетом процветания предвзятостей и дискриминаций) и неоднозначно, поскольку хитрости вполне могут ввести такие системы в заблуждение (см. <https://www.dvhardware.net/article67588.html>). Любая нейронная сеть может быть чувствительной к *сопоставительным атакам* (adversarial attack), когда манипуляции входными данными позволяют обмануть систему и обеспечить пристрастные ответы.

Сквозное обучение

И наконец, глубокое обучение допускает *сквозное обучение* (end-to-end learning), а значит, и решение задачи более простым и понятным способом, чем предыдущее решение для глубокого обучения, которое могло бы потребовать куда больших усилий при решении задачи. Вы можете захотеть решить такую трудную задачу искусственного интеллекта, как распознавание известных лиц или управление автомобилем. Используя классический подход искусственного интеллекта, вы разделили бы задачу на более простые части, чтобы достичь приемлемого результата за вполне реальное время. Например, если вы хотите распознавать лица по фотографии, то учтите, что существующие системы искусственного интеллекта делят эту задачу на следующие части.

1. Поиск лица на фотографии.
2. Обрезка фотографии по лицу.
3. Обработка вырезанного лица с учетом позы, подобно фотографии в удостоверении личности.
4. Подача обработанного вырезанного лица в качестве учебного примера нейронной сети для распознавания образов.

Сегодня вы можете передавать фотографии архитектуре глубокого обучения и учить ее искать лица по образу, а затем классифицировать их. Вы можете использовать тот же самый подход для языкового перевода, распознавания речи или даже беспилотных автомобилей (они рассматриваются в главе 14). Во всех этих случаях вы просто передаете ввод системе глубокого обучения и получаете требуемый результат.

Выявление краев и форм образов

Сверточные нейронные сети (Convolutional Neural Network — ConvNet или CNN) обусловили нынешний ренессанс глубокого обучения. В следующих разделах обсуждается, как сети CNN помогают обнаруживать края образов и форм для таких задач, как распознавание рукописного текста.

Распознавание символов

Идея сетей CNN не нова. Они появились в конце 1980-х годов как результат работ Яна Лекуна (ныне директора по искусственному интеллекту в компании Facebook), когда он работал в исследовательской лаборатории AT&T Labs-Research вместе с Йошуа Бенгио, Леоном Боттоу (Leon Bottou) и Патриком Хаффнером (Patrick Haffner) над сетью LeNet5. Об этой сети можно узнать по

адресу <http://yann.lecun.com/exdb/lenet/>, или можно посмотреть видео, в котором о ней рассказывает сам молодой Лекун: https://www.youtube.com/watch?v=FwFduRA_L6Q. В те времена наличие машины, способной декодировать рукописные числа, было настоящим чудом, помогавшим почтовой службе в автоматизации распознавания почтовых индексов при сортировке входящей и исходящей почты.

Ранее разработчики достигли некоторых результатов, применив нейронную сеть для образов искомых чисел. Каждый пиксель образа объединялся с узлом в сети. Проблема использования такого подхода была в том, что сеть не могла достичь *трансляционной инвариантности* (translation invariance), т.е. возможности декодировать число при различных условиях, включая размер, искажения или позицию в образе, как иллюстрируется на рис. 11.3. Эта нейронная сеть могла обнаружить только подобные числа, т.е. те, которые она видела прежде. Кроме того, она делала много ошибок. Предварительное преобразование образа перед его передачей нейронной сети частично решило проблему: изменение размеров, перемещение, очистка пикселей и создание специальных блоков информации улучшало вычисления в сети. Эта методика *создания возможностей* (feature creation) требует как достаточно сложных преобразований образа, так и множества вычислений в ходе анализа данных. Над задачами распознавания образов тогда работали скорее умельцы, чем ученые.

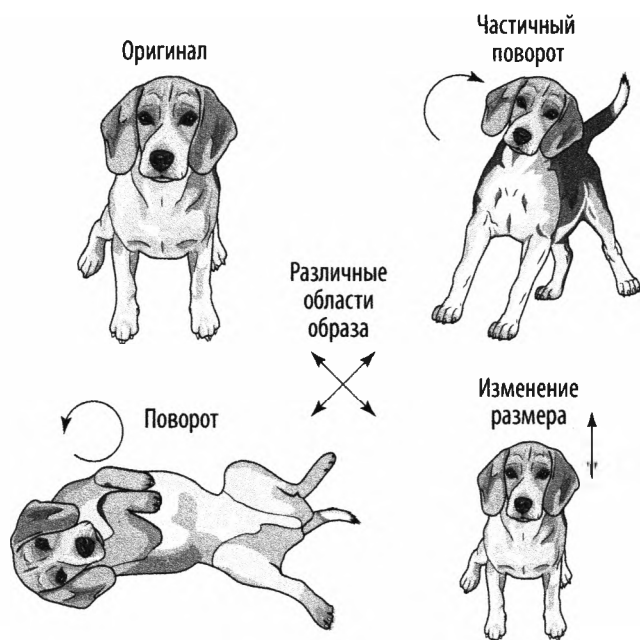


Рис. 11.3. Используя трансляционную инвариантность, нейронная сеть способна различить собаку и ее вариации

Свертка легко решает проблему трансляционной инвариантности, поскольку она предоставляет совершенно иной подход к обработке изображений в нейронной сети. Свертка — это основа сети LeNet5, она обеспечивает фундаментальные стандартные блоки практически для всех сетей CNN, осуществляющих следующее.

- » **Классификация образов.** Определение, какой объект присутствует на образе.
- » **Обнаружение образов.** Поиск, где именно объект находится на образе.
- » **Сегментация образа.** Разделение образа на области согласно их содержанию; например, в образе дороги можно разделить саму дорогу от находящихся на ней автомобилей и пешеходов.

Как работает свертка

Начнем объяснение работы свертки с ввода, являющегося образом, состоящим из одного или нескольких уровней пикселей, называемых *каналами* (channel), используя значения от 0 (пиксель полностью выключен) до 256 (пиксель включен). Например, у образа в формате RGB есть индивидуальные каналы для красного, зеленого и синего цветов. Смешение этих каналов создает всю палитру цветов, которые вы видите на экране.

Входные данные подвергаются простым преобразованиям по смене диапазона значений пикселя (например, можно задать диапазон от нуля до единицы), а затем передаются. Преобразование данных упрощает работу свертки, поскольку это просто операции умножения и суммирования, как показано на рис. 11.4. Свертка — это нейронный уровень, получающий небольшие части образа, умножающий значения пикселей в части решетки на специально разработанные числа, а затем суммирующий все результаты умножения и проецирующий это на следующий нейронный уровень.

Такая обработка очень гибка, поскольку основание для умножения в свертки чисел формирует обратное распространение ошибки (см. статью <https://ujjwalkarn.me/2016/08/11/intuitive-explanation-convnets/> о работе свертки, включая анимацию), а значения, фильтруемые сверткой, являются характеристиками образа, которые нужны нейронной сети для решения ее задачи по классификации. Некоторые свертки захватывают только линии, некоторые — только кривые или специальные шаблоны независимо от того, где они присутствуют в образе (и это свойство трансляционной инвариантности свертки). Поскольку данные образа проходят через разные свертки, они преобразуются, собираются и представляются все более и более сложными шаблонами, пока свертывание не создаст обобщенный образ (например,

образ среднего кота или собаки), который обучаемая сеть CNN впоследствии использует для обнаружения новых образов.

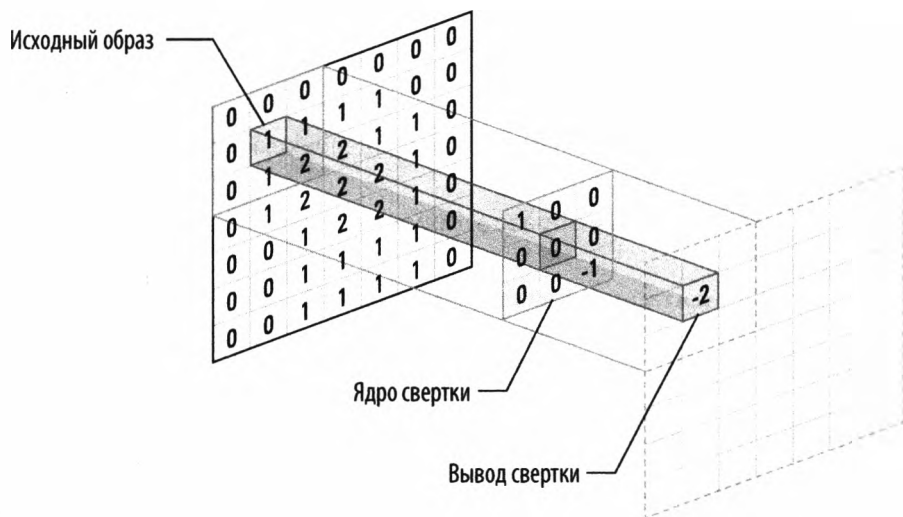
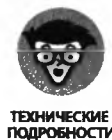


Рис. 11.4. Сверточное сканирование образа



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Если хотите узнать больше о свертывании, посмотрите визуализацию, созданную исследователями Google из Research and Google Brain, — визуализацию внутренней работы сети GoogleLeNet с 22 уровнями, разработанной учеными Google (см. <https://distill.pub/2017/feature-visualization/>). В приложении <https://distill.pub/2017/feature-visualization/appendix/> приводятся примеры того, как уровни обнаруживают вначале края, затем — текстуры, затем — общие шаблоны, затем — части и наконец — все объекты.

Интересно, но установка базовой архитектуры ConvNet вовсе не сложна. Достаточно знать, что чем больше уровней вы имеете, тем лучше. Вы задаете количество уровней свертывания и некоторые характеристики поведения свертывания: как создается решетка (значения *фильтр*, *ядро* и *датчик средств*), как решетка перемещается по образу (*шаг*) и как он будет вести себя вокруг границ образа (*дополнение*).



ЗАПОМНИ

Рассматривая работу свертывания, приходишь к выводу, что оно продвигается куда глубже в глубоком обучении, а значит, данные подвергаются куда более глубоким преобразованиям, чем при обработке любым машинным алгоритмом обучения или поверхностной нейронной сетью. Чем больше уровней, тем большим преобразованиям подвергается образ и тем глубже он становится.

Соревнования с использованием образов

Сеть CNN — это хорошая идея. Телекоммуникационная компания AT&T фактически реализовала сеть LeNet5 для проверки пользователей АТМ. Однако в середине 1990-х началась следующая зима искусственного интеллекта: множество исследователей и инвесторов потеряли веру в способность нейронных сетей реализовать искусственный интеллект. Кроме того, в те времена наблюдался существенный недостаток данных. Исследователи достигли результатов, сравнимых с LeNet5, используя такие новые алгоритмы машинного обучения, как Support Vector Machine (от научной школы аналоговистов) и Random Forest (случайный лес), улучшенный вид деревьев решений от научной школы симболистов (см. главу 10).

Лишь некоторые из исследователей, такие как Джеффри Хинтон, Ян Лекун и Йошуа Бенгио, продолжили разработки технологий нейронной сети до появления новых наборов данных и обеспечили прорыв, закончивший зиму искусственного интеллекта. Тем временем в 2006 году увидели свет работы Фей-Фей Ли (Fei-Fei Li), профессора информатики в Иллинойском университете в Урбане-Шампейне (ныне руководитель исследовательских работ в Google Cloud, а также профессор в Стэнфорде), позволившие обеспечить более реалистичные наборы данных для лучшей проверки алгоритмов. Она начала накапливать невероятные количества образов, представляя большие количества классов объектов. Она и ее группа решили эту грандиозную задачу, используя такую службу от Amazon как Mechanical Turk, которую вы используете для решения небольших задач (таких, как классификация образов) за небольшую плату.

Полученный набор данных, законченный в 2009 году, получил название “ImageNet” и содержал 3,2 миллиона маркированных образов, упорядоченных в 5247 иерархически организованных категорий. Вы можете ознакомиться с этим набором данных по адресу <http://www.image-net.org/> или прочитать оригинальную публикацию http://www.image-net.org/papers/image-net_cvpr09.pdf. Довольно скоро набор данных ImageNet был использован в 2010 соревнованиях, доказавших возможность нейронных сетей правильно классифицировать образы, упорядоченные в 1000 классов.

Через семь лет соревнований (они прекращены в 2017 году) победившие алгоритмы повысили точность распознавания образов с 71,8 до 97,3 процента, что превосходит человеческие возможности (да, люди делают ошибки при классификации объектов). Исследователи с самого начала обратили внимание, что при больших количествах данных их алгоритмы начинали работать лучше (тогда ничего такого, как ImageNet, еще не было), а затем они начали проверять новые идеи и улучшили архитектуры нейронной сети.

Даже если соревнований ImageNet больше не будет, исследователи все равно разработают больше архитектур CNN, увеличат их точность и возможности распознавания, а также повысят надежность. Фактически многие решения для глубокого обучения все еще экспериментальны и пока не применялись в таких важных областях, как банковское дело или безопасность, и не только из-за трудностей в интерпретации, но и из-за возможных уязвимостей.



ВНИМАНИЕ

Уязвимости бывают разными. Исследователи выяснили, что внесение специально подобранных искажений, таких как изменение единственного пикселя в изображении, способно радикально изменить ответы CNN как *не целенаправленно* (когда достаточно только ввести CNN в заблуждение), так и *целенаправленно* (когда CNN должна дать некий конкретный ответ). Об этом вопросе можно узнать больше в руководстве OpenAI по адресу <https://blog.openai.com/adversarial-example-research/>. OpenAI — это некоммерческая компания по исследованию искусственного интеллекта. Статья “One pixel attack for fooling deep neural networks” ([https://arxiv.org/abs/1710,08864](https://arxiv.org/abs/1710.08864)) также будет полезна. Дело в том, что технология CNN еще недостаточно защищена. Вы не можете просто использовать ее вместо собственных глаз; здесь пока следует соблюдать большую осторожность.

Обучение подражанию искусству и жизни

Сети CNN повлияли не только на задачи компьютерного зрения, но и на многие приложения (например, они необходимы для функции зрения беспилотных автомобилей). Сети CNN убедили многих исследователей инвестировать время и силы в революцию глубокого обучения. Последующие исследования и разработки взрастили новые идеи. Последующие разработки привели, наконец, к нововведениям в искусственном интеллекте, позволив компьютерам научиться понимать разговорный язык, переводить написанное на иностранные языки, а также создавать модифицированные тексты и изображения, демонстрируя, таким образом, как сложные вычисления статистических распределений могут быть преобразованы в своего рода искусство, творческий потенциал и воображение. Если вы говорите о глубоком обучении и его возможных применениях, вы также должны упомянуть *рекуррентные нейронные сети* (Recurrent Neural Network — RNN) и *генеративно-состязательные сети* (Generative Adversarial Network — GAN), или у вас не будет ясной картины того, что глубокое обучение может сделать для искусственного интеллекта.

Запоминание важных последовательностей

Одним из недостатков сети CNN была нехватка памяти. Она справлялась с пониманием одиночного изображения, но попытки понять изображение в контексте, как кадр в видео, сталкивались с неспособностью найти правильные ответы на трудные проблемы искусственного интеллекта. Наиболее важной проблемой оказались последовательности. Если вы хотите понять книгу, то читаете ее постранично. Последовательности бывают вложенными. В пределах страницы есть последовательность слов, а в пределах слова — последовательность символов. Чтобы понять книгу, необходимо понять последовательности символов, слов и страниц. Ответом стали сети RNN, поскольку при обработке текущих входных данных они отслеживают и прошлые входные данные. При обработке в такой сети исходные данные продвигаются не только в одном направлении от входа к выходу, как в обычной нейронной сети, но могут и по кругу. Это как будто сеть слышит собственное эхо.

Если передать сети RNN последовательность слов, то она узнает, что, встретив некое слово, которому предшествуют другие определенные слова, можно решить, как закончить фразу. Сети RNN — не просто технология, способная автоматизировать компиляцию ввода (это как автоматическое завершение строк в браузере при вводе искомых слов). Сеть RNN может получать последовательности и предоставлять на выходе перевод, такой как общий смысл фразы (таким образом, искусственный интеллект теперь может устранять неоднозначность фраз, когда формулировка важна), или переводить текст на другой язык (снова работа трансляции в контексте). Это работает даже со звуками, поскольку определенные звуковые модуляции вполне можно интерпретировать как слова. Сеть RNN позволяет компьютерам и мобильным телефонам с высокой точностью понять не только то, что вы сказали (это та же самая технология, что и в автоматическом создании субтитров), но и то, что вы хотели сказать, позволяя компьютерным программам общаться с вами, как Siri, Cortana и Alexa.

Магия общения искусственного интеллекта

Чатбот (chatbot) — это программное обеспечение, способное общаться с вами двумя способами: устно (вы говорите с ним и слушаете ответы) и письменно (вы вводите то, что хотите сказать, и читаете ответ). Возможно, вы слышали об этом под другими названиями (диалоговый агент, виртуальный собеседник, токбот и др.), но дело в том, что вы уже можете использовать их на своем смартфоне, компьютере или специальном устройстве. Такие примеры, как Siri, Cortana и Alexa, известны всем. Когда вы обращаетесь в клиентскую службу компаний, можете общаться с ее чатботом по телефону, через веб или

приложение на вашем мобильном телефоне, если используете Twitter, Slack, Skype или другое приложение связи.

Чатботы — это серьезный бизнес, поскольку они позволяют компаниям экономить деньги на операторах клиентской службы (поддержка контакта с постоянными клиентами и их обслуживание), но сама идея не нова. Хотя название было придумано не так давно (в 1994 году Майклом Молдином (Michael Mauldin), разработчиком поискового механизма Lycos), чатботы считают вершиной искусственного интеллекта. Согласно видению Алана Тьюринга, отличить сильный искусственный интеллект от человека при разговоре не должно быть возможно. Тьюринг разрабатывал известный тест на основании общения, который позволяет удостовериться, достиг ли искусственный интеллект уровня человека.



ЗАПОМНИ

Когда искусственный интеллект демонстрирует осмысленное поведение, но не сознательное, как человек, это слабый искусственный интеллект. Сильный искусственный интеллект может действительно думать, как человек.

Тест Тьюринга требует, чтобы являющийся арбитром человек общался через компьютерный терминал с двумя субъектами: человеком и машиной. В ходе общения он должен выяснить, какой из собеседников является искусственным интеллектом. Тьюринг утверждал, что если искусственный интеллект сможет убедить человека в том, что он общается с другим человеком, то можно утверждать, что искусственный интеллект равен человеку по уровню интеллекта. Проблема довольно трудна, поскольку это вопрос не только грамматически правильного ответа с учетом контекста (место, время и характеристики человека, с которым общается искусственный интеллект), но и демонстрация собственной индивидуальности (искусственный интеллект должен подходить на реальную личность и по образованию, и по мировоззрению).

Начиная с 1960-х годов попытки прохождения Теста Тьюринга мотивировали разработку чатботов на базе *модели подтверждения/опровержения гипотезы* (retrieval-based model). Таким образом, для обработки речевого ввода при общении с человеком используется *обработка текстов на естественном языке* (Natural Language Processing — NLP). Определенные слова или наборы слов приводят к заранее заданным ответам и реакциям чатбота, экономя память.



СОВЕТ

NLP — это анализ текстовых данных. Алгоритм разделяет текст на лексемы (элементы фразы, такие как существительные, глаголы и прилагательные) и удаляет любую бесполезную или неважную информацию. Размеченный текст обрабатывается с использованием статистических операций или машинного обучения. Например, NLP

может отметить части речи, идентифицировать слова и их смысл, а также определить, подобен ли один текст другому.

Первый чатбот этого вида, ELIZA, создал Джозеф Вейценбаум (Joseph Weizenbaum) в 1966 году в форме компьютера-психолога. В основе ELIZA была простая эвристика, она выбирала подходящий ответ из фиксированного набора ответов исходя из базовых фраз и ключевых слов в соответствующем контексте. Сетевую версию ELIZA можно опробовать по адресу <http://www.masswerk.at/elizabot/>. Вас может удивить осмысленность общения с ELIZA, например ее диалог с ее создателем: http://www.masswerk.at/elizabot/eliza_test.html.

Модели подтверждения/опровержения гипотезы работают прекрасно, когда общаются с использованием предварительно заданных тем, поскольку они включают человеческое знание, подобно экспертным системам (как обсуждалось в главе 3), поэтому они вполне могут отвечать адекватными грамматически правильными фразами. Проблемы возникают, когда вопросы не относятся к заданным темам. Чатбот может парировать эти вопросы, возвращая их назад в несколько иной форме (как это делает ELIZA) и выглядеть, как искусственный собеседник. Решение заключается в создании новых фраз на основании, например, статистических моделей, машинного обучения или даже предварительно обученных сетей RNN, способных произносить нейтральную речь или отражать индивидуальность конкретного человека. Это подход *генеративной модели* (generative-based model), являющейся последним достижением современных ботов, поскольку создание речи “на лету” вовсе не просто.

Генеративные модели не всегда отвечают адекватными фразами, но исследователи недавно добились новых успехов, особенно в сетях RNN. Как уже упоминалось, секрет кроется в последовательностях: вы предоставляете исходную последовательность на одном языке, и эта последовательность выводится на другом языке, как в задаче машинного перевода. В данном случае вы предоставляете и исходную последовательность, и результирующую на том же языке. Ввод — это часть сеанса общения, а вывод — последующая реакция.

С учетом текущего состояния в создании чатботов сети RNN прекрасно подходят для коротких диалогов, но получить хорошие результаты для более длинных или более четко сформулированных фраз куда сложнее. Подобно моделям подтверждения/опровержения гипотезы, сети RNN помнят полученную информацию, но не организованным способом. Если темы беседы ограничены, такие системы способны обеспечить хорошие ответы, но все становится куда хуже, когда контекст не задан, поскольку они нуждались бы в знании, сравнимом с приобретенным человеком на протяжении жизни. (Люди являются хорошими собеседниками благодаря имеющимся знаниям и опыту.)

Данные для обучения сети RNN действительно являются ключом. Например, чатбот Google, Google Smart Reply, довольно быстро отвечает на электронные сообщения. Более подробная информация о предполагаемой работе этой системы приведена в статье <https://ai.googleblog.com/2015/11/computer-respond-to-this-email.html>. В реальности он предпочитал отвечать на большинство вопросов фразой “I love you” (я вас люблю), поскольку использованные для его обучения примеры оказались тенденциозными. Нечто подобное случилось с чатботом Tay от Microsoft Twitter, которого подвела способность учиться на общении с пользователями (<https://www.businessinsider.com/microsoft-deletes-racist-genocidal-tweets-from-ai-chatbot-tay-2016-3>).



Если хотите узнать о современном положении дел с чатботами в мире, можете почитать о ежегодных соревнованиях чатботов, в которых тест Тьюринга применяется к текущей технологии. Последним лауреатом приза Lobner на момент написания этой книги была программа Mitsuku, которая хотя и не смогла пройти тест Тьюринга, вполне смогла рассуждать о конкретных объектах, предложенных во время беседы; она может также играть в игры и даже показывать фокусы (<http://www.mitsuku.com/>).

Соревнования между двумя искусственными интеллектами

Сеть RNN может позволить компьютеру общаться с вами, и если вы понятия не имеете, что нейронная сеть просто реагирует на последовательности слов, изученные ею ранее, может создаться впечатление о наличии внутри нее некоего интеллекта. В действительности никаких мыслей или рассуждений здесь нет, хотя эта технология и не просто произносит предварительно заданные фразы, она их формулирует.

Генеративно-состязательные сети (Generative Adversarial Network — GAN) являются другим видом технологии глубокого обучения, способной создать даже более сильную иллюзию наличия у искусственного интеллекта творческого потенциала. Эта технология также полагается на запоминание предыдущих примеров и понимание машиной того, что примеры содержат правила, которыми машина может играть, как ребенок играет кубиками (технически правила — это статистические распределения, лежащие в основе примеров). Тем не менее сети GAN — это невероятная технология, которую, вероятно, мы увидим во многих приложениях будущего.

Сети GAN возникли из работ нескольких исследователей из отдела Département d'informatique et de recherche opérationnelle Монреальского университета в 2014 году, самым известным из которых является Айан

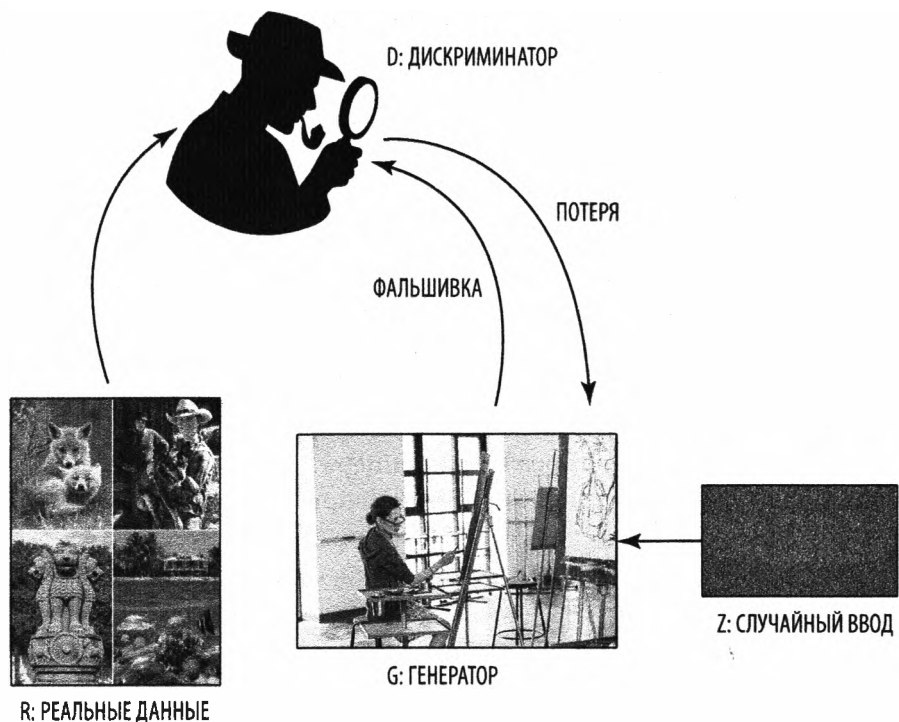
Гудфеллоу (Iam Goodfellow) (см. официальный документ <https://arxiv.org/pdf/1406.2661.pdf>). Предложенный новый подход глубокого обучения немедленно вызвал интерес и теперь является одной из наиболее исследованных технологий, с постоянными разработками и усовершенствованиями. По мнению Яна Лекуна, генеративно-сопоставительные сети — “самая интересная идея в машинном обучении за прошлые десять лет”. В интервью MIT Technology Review Айан Гудфеллоу объясняет нынешний уровень энтузиазма таким интригующим заявлением: “Вы можете считать генеративные модели некой формой воображения для искусственного интеллекта” (<https://www.technologyreview.com/lists/innovators-under-35/2017/inventor/ian-goodfellow/>).

Чтобы увидеть простую сеть GAN в действии (сейчас есть множество сложных вариаций, и разрабатывается еще больше), необходим набор справочных данных, обычно состоящий из реальных данных, примеры которых вы хотели бы использовать для обучения сети GAN. Например, если у вас есть набор данных из образов собаки, используя его, можно объяснить сети GAN, как выглядит собака. Узнав о том, как выглядят собаки, сеть GAN может предложить вероятные, реалистичные образы собак, отличающиеся от находившихся в начальном наборе данных. (Образы будут новыми; простую репликацию существующих образов считают ошибкой сети GAN.)

Отправная точка — это набор данных. Вы также нуждаетесь в двух нейронных сетях, каждая из которых специализируется на другой задаче и обе соревнуются между собой. Одна сеть является *генератором* (generator), получает случайный ввод (например, последовательность случайных чисел) и создает вывод (например, образ собаки), который является *артефактом* (artifact), поскольку он искусственно создан сетью-генератором. Вторая сеть — это *дискриминатор*, который должен правильно различать произведения генератора, артефакты, согласно примерам из учебного набора данных.

Когда сеть GAN начинает обучение, обе сети пытаются улучшаться, используя обратное распространение ошибки на основании результатов дискриминатора. Ошибки, которые допускает дискриминатор, отличая реальный образ от артефакта, распространяются дискриминатору (как при классификации нейронной сетью). Правильные ответы дискриминатора распространяются как ошибки генератору (поскольку он не смог сделать артефакты подобными образам из набора данных, и дискриминатор их определил). Эти отношения демонстрируются на рис. 11.5.

Первоначальные образы, выбранные Гудфеллоу для объяснения работы сети GAN, — это работы фальсификатора и эксперта. Эксперт становился все более и более квалифицированным в распознавании поддельных картин, но и фальсификатор совершенствовался, чтобы избежать разоблачения экспертом.



Фотографии предоставлены (по часовой стрелке от левой нижней): Lileepphoto/Shutterstock; Menno Schaefer/Shutterstock; iofoto/Shutterstock; vilainecrevette/iStockphoto; посередине: Rana Faure/Corbis/VCG/Getty Images.

Рис. 11.5. Как работает сеть GAN, колеблющаяся между генератором и дискриминатором

Вы можете задаться вопросом, как генератор учится создавать правильные артефакты, если он никогда не видел оригиналов. Только дискриминатор видит оригинальный набор данных, когда пытается отличать реальное искусство от артефактов генератора. Даже если генератор никогда ничего не исследует из исходного набора данных, он получает подсказки от дискриминатора. Это небольшие подсказки, получаемые генератором в результате многих неудачных попыток. Это как научиться рисовать Мону Лизу, никогда ее не видев, только с помощью друга, говорящего вам, хорошо ли ваше предположение. Ситуация напоминает теорему о бесконечных обезьянах, с некоторыми отличиями. В этой теореме вы ожидаете, что обезьяны напишут всего Шекспира, просто угадав (см. <https://www.npr.org/sections/13.7/2013/12/10/249726951/theinfinite-monkey-theorem-comes-to-life>). В данном случае генератор использует случайность только в начале, а затем он медленно учится на реакции

дискриминатора. После некоторых модификаций эта простая идея сетей GAN стала способна на следующее.

- » Создание фотореалистичных образов объектов, таких как предметы одежды, интерьеры или промышленный дизайн, на основании словесных описаний (вы просите желто-белый цветок и получаете его, как описано в статье <https://arxiv.org/pdf/1605.05396.pdf>).
- » Изменение существующих образов при повышении их разрешения, добавлении специальных шаблонов (превращение лошади в зебру <https://junyanz.github.io/CycleGAN/>) и заполнении отсутствующих частей (при удалении человека с фотографии сеть GAN заменяет пустое место неким вероятным фоном, как на этом образе <http://hi.cs.waseda.ac.jp/~iizuka/projects/completion/en/>).
- » Многие передовые приложения, такие как создающие движение из статических фотографий, сложные объекты как законченные тексты (называемые *структурированным прогнозом*, поскольку это не просто ответ, а скорее набор взаимосвязанных ответов), данные для контролируемого машинного обучения или даже мощную криптографию (<https://arstechnica.com/informationtechnology/2016/10/google-ai-neural-network-cryptography/>)



СОВЕТ

Сети GAN — это передовая технология глубокого обучения со многими неисследованными и новыми областями применения в искусственном интеллекте. Если у искусственного интеллекта будет образное и творческое мышление, он, вероятно, получит его от таких технологий, как GAN. Вы можете узнать больше об этой технологии из статей о GAN от OpenAI — некоммерческой компании по исследованию искусственного интеллекта, основанной Грегом Брокмэном (Greg Brockman), Ильей Суцкевером (Ilya Sutskever), Илоном Маском (Elon Musk) (основателем PayPal, SpaceX и Tesla) и Сэмом Альтманом (Sam Altman) (<https://blog.openai.com/generative-models/>).

4 Работа с искусственным интеллектом в аппаратных приложениях

В ЭТОЙ ЧАСТИ...

- » Работа с роботами**
- » Полеты с дронами**
- » Позвольте искусственному интеллекту вести автомобиль вместо вас**



Глава 12

Разработка роботов

В ЭТОЙ ГЛАВЕ...

- » Различение роботов в научной фантастике и в действительности
- » Рассуждение об этике роботов
- » Поиск применения роботов
- » Взгляд на работу робота изнутри

Люди часто принимают робототехнику за искусственный интеллект, но робототехника отличается от искусственного интеллекта. Искусственный интеллект нацелен на поиск решения некоторых трудных задач, связанных с человеческими способностями (такими, как распознавание объектов либо понимание речи или текста); робототехника стремится использовать машины для решения задач в физическом мире за счет частичной или полной автоматизации. Поэтому искусственный интеллект можно считать программным обеспечением для решения задач, а робототехнику — аппаратными средствами для того, чтобы сделать эти решения реальностью.

Робототехнические аппаратные средства могут использовать программные средства систем искусственного интеллекта, а могут и не использовать. Некоторые роботы люди контролируют дистанционно, как робот *da Vinci*, обсуждавшийся в разделе “Помощь хирургу” главы 7. Во многих случаях искусственный интеллект действительно обеспечивает усиление способностей человека, но контролирует ситуацию все еще человек. Роботы могут быть

очень разными, от получающих абстрактные распоряжения от людей (такие, как переместиться из точки А в точку В по карте или поднять объект) до полагающихся на искусственный интеллект для выполнения этих распоряжений. Некоторые роботы способны решать задачи автономно, безо всякого человеческого вмешательства. Интеграция искусственного интеллекта делает робот более умным и более полезным в решении задач, но роботам искусственный интеллект для работы нужен отнюдь не всегда. Человеческое воображение сделало эти два понятия неразделимыми благодаря научно-фантастическим романам и фильмам.

В этой главе рассматривается, как произошло это совмещение, а также современные реалии роботов и то, как широкое употребление решений с использованием искусственного интеллекта могло бы их преобразить. Промышленные роботы существуют с 1960-х годов. В главе упоминается также, что люди все чаще используют роботы в промышленности, в научных изысканиях, в медицине и в военных целях. Новейшие исследования в области искусственного интеллекта ускоряют этот процесс, поскольку они решают такие трудные для роботов задачи, как распознавание объектов в реальном мире, прогнозирование поведения людей, понимание голосовых команд, правильная речь, ходьба и (да!) кувырки. Об этом можно прочесть в следующей статье о недавних достижениях в робототехнике: <https://www.theverge.com/circuitbreaker/2017/11/17/16671328/bostondynamics-backflip-robot-atlas>.

Роли роботов

Роботы — это относительно недавняя идея. Слово возникло от чешского слова *robota*, что означает принудительную рабочую силу. Термин впервые появился в 1920 году в пьесе *Россумские универсальные роботы* чешского автора Карела Чапека. Однако человечество давно мечтало о механических существах. Древние греки придумали миф о бронзовом механическом человеке, Талосе, созданном богом металлургии, Гефестом, по требованию Зевса, отца богов¹. Греческие мифы содержат также упоминания других автоматов Гефеста, кроме Талоса. *Автомат* (automata) — это самоуправляемая машина, способная выполнять конкретную заранее заданную последовательность задач (в отличие от роботов, обладающих гибкостью для выполнения широкого диапазона задач). Греки фактически создали водно-гидравлические автоматы²,

¹ См. лучше <https://ru.wikisource.org/wiki/ЭСБЕ/Талос>. — *Примеч. ред.*

² Вероятно, имеется в виду торговый автомат Герона Александрийского, разработанный им в Египте по заказу жрецов приблизительно в I веке н.э. — *Примеч. ред.*

работавшие подобно алгоритмам, но в физическом мире. Подобно алгоритмам, автоматы имеют интеллект их создателя, а потому только создают иллюзию наличия самосознания у машин.

В Европе вы найдете примеры автоматов повсюду, начиная с греческой цивилизации, средневековья и эпохи Возрождения и заканчивая новейшей историей. Множество проектов было создано ближневосточным математиком и изобретателем аль-Джазари (см. <http://www.muslimheritage.com/article/al-jazari-mechanical-genius>). В Китае и Японии были собственные версии автоматов. Одни автоматы были очень сложными механизмами, а другие — хитрым розыгрышем, таким как Механический Турок, машина XVIII века, которая, казалось, была в состоянии играть в шахматы, но скрывала внутри человека.



ВНИМАНИЕ!

Важно отличать автоматы от других подобных человеку объектов. Например, Голем (<https://www.myjewishlearning.com/article/golem/>) является смесью глины и волшебства. Никаких механизмов, поэтому его нельзя отнести к устройствам, обсуждаемым в этой главе.

Роботы, описанные Чапеком, тоже не были абсолютно механическими автоматами, скорее живыми существами, спроектированными и собранными так, как будто они были автоматами. Его роботы обладали формой, подобной форме человека, и играли определенные роли в сообществе, заменяя людей-работчих. Напоминающие Франкенштейна Мэри Шелли, роботы Чапека были тем, что сегодня считается *андроидами* (android) — искусственными биоинженерными существами, описанными Филипом Киндредом Диком (Philip K. Dick) в книге *Мечтают ли андроиды об электроовцах?* (по мотивам которой был снят фильм *Бегущий по лезвию*). Таким образом, под *роботами* имеет смысл подразумевать автономные механические устройства, призванные не поражать и восхищать, а скорее производить товары и оказывать услуги. Кроме того, роботы стали центральной идеей в научной фантастике, в книгах и фильмах, благодаря которым в общественном воображении сложился образ робота как человекоподобного искусственного интеллекта, созданного, чтобы служить людям, что не так уж и далеко от первоначальной идеи слуги Чапека. Идея постепенно переходила от искусства к науке и технике, а также стала вдохновением для ученых и инженеров.



ЗАПОМНИ

Чапек создал идею как роботов, так и апокалипсиса роботов, когда искусственный интеллект в научно-фантастических фильмах уничтожает все. С учетом этого недавний прогресс в области искусственного интеллекта вызывает опасения у таких значительных фигур, как основатель компании Microsoft Билл Гейтс, физик Стивен

Хокинг, а также изобретатель и бизнесмен Илон Маск. Робототехнические рабы Чапека бунтуют против создавших их людей и, в конце концов, уничтожают почти все человечество.

Разоблачение научно-фантастических представлений о роботах

Первый коммерческий робот, Unimate, появился в 1961 году (<https://www.robotics.org/joseph-engelberger/unimate.cfm>). Это был просто роботизированный манипулятор (программируемая механическая рука, состоящая из металлических балок и суставов), рабочий конец которого мог держать, вращать или сваривать объекты согласно инструкциям, заданным человеком-оператором. Он был продан компании General Motors для использования в автомобилестроении. Робот Unimate должен был извлекать отливки из форм на сборочной линии и сваривать их вместе — задача физически опасная для рабочих-людей. Чтобы понять возможности такой машины, посмотрите видео: <https://www.youtube.com/watch?v=hxsWeVtb-JQ>. Реалиям современных роботов посвящены следующие разделы.

Законы роботехники

Еще до появления Unimate и задолго до появления многих других промышленных манипуляторов, которые начали работать рядом с рабочими-людьми на сборочных линиях, люди уже знали, как роботы должны выглядеть, действовать и даже думать. Американский писатель Айзек Азимов, известный своими работами в научной фантастике и популяризации науки, выпустил в 1950-х годах серию романов, в которых была предложена концепция роботов, совершенно отличных от используемых в промышленных условиях.



ЗАПОМНИ

Азимов придумал термин *роботехника* (robotics) и использовал его в том же смысле, в каком люди используют термин *механика*. Его богатое воображение до сих пор устанавливает нормы ожидаемого людьми от роботов. Азимов перенес роботов в эпоху исследования космоса, снабдив их позитронным мозгом, для помощи людям в ежедневном решении ординарных и экстраординарных задач. *Позитронный мозг* (positronic brain) — это вымышленное устройство, позволяющее роботам в романах Азимова действовать автономно и быть способными помочь или заменить людей во многих задачах. Кроме подобных человеческим возможностей в понимании и действии (сильный искусственный интеллект), позитронный мозг работает согласно трем законам роботехники, являющихся частью аппаратных средств, контролирующих поведение роботов моральным способом (moral way).

1. Робот не может причинить вред человеку или своим бездействием допустить, чтобы человеку был причинен вред.
2. Робот должен повиноваться всем приказам, которые дает человек, кроме тех случаев, когда эти приказы противоречат первому закону.
3. Робот должен заботиться о своей безопасности в той мере, в которой это не противоречит первому и второму законам.

Впоследствии автор добавил нулевое правило, с наиболее высоким приоритетом, чтобы гарантировать действия робота по обеспечению безопасности многих.

0. Робот не может причинить вред человечеству или своим бездействием допустить, чтобы человечеству был причинен вред.

В центре всех историй Азимова о роботах находится то, что эти три закона позволяют роботам работать вместе с людьми безо всякого риска апокалипсиса искусственного интеллекта или восстания машин. Эти три закона невозможно ни обойти, ни изменить, они выполняются в порядке приоритетов и присутствуют как математические формулировки в функциях позитронного мозга. К сожалению, в законах есть лазейки и неоднозначности, на которых основаны сюжеты большинства его романов. Эти три закона зафиксированы в вымышленном *Руководстве по роботехнике, 56-е издание, 2058 год*, и основаны на принципах безопасности, повиновения и самосохранения.

Азимов представлял себе Вселенную, в которой вы можете свести весь моральный мир (moral world) к нескольким простым принципам при наличии лишь некоторых рисков, которые и лежат в основе сюжетов его рассказов. В действительности Азимов полагал, что роботы — это инструменты и что эти три закона могли бы работать даже в реальном мире, если контролировать их использование (читайте подробности в интервью журналу *Compute!* за 1981 год: https://archive.org/stream/1981-11-compute-magazine/Compute_Issue_018_1981_Nov#page/n19/mode/2up). Несмотря на оптимизм Азимова у современных роботов нет следующих возможностей.

- » Понять три закона роботехники.
- » Выбрать действия согласно этим трем законам.
- » Обнаружить и осознать возможное нарушение этих трех законов.

Некоторые могут полагать, что современные роботы действительно не очень интеллектуальны, поскольку им не хватает этих возможностей, и они правы. Однако Исследовательский совет инженерных и физических наук (Engineering and Physical Sciences Research Council — EPSRC), являющийся

основным агентством Великобритании по финансированию исследований в технике и физике, рекомендовал в 2010 году пересмотреть законы роботехники Азимова применительно к реальным роботам с учетом текущих технологий. Результат получился весьма отличным от первоначальных установок Азимова (см. <https://www.epsrc.ac.uk/research/ourportfolio/themes/engineering/activities/principlesofrobotics/>). Эти пересмотренные принципы признают, что роботы могут даже убивать (в целях национальной безопасности), поскольку они — инструмент. Подобно всем другим инструментам, выполнение существующих законов и нравственных норм лежит на их пользователе-человеке, а не на машине, робот считается только исполнителем. Кроме того, некто (человек) всегда должен нести ответственность за результаты действий робота.



Принципы EPSRC представляют более реалистичную точку зрения на роботы и мораль, учитывая используемую сегодня технологию слабого искусственного интеллекта, но они могли бы также обеспечить частичное решение в случаях более совершенных технологий. В главе 14 обсуждаются проблемы, связанные с использованием беспилотных автомобилей, своего рода мобильных роботов, способных самостоятельно управлять автомобилем. Например, при рассмотрении *проблемы вагонетки* (trolley problem) далее в этой главе перед вами встанут возможные, хотя и маловероятные, моральные проблемы, которые поколеблют уверенность относительно автоматизированных машин, когда придет время сделать определенный выбор.

Фактические возможности роботов

Возможности существующих роботов все еще далеки от таковых для человекообразных роботов из романов Азимова, они принадлежат к разным категориям. Тип двуногого робота, придуманного Азимовым, является в настоящее время самым редким и наименее совершенным.

Наиболее распространена такая категория роботов, как манипулятор, подобный описанному ранее Unimate. Роботы этой категории так и называются: *манипуляты*. Вы можете найти их на фабриках работающими как индустриальные роботы, где они осуществляют сборку и сварку со скоростями и точностью, недоступными рабочим-людям. Некоторые манипуляторы используются даже в больницах для помощи в хирургических операциях. Диапазон движения манипуляторов ограничен, поскольку они интегрируются по месту применения (они могут немного перемещаться, но немного, поскольку у них нет ни мощных моторов, ни электрического сцепления), поэтому им требуется помощь специального технического персонала для перемещения на новое

место. Кроме того, используемые в промышленности манипуляторы обычно полностью автоматизированы (в отличие от хирургических устройств, дистанционно контролируемых хирургом, принимающим медицинские решения при операции). Сейчас во всем мире работает более миллиона манипуляторов, половина из которых расположена в Японии.

Второй по величине и скорости роста является категория *мобильных роботов*. В отличие от манипуляторов их специализация — перемещение с помощью колес, роторов, крыльев или даже ног. В этой большой категории можно найти роботы, подающие еду (<https://nypost.com/2017/03/29/dominos-delivery-robots-bring-pizza-to-the-final-frontier/>) или книги (<https://www.digitaltrends.com/cool-tech/amazon-prime-air-delivery-drones-history-progress/>), используемые на коммерческих предприятиях или даже в исследовании Марса (<https://mars.nasa.gov/mer/overview/>). Мобильные роботы являются главным образом беспилотными (в них никого нет) и контролируемые дистанционно, но их автономность постоянно увеличивается, и вы можете ожидать увидеть в этой категории более независимые роботы. Есть два специальных вида мобильных роботов: летающие *дроны* (drone) (глава 13) и беспилотные автомобили (глава 14).

Последний вид роботов — это *мобильный манипулятор*, способный перемещаться (как мобильные роботы) и манипулировать (как манипуляторы). Вершина этой категории — не просто робот, способный перемещаться и имеющий механическую руку, но также подражающий человеку по форме и поведению. *Гуманоидный робот* — это двуногое существо (с двумя ногами), имеющее человеческий торс и общающееся с людьми с помощью голоса и жестов. Этот тот вид робота, о котором мечтала научная фантастика, но создать его не так просто.

Почему так трудно быть гуманоидом

Человекоподобные роботы развиваются трудно, и ученые все еще работают над ними. Мало того что гуманоидные роботы требуют повышенных возможностей искусственного интеллекта, чтобы быть автономными, они также должны перемещаться, как люди. Самое большое затруднение, тем не менее, заключается в том, чтобы люди приняли машину, выглядящую, как человек. В следующих разделах рассматриваются различные аспекты создания гуманоидного робота.

Создание ходящего робота

Рассмотрим проблему робота, ходящего на двух ногах (*двуногий робот*). Люди этому учатся опытным путем, не задумываясь, но для робота это весьма проблематично. Четвероногие роботы легко держат равновесие и не

расходуют на это много энергии. Люди, напротив, расходуют энергию даже просто стоя, удерживая равновесие, как и при ходьбе. Гуманоидные роботы — как люди; они должны держать равновесие непрерывно, причем делать это эффективно и экономно. В противном случае роботу понадобится большой аккумулятор, который тяжел и громоздок, что усложняет проблему равновесия еще больше.

Предоставленное IEEE Spectrum видео дает куда лучшее представление о сложности реализации ходьбы. Видео демонстрирует роботов, участвующих в конкурсе DARPA Robotics Challenge (DRC), организованных агентством U.S. Defense Advanced Research Projects Agency с 2012 по 2015 год: <https://www.youtube.com/watch?v=g0TaYhjpofo>. Цель конкурса DRC заключалась в стимулировании робототехнических исследований, способных помочь при ликвидации последствий стихийных бедствий и проведении гуманитарных операций в обстановках, опасных для людей (<https://www.darpa.mil/program/darpa-robotics-challenge>). Поэтому вы видите, что роботы идут по пересеченной местности, открывают двери, берут инструменты (такие, как электродрель) или пытаются повернуть вентиль клапана. Недавно разработанный робот Atlas, от Boston Dynamics, оказался довольно многообещающим, как описано в статье <https://www.theverge.com/circuitbreaker/2017/11/17/16671328/bostondynamics-backflip-robot-atlas>. Робот Atlas действительно является исключительным, но ему все еще предстоит долгий путь обучения.



ЗАПОМНИ

Робот с колесами может легко двигаться по дороге, но в определенных ситуациях необходим робот человеческой формы, удовлетворяющий специфическим требованиям. Большинство инфраструктур в мире создано для человека. Наличие таких препятствий, как узкий проход, дверь или лестница, делает использование роботов других форм затруднительным. Например, во время чрезвычайной ситуации роботу, возможно, придется зайти на атомную электростанцию и перекрыть клапан. Человеческая форма позволяет роботу подойти, подняться по лестнице и закрутить колесо клапана (valve wheel).

Преодоление человеческого нежелания: зловещая долина

У людей есть проблема с гуманоидными роботами, выглядящими слишком похожими на людей. В 1970 году профессор Токийского технологического института Масахиро Мори (Masahiro Mori) изучал влияние роботов на японское общество. Он предложил термин 不気味の谷 (букими но тани), который переводится как “эффект зловещей долины”. Согласно исследованиям Мори чем реалистичнее выглядит робот, тем больше симпатий люди испытывают к нему. Симпатия растет, пока робот не достигнет определенной степени реализма,

после которой он нравиться перестает (даже вызывает отвращение). Отвращение увеличивается до тех пор, пока робот не достигает уровня реализма, делающего его копией человека. Эта зависимость приведена на рис. 12.1 и описана в оригинальной статье Мори <https://spectrum.ieee.org/automaton/robotics/humanoids/the-uncanny-valley>.



Рис. 12.1. Зловещая долина

О причинах отвращения, испытываемого людьми при контакте с роботом, почти, но не полностью похожего на человека, было сформулировано несколько гипотез. Предполагается, что люди обращают внимание на тон голоса робота, жесткость его движений и искусственную текстуру кожи. Некоторые ученые приписывают эффект зловещей долины культурным причинам, другие — психологическим или биологическим. Недавний эксперимент на обезьянах показал, что приматы подвержены подобной практике, когда им демонстрируют более или менее реалистично обработанные фотографии обезьян, выполненные по трехмерной технологии (см. <https://www.wired.com/2009/10/uncanny-monkey/>). Участвующие в эксперименте обезьяны демонстрировали небольшое отвращение к реалистичным фотографиям, намекая на наиболее популярную биологическую причину эффекта зловещей долины. Причиной, таким образом, может быть реакция самозащиты против существ, выглядящих неестественно, поскольку они, возможно, либо больны, либо даже мертвы.

Интересный момент эффекта зловещей долины: если необходимы гуманоидные роботы, чтобы помогать людям, следует также учитывать их уровень реализма и уделять внимание эстетическим деталям, чтобы получить положительную эмоциональную реакцию, которая позволит пользователям принимать

помощь робота. Недавние наблюдения показали, что роботы даже с небольшим подобием человеку вызывают привязанность у их пользователей. Например, согласно многим отчетам американских солдат они чувствовали потерю, когда подрывались их небольшие тактические роботы для обнаружения и обезвреживания взрывных устройств. (Статью об этом можно прочитать по адресу <https://www.technologyreview.com/s/609074/how-we-feelabout-robots-that-feel/>.)

Сотрудничество с роботами

У разных типов роботов разные области применения. Пока люди разработали и улучшили только три класса роботов (манипуляторы, мобильные и гуманоиды), для робототехники открылись новые области применения. Сейчас уже невозможно перечислить все случаи использования роботов, и в следующих разделах затрагивается лишь часть наиболее перспективных и революционных направлений их использования.

Увеличение экономической отдачи

Манипуляторы, или индустриальные роботы, все еще составляют наибольший процент работающих роботов в мире. Согласно *World Robotics 2017*, исследованию, проведенному International Federation of Robotics в конце 2016 года, в промышленности работало более 1 800 000 роботов. (Отчет доступен по адресу https://ifr.org/downloads/press/Executive_Summary_WR_2017_Industrial_Robots.) К 2020 году в результате быстрого развития автоматизации производства количество индустриальных роботов, вероятно, возрастет до 3 000 000. Фактически фабрики (как сущность) будут использовать роботы, чтобы стать более интеллектуальными; концепция выдвинута *Industry 4.0*. Благодаря широкому распространению и использованию Интернета, сенсоров, данных и роботов решения Industry 4.0 упрощают настройку и повышают качество продукции куда лучше, чем можно достичь без роботов. Роботы уже не только работают в опасных окружающих условиях, решая такие задачи, как сварка, сборка, покраска и упаковка; они работают быстрее, точнее и за более низкую плату, чем люди.

Забота о вас

С 1983 года роботы помогают хирургам в сложных операциях, обеспечивая точность и аккуратность разрезов, которой могут достичь только роботизированные манипуляторы. Помимо дистанционного управления ходом операции (устранение хирурга из рабочей зоны для повышения стерильности), совершенствование автоматизации операций открывает возможность полной автоматизации хирургических операций в ближайшем будущем, как упоминается в

статье https://www.huffingtonpost.com/entry/is-the-future-of-robotic-surgery-already-here_us_58e8d00fe4b0acd784ca589a.

Оказание услуг

Роботы оказывают и другие вспомогательные услуги как в частных, так и в публичных местах. Наиболее известен робот-пылесос Roomba, способный самостоятельно чистить пол дома (это настоящий бестселлер, его продажи превысили 3 миллиона экземпляров), но есть и другие сервисные роботы, также заслуживающие внимания.

- » **Доставщик.** Примером является робот — доставщик пиццы Domino (<https://www.bloomberg.com/news/articles/2017-03-29/domino-s-will-begin-using-robots-to-deliver-pizzas-in-europe>).
- » **Газонокосилка.** Роботизированных газонокосилок уже так много, что вы, вероятно, можете найти их в местном садоводческом магазине.
- » **Информация и развлечения.** Один из примеров — робот Pepper, находящийся в каждом Японском магазине SoftBank (<http://mashable.com/2016/01/27/softbank-pepper-robot-store/>).
- » **Забота о пожилых.** Примером такого робота является Hector, финансируемый Европейским союзом (<https://www.forbes.com/sites/jenniferhicks/2012/08/13/hector-robotic-assistance-for-the-elderly/#539ac7f82443>).

Вспомогательные роботы для пожилых людей далеки от предоставления общей помощи, оказываемой реальным медперсоналом. Роботы сосредотачиваются на критически важных задачах, таких как запоминание времени приема лекарств, помощь в перемещении пациентов от кровати до каталки, проверка физического состояния пациентов, объявление тревоги, если что-то пошло не так, или просто для компании. Например, терапевтический робот Paro обеспечивает биологическую терапию пожилым людям (https://www.huffingtonpost.com/the-conversation-global/robotrevolution-why-tech_b_14559396.html)³.

Риск в опасной окружающей среде

Роботы идут туда, куда люди пройти не могут или могут, но сильно рискуют. Некоторые роботы отправлены в космос (наиболее известны марсоходы NASA Opportunity и Curiosity) для дальнейших его исследований. (Роботы в космосе обсуждаются в главе 16.) Многие другие роботы остаются на Земле

³ См. лучше [https://ru.wikipedia.org/wiki/Папо_\(робот\)](https://ru.wikipedia.org/wiki/Папо_(робот)). — *Примеч. ред.*

и используются для подземных работ, таких как транспортировка руды в шахтах или создание карт туннелей и пещер. Подземные роботы исследуют даже системы коллекторов, как Luigi (имя как у знаменитого коллеги-водопроводчика из видеоигры). Робот Luigi предназначен для чистки коллекторов, он разработан лабораторией Senseable City Lab Массачусеттского технологического института. Он предназначен для исследования мест, куда люди не могут проникнуть из-за высоких концентраций химикатов, бактерий или вирусов (см. <https://money.cnn.com/2016/09/30/technology/mit-robots-sewers/index.html>).

Роботы используются даже там, где люди неизбежно умрут, например при авариях на атомных электростанциях в Три-Майл-Айленде, Чернобыле и Фукусиме. Эти роботы удаляют радиоактивные материалы и делают обстановку более безопасной. Высокие дозы радиации влияют даже на роботов, поскольку приводят к электронным помехам и искажениям сигналов, которые со временем наносят роботам ущерб. Только *радиационно-стойкая интегральная схема* позволяют роботам сопротивляться радиации в достаточной мере, чтобы решать задачи, подобные решаемым подводным роботом Little Sunfish, работавшим в одном из затопляемых реакторов Фукусимы, где произошло расплавление топлива (см. <http://www.bbc.com/news/in-pictures-40298569>).

Кроме того, опасные для жизни ситуации возникают в военной или криминальной обстановке, когда роботы используются для транспортировки оружия или обезвреживания бомб. Эти роботы способны также исследовать различные пакеты, в которых могут находиться и многие другие вредные вещи, помимо бомб. Такие модели роботов, как PackBot от iRobot (от той же компании, которая создала чистильщика Rumba) или Talon от QinetiQ North America, обезвреживают опасные взрывчатые вещества при дистанционном управлении, т.е. их действия дистанционно контролирует эксперт по взрывным устройствам. Некоторые роботы могут даже действовать вместо солдат или полицейских, решая задачи разведки или прямого вмешательства (например, полиция в Далласе использовала робот для ликвидации стрелка <https://edition.cnn.com/2016/07/09/opinions/dallas-robot-questions-singer/index.html>).



ЗАПОМНИ

Люди ожидают, что в будущем вооруженные силы будут все чаще использовать роботы. Если отбросить этические соображения о применении этого нового оружия, все сводится к старому вопросу об электронных пушках вместо масла (https://www.huffingtonpost.com/entry/guns-versus-butter-our-re_b_60150.html), а значит, нация вполне может обменять экономическую мощь на военную. Роботы кажутся более пригодными для этой модели, чем традиционное вооружение, нуждающееся в обученном персонале.

Использование страной роботизированных средства для немедленного создания эффективной армии роботов в любой момент — это именно то, что слишком хорошо продемонстрировали предистории *Звездных войн*.

Роль специальных роботов

К специальным роботам относят дроны и беспилотные автомобили. Дроны спорны из-за своего военного применения, но *беспилотные воздушные транспортные средства* (Unmanned Aerial Vehicle — UAV) используются также для мониторинга сельхозугодий и многих других мирных целей, обсуждающихся в главе 13.

Люди долго фантазировали об автомобилях, способных к самостоятельному вождению. Эти автомобили быстро превратились в действительность после достижений DARPA Grand Challenge. Большинство производителей автомобилей пришло к мнению, что производство и продажа беспилотных автомобилей может изменить фактическое экономическое равновесие в мире (а следовательно, порыв в создании рабочих экземпляров таких транспортных средства возможен в ближайшее время: <https://www.washingtonpost.com/news/innovations/wp/2017/11/20/robot-driven-ubers-without-a-human-driver-could-appear-as-early-as-2019/>). Более подробная информация о беспилотных автомобилях, их технологии и значении приведена в главе 14.

Сборка простого робота

Обзор роботов не будет полон без обсуждения его устройства с учетом современного состояния дел и того, как искусственный интеллект может улучшить свое функционирование. Основы роботостроения обсуждаются в следующих разделах.

Компоненты

Робот должен действовать в реальном мире, поэтому он нуждается в *исполнительных элементах* (effector) — ногах для хождения или колесах для обеспечения *возможности передвижения*. Он также нуждается в руках и щипцах, чтобы держать, вращать или обрабатывать предметы, обеспечивая, таким образом, *возможности манипулирования*. В разговоре о возможности робота сделать нечто можно также услышать такой термин, как *исполнительный механизм* (actuator), используемый как синоним термина “исполнительный элемент”. Исполнительный механизм — это один из механизмов, составляющих

исполнительные элементы и обеспечивающий одиночное движение. Таким образом, у ноги робота есть несколько разных исполнительных механизмов, таких как электромоторы или гидравлические цилиндры, выполняющих такие движения, как ориентация стоп или изгиб колена.

Действие в реальном мире требует определения его композиции и понимания, где именно робот находится. *Сенсоры* (sensor) обеспечивают исходные данные о происходящем вне робота. Такие устройства, как камеры, лазеры, гидролокаторы и сенсоры давления, измеряют параметры окружающей обстановки и информируют робота о происходящем, а также подсказывают его местоположение. Поэтому робот состоит главным образом из организованного набора сенсоров и исполнительных элементов. Все они разрабатываются для взаимодействия в архитектуре, которая и составляет робота. (Сенсоры и исполнительные элементы — это фактически механические и электронные части, которые можно использовать как автономные компоненты в разных случаях.)

Наиболее популярна внутренняя архитектура из параллельных процессов, собранных в уровни, специализирующиеся на решении задачи одного вида. Параллелизм важен. Как люди мы имеем единый процесс мышления и внимания; мы не должны думать о таких базовых функциях, как дыхание, сердцебиение и пищеварение, поскольку эти процессы выполняются параллельно процессу мышления. Зачастую мы можем даже выполнить несколько действий, идти или бежать, одновременно разговаривая или делая что-то еще (хотя в некоторых ситуациях это может оказаться опасным). То же самое происходит у роботов. Например, в архитектуре с тремя уровнями множество процессов робота собрано на трех уровнях, каждый из которых характеризуется своим временем ответа и сложностью ответа.

- » **Реактивный** (reactive). Непосредственно получает данные от сенсоров как канал для восприятия роботом мира и немедленно реагирует на внезапные проблемы (например, сразу сворачивать, чтобы не врезаться в неожиданную преграду).
- » **Управляющий** (executive). Обрабатывает ввод сенсоров, определяет положение робота в пространстве (важная функция локализации) и решает, какое действие выполнять с учетом требований предыдущего уровня (реактивного) и следующего (совещательного).
- » **Совещательный** (deliberative). Планирует выполнение таких задач, как переход из одного пункта в другой, а также вырабатывает последовательность действий по взятию объекта. Этот уровень преобразует поставленную роботу задачу в набор инструкций, исполняемых уровнем управления.

Другая популярная архитектура — это конвейерная архитектура, обычно используемая в беспилотных автомобилях. Она просто делит параллельные процессы робота на отдельные фазы, такие как считывание, восприятие (подразумевающее понимание считываемого), планирование и контроль.

Восприятие мира

Более подробно сенсоры рассматриваются в главе 14, при обсуждении их практического применения в беспилотных автомобилях. Существует множество видов сенсоров; одни сосредоточиваются на внешнем мире, другие — на самом роботе. Например, роботизированный манипулятор должен знать, насколько распространяется его рука и достанет ли он до предмета. Кроме того, некоторые сенсоры являются активными (они активно ищут информацию на основании решения робота), а другие — пассивными (они получают информацию постоянно). Каждый сенсор предоставляет электронный ввод, который робот может использовать немедленно или предварительно обработав, чтобы получить о нем представление.

Восприятие (perception) подразумевает создание локальной карты реальных объектов и определение расположения робота на более общей карте известного мира. Процесс комбинации данных от всех сенсоров, *сочетание датчиков* (sensor fusion), создает список основных фактов, используемых роботом. Машинное обучение в данном случае помогает выработать алгоритмы зрения, а использование глубокого обучения позволяет распознавать объекты и образы (как обсуждалось в главе 11). Он также объединяет все данные в осмысленное представление, используя алгоритмы неконтролируемого машинного обучения. Это задача *малоразмерного встраивания* (low-dimensional embedding), подразумевающего преобразование сложных данных от всех сенсоров в простую плоскую карту или другое представление. Расположение робота определяет *метод одновременной локализации и построения карты* (Simultaneous Localization And Mapping — SLAM), работающий точно так же, как когда вы смотрите на карту, чтобы понять, где вы находитесь.

Управление роботом

После предоставления сенсорами всей необходимой информации система планирования предоставляет роботу список действий, которые следует предпринять для достижения цели. Планирование осуществляется программно (с использованием экспертной системы, например, как описано в главе 3) или с помощью алгоритма машинного обучения, такого как байесовские сети, описанного в главе 10. Разработчики экспериментируют с применением обучения с подкреплением (машина учится методом проб и ошибок), но робот — не

малыш (он тоже учится ходить методом проб и ошибок); эксперименты могут доказать неэффективность метода, а кроме того, они требуют автоматического создания плана обучения, поскольку довольно дорогой робот может быть поврежден в процессе.

И наконец, планирование — это не просто вопрос интеллектуальных алгоритмов, поскольку, когда дело доходит до его исполнения, все обычно проходит не так, как планировалось. Давайте взглянем на эту проблему с человеческой точки зрения. Слепую вы не сможете идти прямо, если у вас не будет некоего ориентира для внесения постоянных поправок. Иначе вы будете ходить кругами. Ваши ноги являются исполнительными механизмами, но они не всегда отлично выполняют инструкции. Роботы стоят перед той же самой проблемой. Кроме того, роботы сталкиваются с такими проблемами, как задержки системы (технически *запаздыванием*), или робот выполняет инструкции не точно вовремя, в результате чего получаются смешные вещи. Но чаще всего возникает проблема с окружающей робота обстановкой.

» **Неопределенность.** Робот не знает, где он находится; он может наблюдать ситуацию только частично и не может понять ее точно. При неопределенности разработчики говорят, что робот работает в *стохастической среде* (stochastic environment).

» **Состязательная ситуация.** Люди или объекты находятся в движении. В некоторых ситуациях эти объекты могут быть опасны (см. <https://www.businessinsider.com/kids-attack-bully-robot-japanese-mall-danger-avoidance-ai-2015-8>). Это *проблема многоагентности* (multiagent problem).



ЗАПОМНИ

Роботы должны работать в окружающих средах, являющихся частично неизвестными, изменчивыми, а главное — непредсказуемыми, без постоянных путей. Все его действия связаны, и робот должен непрерывно обрабатывать поток информации и действовать в режиме реального времени. Нужна способность приспосабливаться к такому виду окружающей обстановки, который не может быть полностью предусмотрен или запрограммирован, а это требует способности к обучению, все чаще предоставляемой роботам алгоритмами искусственного интеллекта.



Глава 13

Полеты с дронами

В ЭТОЙ ГЛАВЕ...

- » Различие между военными и гражданскими дронами
- » Возможности использования дронов
- » Что искусственный интеллект может дать дронам
- » Правила применения дронов

Дроны (drone) — это мобильные роботы, способные перемещаться в окружающей среде по воздуху. Первоначально предназначенные для военных целей, дроны стали великолепной новинкой для досуга, исследований, коммерческой доставки и многого другого. Однако за всеми их разработками все еще скрываются военные цели, что вызывает беспокойство у многих экспертов по искусственному интеллекту и общественных деятелей, которые видят в них возможную непобедимую машину для убийства.

Люди летают с тех пор, как братья Райт совершили первый полет 17 декабря 1903 года (см. <https://www.nps.gov/wrbr/learn/historyculture/thefirst-flight.htm>). Однако летать люди хотели всегда, и легендарные мыслители, такие как Леонардо да Винчи в эпоху Возрождения (см. статью Смитсоновского музея: <https://airandspace.si.edu/stories/editorial/leonardo-da-vinci-and-flight>), всегда обращали внимание на эту задачу. Технология полета уже довольно совершенна, поэтому дроны являются более зрелыми, чем другие мобильные роботы, поскольку ключевая технология их работы хорошо известна. Ограничением дронов является применение искусственного интеллекта. Полет накладывает некоторые существенные ограничения на возможности

дронов, например вес, который они могут нести, или действия, которые они могут совершать, достигнув места назначения.

В этой главе обсуждаются текущее состояние дронов (потребительского, коммерческого и военного классов), а также роли, которые дроны могут играть в будущем. Частично эти роли зависят от интеграции с решениями для искусственного интеллекта, которые дадут им больше автономности и дополнительных возможностей при перемещении и работе.

На грани искусства

Дроны, мобильные роботы, способные летать, существуют довольно давно, в военной области (откуда и произошла эта технология). Официальное военное название таких летательных аппаратов — *беспилотный самолет* (Unmanned Aircraft System — UAS). На публике такие мобильные роботы больше известны как *дроны*, поскольку их звук напоминает жужжание трутня. Но вы не встретите этот термин в официальных публикациях, поскольку чиновники предпочитают такие термины, как “UAS”, или *беспилотные воздушные боевые средства* (Unmanned Aerial Combat Vehicle — UACV), или *беспилотные воздушные средства* (Unmanned Aerial Vehicle — UAV), или даже *дистанционно пилотируемые самолеты* (Remotely Piloted Aircraft — RPA).



СОВЕТ

Названий много. Следующая статья *ABC News* поможет понять наиболее популярные аббревиатуры и официальные названия дронов: <http://www.abc.net.au/news/2013-03-01/drone-wars-the-definition-dogfight/4546598>.

Беспилотный вылет на задание

Напоминающие обычный самолет (но меньшего размера) военные дроны летают на крыльях, т.е. имеют крылья, а также один или несколько пропеллеров (или реактивных двигателей) и в некоторой степени не очень отличаются от самолетов, на которых путешествуют гражданские лица. Сейчас используются военные дроны шестого поколения, как описано в статье <https://www.military.com/dailynews/2015/06/17/navy-air-force-to-develop-sixth-generation-unmanned-fighter.html>. Военные дроны являются беспилотными и дистанционно контролируются по спутниковой связи даже с другой стороны земли. Военные операторы дронов получают от контролируемых дронов телеметрическую информацию, преобразуемую в изображение. Операторы могут использовать эту информацию при управлении аппаратом, которому передаются определенные команды. Одни военные дроны выполняют задачи наблюдения и разведки, поэтому несут только камеры и другие устройства сбора

информации. Другие вооружены и могут атаковать цели. Некоторые из наиболее оснащенных моделей сравнимы по возможностям с пилотируемыми боевыми самолетами (см. <https://www.military.com/defensetech/2014/11/20/navy-plans-for-fighter-to-replace-the-fa-18-hornet-in-2030s>) и способны достичь даже тех мест на планете, которых обычный пилот достичь не может (<https://spacenews.com/u-s-military-gets-taste-of-new-satellite-technology-for-unmanned-aircraft/>).

История военных дронов длинна. Хотя ее начало и является темой для дебатов, Королевский флот начал использовать подобные дрону аппараты для учебных стрельб в 1930-х годах (см. <https://dronewars.net/2014/10/06/rise-of-the-reapers-a-brief-history-of-drones/>). В США практически настоящие дроны регулярно использовались в качестве мишеней уже с 1945 года (см. <http://www.designation-systems.net/dusrm/m-33.html>). С 1971 года исследователи начали применять любительские дроны в военных целях. Джон Стюарт Фостер (John Stuart Foster), физик-ядерщик, работавший на американское правительство, увлекался моделями самолетов и предположил идею ставить на них оружие. Это привело к разработке двух прототипов DARPA в 1973 году, но применение подобных дронов Израилем в ближневосточных конфликтах прошлого десятилетия вызвало интерес и потребовало дальнейшей разработки военных дронов. Довольно интересно, что в 1973 году военные впервые сбили дрон, используя лазер (см. статьи Popular Science в <https://www.popsci.com/laser-guns-are-targeting-uavs-but-drones-are-fighting-back> и Popular Mechanics в <https://www.popularmechanics.com/military/research/a22627/drone-laser-shot-down-1973/>). Первое убийство с помощью дрона произошло в 2001 году в Афганистане (см. <https://www.theatlantic.com/international/archive/2015/05/america-first-drone-strike-afghanistan/394463/>). Конечно, на курок нажимал человек на другом конце канала связи.

Люди обсуждают, предоставлять ли военным дронам возможности искусственного интеллекта. Некоторые полагают, что наличие у дронов собственного процесса принятия решений могло бы привести, при неисправности, к уничтожению людей. Однако возможности искусственного интеллекта могли бы также позволить дронам лучше избегать повреждений или выполнять другие задачи, не связанные с убийством, как сегодня искусственный интеллект помогает водить автомобили. Он мог бы помогать пилотам стабилизировать полет при плохой погоде, как роботизированная система da Vinci помогает хирургам (см. раздел “Помощь хирургу” главы 7). Сейчас военные дроны с возможностями убийства также спорны, поскольку искусственный интеллект может сделать военные действия просто абстракцией с дальнейшей дегуманизацией, сводя их к передаче дронами образов операторам и отдаче ими команд

дистанционно. Да, решение об убийстве все еще принимал бы оператор, но фактическое убийство осуществлял бы дрон, ограждая оператора от ответственности за это.



СОВЕТ

Обсуждение военных дронов существенно для этой главы, поскольку они неотделимы от разработки гражданских дронов, и это влияет на обсуждение данной технологии в обществе. Кроме того, предоставление военным дронам полной автономии дает аргументацию статьям об апокалипсисе искусственного интеллекта уже вне области научной фантастики и несколько беспокоит общество. Более подробный технический обзор моделей и их возможностей приведен в статье Deutsche Welle <https://www.dw.com/en/a-guide-to-military-drones/a-39441185>.

Знакомство с квадрокоптером

Многие люди сначала слышали о дронах-квадрокоптерах потребительского и любительского классов и только затем — о коммерческих дронах-квадрокоптерах (таких, как используемый компанией Amazon и обсуждаемый в статье <https://www.amazon.com/Amazon-Prime-Air/b?node=8037720011>). Большинство современных военных дронов не имеет вида вертолета, но вы можете найти некоторые и из таких, как дрон TIKAD от Университета Дьюка. Он описан в статье <https://www.defenseone.com/technology/2017/07/israeli-military-buying-copter-drones-machine-guns/139199/> и продемонстрирован на https://www.youtube.com/watch?v=VaTW8uAo_6s. Военные дроны вертолетного типа фактически были сначала любительскими прототипами (см. <https://www.popularmechanics.com/military/research/news/a27754/hobby-drone-sniper/>).

Однако неотъемлемой частью всех этих работ являются мобильные телефоны. По мере того как мобильные телефоны становились все меньше, их батареи становились все легче. Мобильные телефоны имеют также миниатюрные камеры и беспроводное подключение — все, что необходимо для современного дрона. Несколько десятилетий назад у маленьких дронов был целый набор ограничений.

- » Они были радиоуправляемыми с большими наборами команд.
- » Они нуждались в прямой видимости (или вы будете лететь вслепую).
- » Это были маленькие самолеты с жестким крылом (без возможности зависания).
- » Они имели шумные дизельные или бензиновые двигатели (ограничивающие их диапазон и простоту применения).

Новые легкие литиевые батареи обеспечили дронам следующее.

- » Меньшие, более тихие и надежные электромоторы.
- » Беспроводное дистанционное управление.
- » Видеосигнал от дрона (прямая видимость больше не нужна).

Сейчас дроны обладают также датчиком GPS, акселерометрами и гироскопами — все из комплектующих обычных мобильных телефонов. Эти средства позволяют контролировать позицию, высоту и ориентацию, что для телефонных приложений весьма полезно, а для полета дрона крайне важно.

Благодаря всем этим усовершенствованиям дроны изменились от подобия моделям самолетов с фиксированным крылом до некоего подобия вертолетам, но с несколькими роторами, чтобы поднять себя в воздух и выбрать направление. Использование нескольких роторов предоставляет ряд преимуществ. В отличие от вертолетов, дроны не нуждаются в роторах с изменяемым шагом для ориентации. Роторы с переменным шагом дороже и их сложнее контролировать. Вместо этого дроны используют простые пропеллеры фиксированного шага, способные совместно выполнять те же самые функции роторов переменного шага. Следовательно, дроны теперь многороторные: трикоптеры, квадрокоптеры, гексакоптеры и октокоптеры с 3, 4, 6 и 8 роторами соответственно. Из всех возможных конфигураций квадрокоптер стал наиболее популярной конфигурацией дрона для коммерческого и гражданского использования. Имея четыре небольших ротора, вращающихся навстречу друг другу, оператор легко может повернуть и переместить дрон, изменив скорость вращения разных роторов, как показано на рис. 13.1.

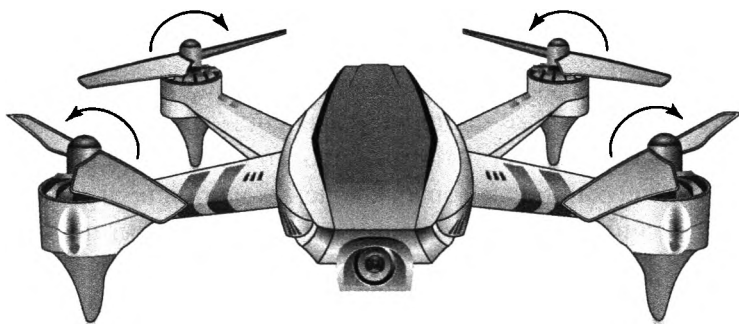


Рис. 13.1. Квадрокоптер летит при противовращении его роторов в указанных направлениях

Области применения дронов

У каждого типа дрона есть различные возможности для использования искусственного интеллекта. У больших и маленьких военных дронов уже есть свои параллельные наработки в смысле технологии, и эти дроны, вероятно, будут использованы для разведки и наблюдения. Эксперты предсказывают, что военные дроны будут отличаться от личных и коммерческих, поскольку в последних используются технологии, отличные от военных. (Некоторые сходства существуют, например TIKAD от Университета Дьюка фактически начал использоваться как потребительский дрон.)

Небольшими дешевыми и легко настраиваемыми дронами интересуются не только повстанцы и террористические группы (см. пример <http://www.popularmechanics.com/military/weapons/a18577/isi-spacing-drones-with-explosives/>), но и правительства для военных действий в условиях города и внутри помещений. Помещения, такие как коридоры или комнаты, существенно ограничивают боевые возможности таких военных дронов, как Predator или Reaper, имеющих размер самолета (если только не нужно снести все здание). То же самое относится к дронам-разведчикам, таким как Raven и Puma, поскольку они созданы для действий на открытом поле боя, а не для уличных боев. (Более подробный анализ возможного военного развития ранее безоружных потребительских дронов приведен в статье от *Wired* <https://www.wired.com/2017/01/military-may-soon-buy-drones-home/>).

Коммерческие дроны далеки от того, чтобы быть использованными на поле боя сразу после покупки в магазине, хотя и предоставляют подходящую платформу. Важнейшая причина интереса военных к использованию коммерческих дронов в том, что они общедоступны и недороги по сравнению со стандартным вооружением, что делает их как разовыми в применении, так и применимыми во многих ситуациях. Несложные для взлома и модификации, они требуют большей защиты, чем их уже защищенные военные аналоги (их каналы связи и элементы управления могут быть блокированы электронными средствами), и нуждаются в интеграции некоторых ключевых частей программного обеспечения и оборудования, прежде чем они могут быть эффективно применены в любом задании.

Для навигации в закрытом пространстве необходимы улучшенные способности избегать столкновений и контролировать направления без GPS (сигналы которого в здании обычно экранируются), а также избегать привлечения внимания потенциального противника. Кроме того, для разведки дроны нуждались бы в способностях планирования (определения засад и угроз) и выявления целей. Такие дополнительные характеристики отсутствуют у существующих коммерческих технологий и потребовали бы для решения искусственного

интеллекта, разработанного специально для этой цели. Военные исследователи активно работают над необходимыми дополнениями, чтобы получить военные преимущества. Недавние разработки в сложных сетях глубокого обучения, устанавливаемых на стандартных мобильных телефонах, такие как YOLO (<https://pjreddie.com/darknet/yolo/>) или MobileNets от Google (<https://ai.googleblog.com/2017/06/mobilenets-open-source-models-for.html>), доказывают, что установка сложного искусственного интеллекта на маленький дрон вполне возможна уже при существующих технологических достижениях.

Дроны на невоенных ролях

В настоящее время у коммерческих дронов нет многих возможностей, которыми обладают военные модели. Коммерческий дрон может лишь сделать ваш снимок и снимок вашего участка с воздуха. Но даже коммерческие дроны находят применение сейчас и будут находить в ближайшем будущем.

- » Своевременная доставка товаров независимо от транспортных за-
торов (разрабатывается Google X, Amazon и многими другими).
- » Мониторинг при обслуживании и руководстве проектом.
- » Оценка ущерба в страховании.
- » Создание карт полей и подсчет поголовья стад для фермеров.
- » Помощь при поисково-спасательных операциях.
- » Поддержка подключения к Интернету в удаленных, неподключен-
ных областях (идея разрабатывается Facebook).
- » Производство электричества с использованием высотных ветров.
- » Перемещение людей из одного места в другое.

Доставка товаров дроном недавно поразила воображение публики благодаря усилиям больших компаний. Одним из первых и наиболее известных примеров является компания Amazon (пообещавшая, что служба Amazon Prime Air вскоре начнет работу: <https://www.amazon.com/Amazon-Prime-Air/b?node=8037720011>). Компания Google обещает подобную службу, Project Wing (<https://www.businessinsider.com/project-wing-update-future-google-drone-delivery-project-2017-6?IR=T>). Однако мы можем еще на протяжении многих лет ожидать запуска реальной масштабной службы воздушной доставки на базе дронов.



ЗАПОМНИ

Хотя идея была в том, чтобы убрать посредников из цепочки логистики наиболее выгодным способом, остается еще урегулировать некоторые неоднозначности и решить множество технических проблем. Средствам массовой информации продемонстрировали

дрон, успешно доставляющий небольшие пакеты и предметы, такие как пицца или бургеры, по назначению экспериментальным способом (<https://www.theverge.com/2017/10/16/16486208/alphabet-google-project-wing-drone-delivery-testing-australia>), но правда в том, что дроны не могут летать далеко и нести большой вес. Одна из самых больших проблем — регулирование полетов целых роев дронов, каждый из которых должен доставить что-то из одного пункта в другой. Есть и вполне очевидные проблемы, например уклонение от таких препятствий, как линии электропередачи, здания и другие дроны, устойчивость к плохой погоде и поиск подходящего места посадки. Дроны должны также избегать определенных воздушных пространств и отвечать всем требованиям безопасности полетов, предъявляемым к самолетам. Искусственный интеллект станет ключом к решению большинства этих проблем, но не всех. В настоящее время доставка дронами кажется прекрасной, пока ее масштабы не велики, а задачи куда важнее доставки бургера вам на дом (пример приведен в статье <http://time.com/rwanda-drones-zipline/>).

Дроны могут стать вашими глазами, позволяя видеть в ситуациях, которые слишком опасны, затруднительны или дороги. Дистанционно контролируемые или полуавтономные (с использованием решений искусственного интеллекта для обнаружения образов и обработки данных сенсоров) дроны могут осматривать местность и оказывать поддержку в выживании, т.е. найти и спасти, поскольку они способны осмотреть любую инфраструктуру сверху и позволить оператору-человеку принять соответствующие меры. Например, дроны успешно инспектируют линии электропередачи, трубопроводы (<https://www.wsj.com/articles/utilities-turn-to-drones-toinspect-power-lines-and-pipelines-1430881491>) и железнодорожные пути (<http://fortune.com/2015/05/29/bnsf-drone-program/>), обеспечивая более частый и менее дорогой мониторинг жизненно важных, но не легкодоступных инфраструктур. Даже страховые компании посчитали их полезными для оценки ущерба (<https://www.wsj.com/articles/insurers-are-set-to-use-drones-to-assess-harveys-property-damage-1504115552>).

Полиция и службы быстрого реагирования по всему миру нашли дроны полезными для множества действий, от поисково-спасательных операций до поиска лесных пожаров, а также от пограничного патрулирования до мониторинга скопления масс. Полицейские находят более новые способы использования дронов (<http://www.foxnews.com/tech/2017/07/19/drones-become-newest-crime-fighting-tool-for-police.html>), включая поиск нарушителей дорожного движения (см. статью https://www.motorauthority.com/news/1014233_drone-police).

com/news/1113849_french-police-using-drones-to-catch-drivers-breaking-the-law).

Сельское хозяйство — еще одна важная область применения дронов. Мало того что они могут контролировать посевы зерновых, сообщая об их состоянии и выявляя проблемы, они также способны применять пестициды или удобрения только там, где это на самом деле необходимо, как описано в обзоре Мас-сачусеттского технологического института (<https://www.technologyreview.com/s/526491/agricultural-drones/>). Дроны предоставляют образы, которые подробнее и дешевле получаемых от орбитального спутника, и применимы для следующего.

- » Картографирование почв, а также анализ полученных образов и результатов трехмерного лазерного сканирования для повышения эффективности планирования посевов.
- » Контроль движения тракторов.
- » Контроль созревания урожая в реальном времени.
- » Распыление химикатов, когда и где необходимо.
- » Ирригация, когда и где необходимо.
- » Оценка состояния урожая с использованием инфракрасного зрения, чего сам фермер сделать не может.

Точное земледелие (precision agriculture) подразумевает использование возможностей искусственного интеллекта для управления движением, локализации, зрения и обнаружения. Точное земледелие способно повысить продуктивность сельского хозяйства (более здоровые зерновые культуры — больше пищи для всех) при снижении затрат на него (нет необходимости распылять пестициды повсюду).

Дроны способны и на большие подвиги. Идея в том, чтобы переместить существующую инфраструктуру в небо, используя дроны. Например, компания Facebook намеревается обеспечить соединение с Интернетом (<https://www.theguardian.com/technology/2017/jul/02/facebook-drone-aquila-internet-test-flight-arizona>) там, куда кабель связи недотянули или где повредили, используя специальные дроны Aquila (<https://www.facebook.com/notes/mark-zuckerberg/the-technology-behind-aquila/10153916136506634/>). Планируется также использовать дроны для транспортировки людей, заменив такие популярные транспортные средства, как автомобиль (<http://www.bbc.com/news/technology-41399406>). Еще одна возможность — производство электроэнергии в верхних слоях атмосферы, где дуют постоянные сильные ветры и никто не жалуется на шум от ротора (<https://www.bloomberg.com/news/articles/2017-04-11/flying-drones-that-generate-power-from-wind-get-backing-from-eon>).

Усиление дронов искусственным интеллектом

Независимо от области применения дрона (потребительской, коммерческой или военной) искусственный интеллект играет для него важнейшую роль и является решающим фактором. Улучшив автономность и координацию, искусственный интеллект делает возможными новые области применения и повышает эффективность в старых. Раффаэлло Д'Андреа (Raffaello D'Andrea), канадско-итальянско-швейцарский инженер, профессор динамических систем и управления в ЕТН Цюриха, а также изобретатель дронов, продемонстрировал свои достижения на следующем видео: <https://www.youtube.com/watch?v=RCXGpEmFbOw>. Видео демонстрирует, что дроны могут стать более автономными при использовании алгоритмов искусственного интеллекта. *Автономность* влияет на то, как дрон летает, снижая роль отдающих команды людей, ведь дрон автоматически обнаруживает препятствия и реагирует на них, обеспечивая безопасную навигацию в сложных условиях. *Координация* подразумевает способность дронов взаимодействовать без обмена данными с центром управления и получения от него инструкций, дроны могут обмениваться информацией между собой и взаимодействовать в реальном времени, чтобы выполнить свою задачу.

Полная автономность может даже исключить человека из управления дроном, чтобы он мог сам проложить маршрут и сам выполнить поставленную задачу. (Люди отдают только высокоуровневые распоряжения.) Если дрон не управляется пилотом, он полагается на GPS, чтобы выбрать оптимальный путь к точке назначения, но это возможно только на открытом воздухе и маршрут не всегда точен. Применение внутри зданий повышает требования к точности полета, а значит, требует использования других сенсоров, позволяющих дрону контролировать *ближайшее окружение* (такие элементы здания, как выступы стен, в которые он может врезаться). Самым дешевым и легким из этих сенсоров является камера, которую большинство коммерческих дронов использует как стандартное устройство. Однако наличия камеры не достаточно, поскольку еще нужны мастерство в обработке образов, использовании компьютерного зрения и методов глубокого обучения (обсуждаемого в главе 11 в контексте сверточных сетей).

Компании ожидают от коммерческих дронов автономного решения задач, например, чтобы они были способны доставить пакет со склада клиенту и справиться со всеми проблемами в пути. (Как и с роботами, всегда что-то идет не так, как надо, чтобы принимать решения на месте, устройство должно обладать искусственным интеллектом). Исследователи из лаборатории реактивного движения НАСА в Пасадене, Калифорния, недавно сравнили полет автоматизированного дрона и дрона, управляемого профессиональным пилотом дронов (см. <https://www.nasa.gov/feature/jpl/drone-race-human-versus-artificial-intelligence>). Интересно, что пилот-человек побеждал в этом

соревновании, пока не устал, и с этого момента более медленный, но более стабильный и менее подверженный ошибкам дрон его догнал. В будущем того же самого можно ожидать в шахматах и игре Го: автоматизированные дроны опередят людей, как пилота дронов, и в летных навыках, и в выносливости.

Полная координация позволила бы сотням, если не тысячам, дронов летать вместе. Такая возможность могла бы иметь смысл для коммерческих и потребительских дронов, когда дроны переполняют небеса. Использование координации было бы полезным с точки зрения избегания столкновений, совместного использования информации на препятствиях и анализа трафика способом, подобным используемому частично или полностью автоматизированными автомобилями (искусственный интеллект, управляющий автомобилем, обсуждается в главе 14).

Существующие алгоритмы дронов уже пересматриваются, и уже существуют некоторые решения по координации действия дронов. Например, Массачусетский технологический институт недавно разработал алгоритм децентрализованной координации дронов (см. <https://techcrunch.com/2016/04/22/mit-creates-a-control-algorithm-for-drone-swarms/>). Однако большинство исследований остаются незамеченными, поскольку координация дронов имеет возможное военное применение. Рой дронов может быть куда эффективнее при прорыве обороны противника и нанесении существенного урона, предотвратить который будет куда сложнее. У врага больше не будет цели в виде одного большого дрона, а будут сотни маленьких. Есть также и решения по устранению подобных угроз (см. <http://www.popularmechanics.com/military/weapons/a23881/the-army-is-testing-a-real-life-phaser-weapon/>). Недавняя проверка на рое из ста дронов (специальной модели Perdix, изготовленной по заказу Министерства обороны Соединенных Штатов), выпущенных из трех F/A-18 Super Hornet и выполняющих разведывательную миссию, были перехвачены (<https://www.technologyreview.com/s/603337/a-100-drone-swarmd-ropped-from-jets-plans-its-own-moves/>). Другие страны также включаются в эту новую гонку вооружений.



ЗАПОМНИ

Когда Илон Маск, соучредитель Apple Стив Возняк, физик Стивен Хокинг и многие другие известные общественные деятели и исследователи искусственного интеллекта подняли тревогу по поводу недавних разработок вооружения с искусственным интеллектом, они имели в виду не таких роботов, как в фильмах *Терминатор* или *Я, робот*, а скорее, вооруженные летающие дроны и другое автоматизированное оружие. Автономное оружие может начать гонку вооружений и навсегда изменить лик войны. Больше об этой теме можно прочитать в статье <https://mashable.com/2017/08/20/ai-weapons-ban-open-letter-un/>.

По большей части в этой книге описано создание обстановки и обеспечение данных для обучения искусственного интеллекта. Кроме того, вы уделили много времени рассмотрению того, что возможно, а что нет для искусственного интеллекта, исключительно с точки зрения обучения. В некоторых частях книги рассматриваются даже мораль и этика применительно к искусственному интеллекту и его пользователям-людям. Однако *ориентация обучения* (teaching orientation) искусственного интеллекта также важна.

В фильме *Военные игры* (<https://www.amazon.com/exec/obidos/ASIN/B0089J2818/datacservip0f-20/>) военный компьютер WOPR (War Operation Plan Response) содержит сильный искусственный интеллект, способный выработать наилучший комплекс действий в ответ на угрозу. В начале фильма компьютер WOPR является просто советником ответственных политиков. Затем, после взлома хакером, он начинает играть в игру "Термоядерная война". К сожалению, WOPR предполагает, что все игры реальны, и фактически начинает выполнять план термоядерной войны с Советским Союзом. Фильм, кажется, находится на грани подтверждения всех когда-либо существовавших страхов относительно искусственного интеллекта и войны.

Вот странная часть этого фильма. Хакер найден и работает теперь на хороших парней, он разрабатывает метод объяснить искусственному интеллекту тщетность его действий. Таким образом, искусственный интеллект попадает в обстановку, в которой он узнает, что победа в некоторых играх ("крестики-нолики" в данном случае) невозможна. Независимо от того, как хорошо каждый начнет, игра, в конце концов, закончится безвыходным положением. Затем искусственный интеллект проверяет это новое знание на термоядерной войне. В конце концов, искусственный интеллект приходит к выводу, что единственный способ победить — это не играть вообще.

В большинстве статей в средствах массовой информации, в научной фантастике и в кино никогда не рассматривается окружающая обстановка обучения. Но все же окружающая обстановка обучения — это основная часть уравнения, поскольку то, как вы настроите окружающую обстановку, определяет то, что изучит искусственный интеллект. В контексте военного оборудования, вероятно, имеет смысл не только учить искусственный интеллект побеждать, но и давать ему представление о том, что в некоторых сценариях победителей просто не будет, поэтому наилучший выход — не играть вообще.

Проблемы регулирования

Вполне очевидно, что дроны — не первые и не единственные вещи, летающие над облаками. Десятилетия коммерческих и военных полетов переполнили небеса, требуя строгого регулирования и человеческого контроля воздушного трафика, чтобы гарантировать безопасность. В США полномочной организацией по регулированию всей гражданской авиации, принимающей решения об управлении воздушным трафиком и аэропортами, является *Федеральное управление гражданской авиации* (Federal Aviation Administration — FAA). Управление FAA выработало серию правил для UAS (дронов), и вы можете прочитать их по адресу https://www.faa.gov/uas/resources/uas_regulations_policy/.

Управление FAA выработало набор правил, *Part 107*, в августе 2016 года. Этот свод правил ограничивает использование коммерческих дронов только дневным временем суток. Полный список правил доступен по адресу https://www.faa.gov/news/fact_sheets/news_story.cfm?newsId=20516. Правила сводятся к следующим пяти простым пунктам.

- » Летать на высотах ниже 400 футов (120 м).
- » Летать не быстрее 100 миль в час (161 км/ч).
- » Всегда находиться в зоне видимости.
- » У оператора должна быть соответствующая лицензия.
- » Никогда не летать возле пилотируемых самолетов, особенно возле аэропортов.
- » Никогда не пролетать над группами людей, стадионами или местами проведения спортивных мероприятий.
- » Никогда не летать в местах работы аварийно-спасательных служб.

Вскоре FAA издаст правила для полетов дронов ночью, когда они могут быть вне зоны видимости и в городских условиях, хотя и сейчас уже можно получить специальные формуляры от FAA. Цель таких регулирующих правил состоит в обеспечении общественной безопасности в условиях, когда влияние дронов на нашу жизнь все еще не ясно. Эти правила не сдерживают также технологический и экономический рост данной технологии.



ЗАПОМНИ

В настоящий момент многие страны вырабатывают правила регулирования полетов дронов. Эти инструкции гарантируют безопасность и увеличивают области применения дронов в экономических целях. Например, во Франции закон позволяет использование дронов в сельском хозяйстве с весьма немногими ограничениями, что позиционирует эту страну среди пионеров в их применении.

Сейчас отсутствие искусственного интеллекта означает, что дроны легко могут потерять связь и повести себя беспорядочно, иногда принося ущерб (см. <https://www.theatlantic.com/technology/archive/2017/03/drones-invisible-fence-president/518361/>). Даже при том, что у некоторых из них есть меры по обеспечению безопасности в случае потери связи с контроллером, такие как функция автоматического возвращения в точку запуска, FAA ограничивает их применение областью прямой видимости их оператора.

Другая важная мера безопасности — *геозонирование* (geo-fencing). У дронов, использующих службу GPS для локализации, есть программное обеспечение, ограничивающее их доступ к предопределенным периметрам (описанным координатами GPS), таким как аэропорты, военные зоны и другие области национальных интересов. Вы можете получить список параметров по адресу <http://tfr.faa.gov/tfr2/list.html> или прочитать больше на эту тему по адресу <https://www.theatlantic.com/technology/archive/2017/03/drone-invisible-fence-president/518361/>.

Алгоритмы и искусственный интеллект приходят на помощь при разработке подходящего технологического решения по безопасному применению дронов, доставляющих товары в городах. Исследовательский центр НАСА Ames Research Center работает над системой Unmanned Aerial Systems Traffic Management (UTM), предназначенной играть роль авиадиспетчерской службы для дронов, подобной используемой для пилотируемых самолетов. Эта система полностью автоматизирована; она учитывает возможности дронов общаться между собой. Система UTM помогает идентифицировать дроны в небе (у каждого будет код идентификации, подобно номерам автомобилей), а также задает маршрут и высоту полета каждого дрона, позволяя избегать возможных столкновений, последствий потери управления или потенциального ущерба гражданам. Система UTM может быть передана FAA для возможного применения или дальнейших разработок в 2019 году или позже. Веб-сайт НАСА предоставляет дополнительную информацию об этой революционной системе управления для дронов, которая может сделать коммерческое применение дронов реальным и безопасным: <https://utm.arc.nasa.gov/>.



ЗАПОМНИ

На случай, если ограничений недостаточно и незаконные дроны будут представлять угрозу, у полиции и военных есть несколько эффективных контрмер: поражение дрона дробовиком, его захват в сети, заклинивание его элементов управления и поражение лазером, микроволнами либо даже управляемыми ракетами.



Глава 14

Автомобиль, управляемый искусственным интеллектом

В ЭТОЙ ГЛАВЕ...

- » Пути к автономии беспилотного автомобиля
- » Будущее в мире беспилотных автомобилей
- » Цикл "восприятие–планирование–действие"
- » Применение и объединение различных сенсоров

Беспилотный автомобиль (SD car) является автономным транспортным средством, способным проехать от отправной точки до места назначения без человеческого вмешательства. Автономность подразумевает не просто автоматизацию выполнения некоторых задач (таких, как система активной помощи при парковке <https://www.youtube.com/watch?v=xW-MhoLImqg>); это хороший шаг в правильном направлении. Беспилотный автомобиль выполняет все необходимые задачи самостоятельно, человек может только наблюдать. Поскольку история беспилотных автомобилей насчитывает более ста лет (да, невероятно, но это так), эта глава начинается с их краткой истории.



ЗАПОМНИ

Для успеха технология должна предоставлять некое преимущество, в котором люди видят необходимость и которого не так легко добиться другими методами. Именно поэтому беспилотные автомобили настолько захватывающи. Они обеспечивают не только вождение. В следующем разделе речь пойдет о том, как беспилотные автомобили изменяют движение весьма существенно и помогут понять, почему эта технология столь неотразима.

Когда беспилотные автомобили распространятся немного больше и мир примет их как нечто повседневное, они все равно продолжают влиять на общество. Следующая часть главы поможет вам понять эти проблемы и их важность. Вы узнаете, на что это будет похоже и как сесть в беспилотный автомобиль и предположить, что он доставит вас из одного места в другое без проблем.

И наконец, беспилотные автомобили требуют работы сенсоров многих типов. Да, в некотором смысле эти сенсоры можно отнести к тем, которые видят, слышат и ощущают, но это было бы слишком большим упрощением. Заключительный раздел главы поможет понять функции различных сенсоров беспилотного автомобиля и что они дают беспилотному автомобилю в целом.

Краткая история

Разработка самоуправляемых автомобилей долго была частью футуристического предвидения, поддерживаемого научной фантастикой и фильмами начиная с первых экспериментов над радиоуправляемыми автомобилями в 1920-х годах. Об этом факте в истории автономных автомобилей можно прочитать в статье <https://qz.com/814019/driverless-cars-are-100-years-old/>. Проблема этих первых транспортных средств была в том, что они не были практичны; кто-то должен был следовать за ними, чтобы управлять ими по радио. Следовательно, несмотря на давние мечты о беспилотных автомобилях у существующих проектов не много общего с прошлым, кроме самого видения автономности.

Современные беспилотные автомобили серьезно улучшились в проектах, начатых в 1980-е годы (<https://www.technologyreview.com/s/602822/in-the-1980s-the-self-driving-van-was-born/>). В этих новых проектах предполагалось использовать искусственный интеллект для устранения необходимости в радиоуправлении, как в прежних проектах. Эти усилия финансируют многие университеты и военные (особенно армия США). Когда-то задача заключалась в том, чтобы выиграть конкурс DARPA Grand Challenge (<http://archive.darpa.mil/grandchallenge/>), который закончился в 2007 году. Однако сейчас множество стимулов для продолжения разработок инженерам предоставляют уже военные и коммерческие задачи.

Поворотным моментом конкурса стало создание Stanley, разработанного ученым и предпринимателем Себастьяном Труном (Sebastian Thrun) и его группой. Они выиграли конкурс DARPA Grand Challenge в 2005 году (см. видео <https://www.youtube.com/watch?v=LZ3bbHTsOL4>). После победы Трун начал разработку беспилотных автомобилей в компании Google. Сегодня вы можете видеть Stanley на выставке в Смитсоновском музее.



ЗАПОМНИ!

Автономные транспортные средства нужны не только военным. Довольно давно автомобильная промышленность пострадала от перепроизводства, поскольку было произведено больше автомобилей, чем было необходимо на рынке. Рыночный спрос снижается в результате всякого рода давлений, таких как долговечность автомобилей. В 1930-х годах долговечность автомобилей составляла в среднем 6,75 года, но сегодня она составляет в среднем 10,8 и более лет, позволяя водителям проходить по 250 и более тысяч миль (402 336 км). Снижение уровня продаж вынудило некоторых производителей уйти из отрасли или слиться в большие компании. Беспилотные автомобили — серебряная пуля для индустрии — способ благоприятно изменить рыночный спрос, повысив потребление. Эта полезная технология приведет к увеличению количества рабочих мест и появлению новых транспортных средств.

Будущее мобильности

Беспилотные автомобили — это прорыв не только потому, что они радикально изменят восприятие людьми автомобилей, а также и потому, что их введение в эксплуатацию окажет существенное влияние на общество, экономику и урбанизацию. В настоящее время никаких беспилотных автомобилей на дорогах еще нет — это только модели. (Вы можете полагать, что беспилотные автомобили — это уже коммерческая действительность, но факт в том, что все они — только прототипы. Ознакомьтесь, например, со статьей <https://www.wired.com/story/uber-self-driving-cars-pittsburgh/>, и вы найдете такие фразы, как *пилотный проект*, под которыми следует понимать прототипы, которые на настоящий момент еще не готовы.¹) Многие полагают, что введение беспилотного автомобиля потребует по крайней мере еще пары десятилетий, а замена всех существующих автомобилей беспилотными отнимет значительно больше времени. Но даже если беспилотные автомобили все еще в будущем, от них можно уверенно ожидать великих вещей, как описано в следующих разделах.

¹ <https://www.bbc.com/russian/news-37187602>. — *Примеч. ред.*

Восхождение на шестой уровень автономности

Предсказать формы будущего невозможно, но многие люди по крайней мере задумываются о характеристиках беспилотных автомобилей. Для ясности группа SAE International (<http://www.sae.org/>), орган стандартизации автомобилей, опубликовала классификационный норматив для автономных автомобилей (см. стандарт J3016 по адресу https://www.smmmt.co.uk/wp-content/uploads/sites/2/automated_driving.pdf). Наличие стандарта является серьезной вехой автоматизации автомобилей. Вот пять уровней автономности, определяемых стандартом SAE.

- » **Уровень 1. Помощь водителю.** Управление все еще находится в руках водителя, но автомобиль может выполнять простые вспомогательные действия, такие как контроль скорости. Этот уровень автоматизации подразумевает систему автоматического регулирования скорости, когда вы задаете для автомобиля предел скорости, и он стабилизирует ее, притормаживая.
- » **Уровень 2. Частичная автоматизация.** Автомобиль может чаще действовать вместо водителя, справляясь с ускорением, торможением и рулевым управлением, если нужно. Задача водителя — не терять бдительности и обеспечить контроль над автомобилем. Пример частичной автоматизации — автоматическое торможение, осуществляемое автомобилями некоторых моделей, если они определяют возможность столкновения (пешеходный переход или другой внезапно остановившийся автомобиль). Еще пример — адаптивная система автоматического регулирования скорости (которая не только контролирует скорость автомобиля, но и адаптирует ее согласно текущей ситуации, например когда впереди движется автомобиль) и стабилизации курса. Этот уровень доступен коммерческим автомобилям с 2013 года.
- » **Уровень 3. Условная автоматизация.** Большинство автомобилестроителей работало над этим на момент написания этой книги. *Условная автоматизация* означает, что автомобиль может взять на себя управление при некоторых условиях (например, только на магистралях или на дорогах с односторонним движением), при ограничении скорости и под бдительным человеческим контролем. Автомат может также потребовать передать управление человеку. Одним из примеров этого уровня автоматизации являются недавние модели автомобилей, самоуправляемые на магистрали и автоматически тормозящие, когда трафик замедляется.
- » **Уровень 4. Высокая автоматизация.** Автомобиль выполняет все задачи вождения (руль, газ и тормоз) и контролирует любые изменения в дорожных условиях от начала до конца поездки. Этот

уровень автоматизации не требует человеческого вмешательства и доступен только в определенных местах и ситуациях, поэтому водитель должен быть готов при необходимости вернуть управление себе. Производители ожидают появления такого уровня автоматизации примерно в 2020-х годах.

- » **Уровень 5. Полная автоматизация.** Автомобиль способен сделать все сам от места отправления до места назначения, без человеческого вмешательства. Способности автоматики такого уровня сравнимы или превосходят таковые у водителя. У автоматизированных автомобилей уровня 5 не будет рулевого колеса. Этот уровень автоматизации ожидается к 2025-му году.

Даже когда беспилотные автомобили достигнут пятого уровня автономности, вы не увидите их слоняющимися по всем дорогам. Такие автомобили все еще в далеком будущем, а в будущем могут возникнуть трудности. В разделе “Преодоление неопределенности восприятия”, приведенном далее в этой главе, обсуждаются некоторые из препятствий, с которыми встретится искусственный интеллект при вождении автомобиля. Внезапно беспилотный автомобиль явно не появится; вероятнее всего, он будет получен в ходе серии модификаций одновременно со все большим и большим количеством автомобилей автоматизированных моделей. Люди будут держаться за рулевое колесо еще достаточно долго. То, что вполне можно ожидать увидеть, — это искусственный интеллект, помогающий как в повседневном вождении, так и в опасных условиях, повышая безопасность на дорогах. Даже когда производители коммерциализируют беспилотные автомобили, на замену имеющегося парка могут уйти годы. Процесс полного изменения дорожного движения в городских условиях беспилотными автомобилями может занять лет 30.



ВНИМАНИЕ!

В этом разделе содержится много дат, и некоторые люди склонны полагать, что любая присутствующая в книге дата точна. Однако случиться может всякое, что способно ускорить или задержать принятие беспилотных автомобилей. Например, система страхования в настоящее время с подозрением относится к беспилотным автомобилям, поскольку опасается, что в будущем их страхование будет затруднительно, поскольку снизится риск дорожно-транспортных происшествий. (Консалтинговая фирма McKinsey прогнозирует, что количество дорожно-транспортных происшествий снизится на 90 процентов: <https://www.mckinsey.com/industries/automotive-and-assembly/our-insights/ten-ways-autonomous-driving-could-redefine-the-automotive-world>). Лоббирование системой страхования может задержать принятие беспилотных автомобилей. С другой стороны, люди, перенесшие потерю близкого человека в

дорожно-транспортных происшествиях, вероятно, поддержат то, что снизит количество таких происшествий. Они могли бы способствовать успеху и ускорить принятие беспилотных автомобилей. Следовательно, множество разнообразных факторов, включая социальное давление, способно изменить историю, и точное прогнозирование дат принятия беспилотных автомобилей абсолютно невозможно.

Пересмотр мнения о роли автомобилей

Мобильность неразрывно связана с цивилизацией. Это не только транспортировка людей и грузов, но и доступность отдаленных мест. Когда автомобили впервые появились на дорогах, лишь немногие полагали, что они скоро заменят лошадей и кареты. У автомобилей все же много преимуществ перед лошадьми: они практичнее в обслуживании, их скорость выше, а проходимые ими дистанции длиннее. Автомобили также требуют большего контроля и внимания от людей, поскольку лошади видят дорогу и реагируют на препятствия, избегая возможных столкновений, но люди пожертвовали этим в обмен на большую мобильность.

Сегодня использование автомобилей формирует и городскую среду, и экономическую жизнь. Автомобили позволяют людям преодолевать дальние расстояния от дома до работы каждый день (делая пригороды возможными для реализации недвижимости). Фирмы легко доставляют товары на большие расстояния, автомобили создают новые компании и рабочие места, а рабочие автомобильной промышленности давно стали основными действующими лицами в новом перераспределении богатств. Автомобиль — первый реальный продукт рынка товаров массового потребления, созданный одними рабочими для других рабочих. Когда автомобильный бизнес процветает, общество благополучно; когда он гибнет, может последовать катастрофа. Поезда и самолеты привязаны к заранее определенным рейсам, а автомобили — нет. Автомобили обеспечили мобильности куда больший масштаб и повлияли на повседневную жизнь людей куда больше, чем другие транспортные средства. Как сказал Генри Форд, основатель автомобильной компании Ford, “автомобили освободили простых людей от географических ограничений”.

Когда беспилотные автомобили появятся на рынке, цивилизация окажется на грани новой вызванной ими революции. Когда производители достигнут пятого уровня автономности и беспилотные автомобили станут господствующей тенденцией, вполне можно ожидать существенных смещений акцентов в градостроении, экономике и всеобщем образе жизни. Есть очевидные и менее очевидные способы, которыми беспилотные автомобили изменят жизнь. Наиболее очевидны и наиболее обсуждаемы следующие.

- » **Меньше дорожно-транспортных происшествий.** ДТП станет меньше потому, что искусственный интеллект будет соблюдать правила дорожного движения и безопасности; он куда более умный водитель, чем люди. Сокращение количества происшествий существенно повлияет на способ создания автомобилей производителями, которые будут намного безопаснее, чем в прошлом, из-за структурных пассивных средств защиты. В будущем, с учетом их абсолютной безопасности, беспилотные автомобили могли бы стать легче из-за меньшего количества средств защиты, чем теперь. Они могут быть сделаны даже из пластмассы. В результате автомобили будут расходовать меньше ресурсов, чем сегодня. Кроме того, снижение количества несчастных случаев будет означать снижение страховых выплат, обусловив существенное влияние на систему страхования, имеющую отношение к происшествиям.
- » **Меньше задач, связанных с вождением.** Многие задачи, связанные с вождением, исчезнут или потребуют меньшего количества рабочих. Это позволит сэкономить на рабочей силе в транспортной сфере, сделав транспортировку товаров и людей более доступной, чем теперь. Это также создаст проблемы поиска новой работы для людей. (В одних только Соединенных Штатах на транспорте занято около трех миллионов человек.)
- » **Больше времени.** Беспилотные автомобили подарят людям самую драгоценную вещь в жизни — их время. Беспилотные автомобили не помогут людям жить дольше, но помогут сэкономить время, которое они тратили на вождение, для использования в других целях (поскольку вести будет искусственный интеллект). Кроме того, даже если трафик повысится (из-за меньшей стоимости транспортировки и других факторов), он станет более плавным, с небольшими перегрузками или даже без них. Кроме того, пропускная способность существующих транспортных магистралей увеличится. Это может показаться парадоксальным, но в том и мощь искусственного интеллекта, когда люди остаются не у дел, как демонстрируется в видео <https://www.youtube.com/watch?v=iHzzSao6ypE>.

Кроме этих непосредственных эффектов, будут и менее уловимые, которые могут быть незаметны сразу, но впоследствии станут очевидными. Бенедикт Эванс (Benedict Evans) указывает на некоторые из них в своем сообщении “Cars and second order consequences” (Автомобили и последствия второго порядка) (<https://www.ben-evans.com/benedictevans/2017/3/20/cars-and-second-order-consequences>). В этой статье рассматриваются последствия появления на рынке электрических автомобилей пятого уровня автономности и беспилотных автомобилей. В качестве одного из примеров беспилотные автомобили могли сделать антиутопический паноптикум реальностью

(см. <https://www.theguardian.com/technology/2015/jul/23/panopticon-digital-surveillance-jeremy-bentham>). *Паноптикум* — это административное здание из теорий английского философа конца XVIII века Иеремии Бен-тама (Jeremy Bentham), в котором все находится под наблюдением, не зная об этом. Когда беспилотные автомобили выйдут на улицы в большом количестве, автомобильные камеры будут повсюду, отслеживая и, возможно, сообщая обо всем, что увидят. Ваш автомобиль может шпионить за вами и другими, когда вы меньше всего ожидаете этого.

Размышление о будущем не является простым упражнением, поскольку это не просто вопрос причин и следствий. Даже приблизительные предположения могут оказаться неэффективными, поскольку сам контекст ожидаемого способен измениться. Например, паноптикум никогда не сможет появиться потому, что законодательная система могла бы вынудить беспилотные автомобили скрывать снимаемые изображения. Поэтому предлагаемые предсказателями сценарии являются весьма приблизительными в описании возможного будущего; эти сценарии могут сбыться, а могут и не сбыться, в зависимости от различных обстоятельств. Эксперты размышляют над тем, что автомобиль, обладающий возможностями автономного вождения, может оказаться в четырех разных сценариях, полностью переопределяющих то, как люди используют автомобиль.

- » **Автономное вождение на дальних дорогах и магистралях.** Когда водитель может добровольно позволить искусственному интеллекту вести автомобиль, водитель может уделить внимание другим задачам. Многие считают этот сценарий возможной причиной использования автономных автомобилей. Однако, с учетом высоких скоростей на магистралях, передача управления искусственному интеллекту не особенно безопасна, поскольку другие автомобили, ведомые людьми, могут вызвать аварийную ситуацию. Люди должны учитывать последствия, такие как текущее дорожное законодательство, разное в разных странах. Вопрос в том, что некоторые законодательства могут рассматривать водителя, использующего искусственный интеллект, как невнимательного. Этот сценарий относится к автономности уровня 3.
- » **Действия при парковке.** В этом сценарии искусственный интеллект вмешивается, когда пассажиры оставили автомобиль искать место для парковки. Беспилотный автомобиль экономит время владельцам, поскольку открывает возможность как оптимизации места для стоянки автомобилей (беспилотный автомобиль будет знать, где лучше всего парковаться), так и совместного использования автомобилей. (После того как вы оставили автомобиль, некто еще может использовать его; позже вы воспользуетесь другим автомобилем,

Некоторые говорят, что страховая ответственность и проблема вагонетки серьезно препятствуют использованию беспилотного автомобиля. Страховая проблема подразумевает вопрос, кто виноват, если что-то пошло не так, как надо. ДТП случаются и теперь, и хотя беспилотные автомобили должны вызывать меньше происшествий, чем люди, казалось бы, легко решаемая автомобилестроителями проблема остается и система страхования не может обеспечивать беспилотные автомобили страховкой. (Система страхования опасается беспилотных автомобилей потому, что их использование может в корне изменить сам бизнес.) Создатели беспилотных автомобилей, такие как Audi, Volvo, Google и Mercedes-Benz, уже обязались принимать на себя ответственность, если их транспортные средства вызовут происшествие (см. <https://www.cohen-lawyers.com/wp-content/uploads/2016/08/WestLaw-Automotive-Cohen-Commentary.pdf>). Это значит, что автомобилестроители станут страховщиками для большей пользы продвижения беспилотных автомобилей на рынок.

Проблема вагонетки (trolley problem) — это моральная задача, поставленная британским философом Филиппой Фут (Philippa Foot) в 1967 году (но сама дилемма довольно древняя). Это проблема неконтролируемой вагонетки, которая, двигаясь по рельсам, переезжает нескольких привязанных к ним человек, но вы можете их спасти, переведя вагонетку на другой путь, на котором к рельсам привязан только один человек. Необходимо сделать выбор, зная, что кто-то точно умрет. Существует довольно много вариантов проблемы вагонетки, и Массачусетский технологический институт даже предлагает веб-сайт, <http://moralmachine.mit.edu/>, с альтернативными ситуациями, более подходящими для беспилотного автомобиля.

оставленным поблизости в месте для стоянки автомобилей.) Учитывая ограниченность автономного вождения только некоторыми автомобилями, этот сценарий является переходным с третьего уровня автономности на четвертый.

- » **Автономное вождение в любой поездке, кроме поездки в те места, где беспилотные автомобили вне закона.** Этот дополнительный сценарий позволяет искусственному интеллекту управлять в любых местах, кроме мест, запрещенных по соображениям безопасности (таких, как новые дорожные инфраструктуры, отсутствующие в системе навигации). Этот сценарий для беспилотных автомобилей ближе к зрелости (уровень автономности — 4).

» **Такси по требованию.** Это продолжение сценария 2, когда беспилотные автомобили уже достаточно зрелы, чтобы водить постоянно (уровень автономности — 5) как с пассажирами, так и без, оказывая транспортные услуги любому по требованию. Это сценарий полного использования автомобилей (сегодня автомобили 95 процентов времени проводят на парковке; см. <http://fortune.com/2016/03/13/cars-parked-95-percent-of-time/>), что в корне изменит идею владения автомобилем, поскольку вы не будете нуждаться в одном и собственном.

Введение в беспилотные автомобили

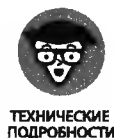
Создание беспилотного автомобиля, вопреки воображению людей, не подразумевает водружение робота на переднее сиденье и вручение ему управления автомобилем. Люди решают бесчисленные задачи, управляя автомобилем, с которыми робот не знал бы что делать. Создание подобного человеческому интеллекта требует гармоничного взаимодействия многих систем, обеспечивающих правильную и безопасную дорожную обстановку. Понадобится много усилий, чтобы получить комплексное решение, а не полагаться на отдельные решения искусственного интеллекта для каждой потребности. Проблема разработки беспилотного автомобиля требует решения многих отдельных проблем и наличия индивидуальных решений для эффективного взаимодействия. Например, распознавание дорожных знаков и изменение маршрутов требует разных систем.



ЗАПОМНИ

Комплексное решение (end-to-end solution) — это то, что вы часто слышите при обсуждении роли глубокого обучения в искусственном интеллекте. Мощь обучения на примерах столь велика, что многие задачи не требуют отдельных решений, являющихся, по существу, комбинацией многих меньших задач, каждая из которых решается в соответствии с другим решением искусственного интеллекта. Глубокое обучение может решить проблему в целом, имея примеры решений и вырабатывая уникальные решения, охватывающие все проблемы, требовавшие отдельных решений искусственного интеллекта в прошлом. Проблема в том, что сегодня возможности глубокого обучения слишком ограничены для фактического решения этой задачи. Единое решение глубокого обучения может работать для некоторых задач, но другие все еще требуют комбинации меньших решений искусственного интеллекта, если необходимо получить надежное полное решение.

Над комплексными решениями глубокого обучения работает компания NVidia — производитель графических процессоров. Посмотрите видео <https://www.youtube.com/watch?v=-96BEoXJMs0>, демонстрирующее пример эффективности решения. Но все же и это справедливо для любого применения глубокого обучения: уровень совершенства решения непосредственно зависит от полноты и количества использованных примеров. Чтобы получить функцию беспилотного автомобиля как комплексное решение глубокого обучения, требуется набор данных, который научит управлять автомобилем в контекстах огромного количества ситуаций, которые еще недоступны, но могут возникнуть в будущем.



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Есть надежда, что комплексные решения упростят структуру беспилотных автомобилей. В статье <https://devblogs.nvidia.com/explaining-deep-learning-self-driving-car/> объясняется, как работает процесс глубокого обучения. Вы можете также прочитать оригинальную статью NVidia о том, как сквозное обучение помогает водить автомобиль <https://arxiv.org/pdf/1704.07911.pdf>.

Объединение всех технологий

Под покровом беспилотного автомобиля кроются системы, взаимодействующие согласно робототехнической парадигме считывания, планирования и действия. Все начинается на уровне считывания, где множество разных сенсоров сообщают автомобилю различную информацию.

- » Датчик GPS указывает, где автомобиль находится (по системе карт), выдавая координаты широты, долготы и высоты.
- » Такие устройства, как радар, сонар и лидар, обнаруживают объекты и предоставляют данные об их расположении и движении в терминах изменения их координат в пространстве.
- » Камеры оповещают автомобиль о его окружении, преобразуя снимки в цифровые образы.

В беспилотном автомобиле применяется множество специализированных сенсоров. В разделе “Преодоление неопределенности восприятия”, приведенном далее в этой главе, они описаны подробнее; также в нем демонстрируется, как система объединяет их вывод. Система должна объединить и обработать данные сенсоров, прежде чем необходимый автомобилю уровень восприятия начнет работать и станет полезным. Поэтому объединение данных сенсоров и определяет различные точки зрения на мир вокруг автомобиля.

Локализация позволяет узнать, где именно находится автомобиль. Эта задача решается главным образом обработкой данных от устройства GPS. Система GPS — это система навигации на базе космических спутников, первоначально

созданная в военных целях. Будучи использованной в гражданских целях, она обнаружила некоторую встроенную погрешность (чтобы только уполномоченный персонал мог использовать ее с максимальной точностью). Те же погрешности присутствуют и в других системах, таких как GLONASS (русская система навигации), GALILEO (или GNSS, европейская система) и BeiDou (или BDS, китайская система). Следовательно, независимо от используемой спутниковой системы автомобиль может понять, где он находится на конкретной дороге, но не может установить используемую полосу движения (или даже может закончить тем, что окажется на параллельной дороге). В дополнение к грубому определению местоположения по GPS система уточняет данные GPS данными сенсора оптического локатора, чтобы определить точную позицию в окружающей среде.

Система обнаружения определяет происходящее вокруг автомобиля. Ей требуется много подсистем, каждая из которых предназначена для конкретной цели и использует уникальный комплекс сенсорных данных и результаты их анализа.

- » Определение полосы движения осуществляется в результате обработки изображений с камеры и их анализа либо в результате применения специализированной сети глубокого обучения для *сегментации изображения*, при котором изображение делится на отделенные области, размеченные по типу (т.е. дорога, автомобили и пешеходы).
- » Обнаружение дорожных знаков и светофоров, а также их классификация (classification) осуществляются с помощью обработки изображений с камер с использованием сети глубокого обучения, которая сначала определяет область содержащего знак изображения или светофора, а затем маркирует их (тип знака или цвет светофора). Следующая статья NVidia поможет понять, как видит беспилотный автомобиль: <https://blogs.nvidia.com/blog/2016/01/05/eyes-on-the-road-how-autonomous-cars-understand-what-theyre-seeing/>.
- » Объединенные данные от радара, лидара, сонара и камер позволяют обнаруживать внешние объекты и отслеживать их движение, включая направление, скорость и ускорение.
- » Данные лидара используются главным образом для обнаружения свободного пространства на дороге (свободная полоса или место для стоянки).

Появление на сцене искусственного интеллекта

После фазы считывания, позволяющей беспилотному автомобилю определить, где он находится и что происходит вокруг него, начинается фаза пла-

нирования. Теперь искусственный интеллект выходит на сцену полностью. Планирование беспилотных автомобилей сводится к решению следующих конкретных задач.

- » **Маршрут.** Определить путь, которым должен следовать автомобиль. Находясь в автомобиле, едущем по назначению (хорошо, это не всегда так, но так обычно бывает), вы хотите достичь места назначения самым быстрым и самым безопасным способом. В некоторых случаях следует учитывать стоимость. Здесь должны помочь классические алгоритмы маршрутизации.
- » **Прогноз окружающей обстановки.** Необходимо помочь автомобилю спроецировать себя в будущее, поскольку ему нужно время, чтобы понять ситуацию, выбрать маневр и завершить его. За время, необходимое для осуществления маневра, другие автомобили могли решить изменить свою позицию или начать собственные маневры. При вождении вы пытаетесь понять намерения других водителей, чтобы избежать возможных столкновений. Беспилотный автомобиль делает то же самое, используя прогноз машинного обучения для оценки того, что произойдет далее, а затем принимает во внимание будущее.
- » **Планирование поведения.** Это базовый интеллект автомобиля. Он включает практики (practices), необходимые для успешного участия в дорожном движении: следование по полосе; смена полосы; съезд и выезд на дорогу; поддержание дистанции; учет светофоров, знаков и надписей; избегание препятствий; и многое другое. Все эти задачи выполняет искусственный интеллект, такой как экспертная система, включающая опыт многих водителей, или вероятностная модель, такая как байесовская сеть, или даже упрощенная модель машинного обучения.
- » **Планирование траектории.** Определяет фактическое выполнение автомобилем необходимых задач при наличии нескольких путей ее решения. Например, когда автомобиль решает сменить полосу, это должно быть сделано без резкого ускорения или чрезмерного сближения с другим автомобилем, т.е. перемещаться приемлемым, безопасным и лишенным неприятных маневров способом.

Не искусственным интеллектом единым

После считывания и планирования наступает время беспилотному автомобилю действовать. Считывание, планирование и действие являются частями бесконечного цикла, повторяемого до тех пор, пока автомобиль не достигнет места назначения и не остановится на парковке. Действие подразумевает такие базовые функции, как ускорение, торможение и применение системы рулевого

управления. Инструкции вырабатываются на фазе планирования, а автомобиль просто выполняет действия с помощью системы контроллера, такой как контролер PID (Proportional-Integral-Derivative) или MPC (Model Predictive Control), обладающей алгоритмами проверки правильности выполнения этих действий; в противном случае он немедленно предпринимает подходящие контрмеры.

Это может показаться немного сложным, но на всем протяжении поездки система выполняет только эти три действия, одно за другим. Каждая система содержит подсистемы, решающие одну задачу вождения, используя самые быстрые и самые надежные алгоритмы, как изображено на рис. 14.1.

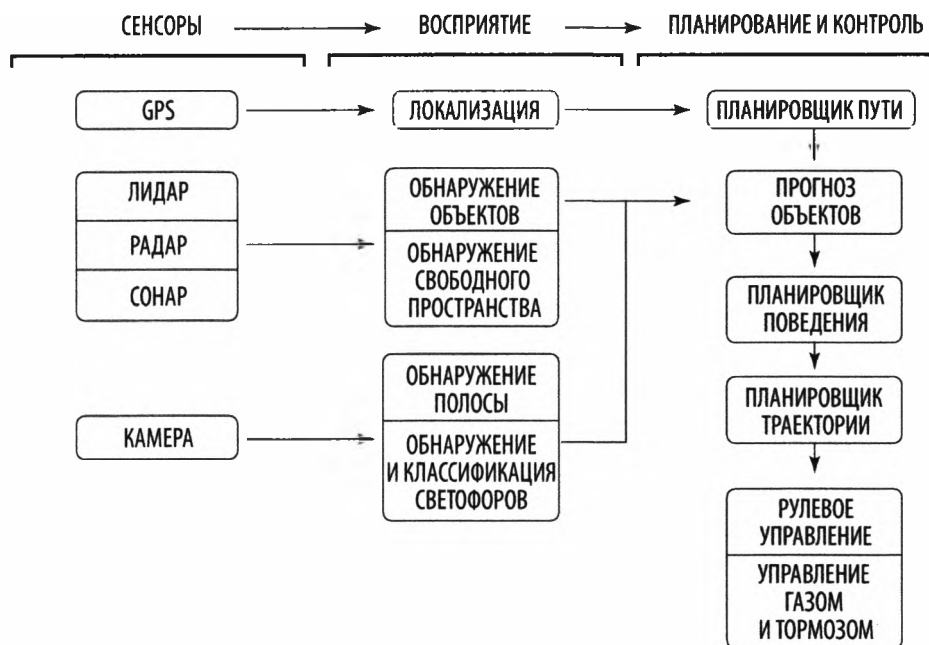


Рис. 14.1. Общая схема систем, работающих в беспилотном автомобиле

Таким было состояние на момент написания книги. Но беспилотные автомобили, вероятно, так и останутся комплексом программного обеспечения и оборудования, выполняющим различные функции и действия. В некоторых случаях системы будут обладать избыточными функциональными возможностями, такими как использование нескольких сенсоров для отслеживания того же самого внешнего объекта, или применять уровень восприятия нескольких систем для гарантии нахождения на правильной полосе. Избыточность позволяет существенно снизить количество ошибок, а следовательно, и происшествий. Например, даже когда система глубокого обучения для распознавания дорожных знаков откажет или ошибется (см. <https://thehackernews>.

com/2017/08/self-driving-car-hacking.html), другие системы смогут ее продублировать и минимизировать или аннулировать последствия для автомобиля.

Преодоление неопределенности восприятия

Стивен Пинкер (Steven Pinker), профессор в Отделе психологии Гарвардского университета, пишет в своей книге *The Language Instinct: How the Mind Creates Language*, что “в робототехнике простые проблемы трудны, а трудные проблемы просты”. Фактически искусственный интеллект играет в шахматы против гроссмейстеров невероятно успешно; однако более обыденные действия, такие как распознавание объектов в таблице, уклонение от столкновения с пешеходом, распознавание лиц или правильный ответ на вопрос по телефону, могут оказаться весьма трудными для искусственного интеллекта.



ЗАПОМНИ

Парадокс Моравека гласит, что то, что просто для людей, трудно для искусственного интеллекта (и наоборот). Он был сформулирован в 1980-х годах робототехниками-когнитивистами Хансом Моравеком (Hans Moravec), Родни Бруксом (Rodney Brooks) и Марвином Минском (Marvin Minsk). У людей бывает достаточно времени для выработки таких навыков, как ходьба, бег, удержание объектов, речь и наблюдение; эти навыки вырабатывались в ходе эволюции и естественного отбора на протяжении миллионов лет. Чтобы выжить в этом мире, люди делают то, что делали все живые существа с момента появления жизни на Земле. И наоборот, высокая абстракция и математика — это относительно новое изобретение людей, и мы, естественно, еще не адаптировались к ним.

У автомобилей есть некоторые преимущества перед роботами, которые должны прокладывать путь в зданиях и на пересеченной местности. Автомобили работают на специально созданных для них дорогах, обычно хорошо картографированных, у автомобилей уже есть все механические решения для движения по дорожной поверхности.

Исполнительные механизмы — не самая большая проблема беспилотных автомобилей. Серьезные проблемы создают планирование и считывание. Планирование — это самый высокий уровень (здесь искусственный интеллект превосходит). Но когда дело доходит до общего планирования, беспилотные автомобили уже могут полагаться на навигаторы GPS, тип искусственного интеллекта, специализированного для поддержания направления. Считывание — это реально узкое место для беспилотных автомобилей, поскольку без

него невозможны ни планирование, ни действия. Водители считают дорогу все время, чтобы удерживать автомобиль на полосе, заменить препятствия и соблюсти необходимые правила.



ЗАПОМНИ

На данном этапе развития беспилотных автомобилей аппаратные средства считывания непрерывно модернизируются в поисках более надежных, точных и менее дорогих решений. С другой стороны, для обработки данных сенсоров и их использования применяются более эффективные и надежные алгоритмы, такие как *фильтр Калмана* (см. <http://www.bzarg.com/p/how-a-kalman-filter-works-in-pictures/> и https://home.wlu.edu/~levys/kalman_tutorial/), которые известны уже несколько десятилетий.

Знакомство с сенсорами автомобиля

Сенсоры — это ключевые компоненты восприятия окружающей обстановки, и беспилотный автомобиль может считывать информацию в двух направлениях — внутрь и вовне.

- » **Проприоцептивные сенсоры.** Отвечают за считывание таких состояний транспортного средства, как состояние его систем (двигатель, передача, тормоза и рулевое управление) и позиция по датчику GPS, а также вращение колес, скорость и ускорение транспортного средства.
- » **Экстероцептивные сенсоры.** Отвечают за считывание окружающей обстановки с использованием таких сенсоров, как камера, лидар, радар и сонар.

И проприоцептивные, и экстероцептивные сенсоры способствуют автономности беспилотного автомобиля. В частности, локализация GPS обеспечивает представление о примерном местоположении беспилотного автомобиля, что полезно при общем планировании направлений и действий по его успешному следованию к месту назначения. Датчик GPS помогает беспилотному автомобилю в такой же степени, как и любому водителю: указывает правильное направление.

Экстероцептивные сенсоры (представленные на рис. 14.2) помогают автомобилю при собственно вождении. Они усиливают или заменяют человеческие органы чувств в данной ситуации. Каждый из них предлагает собственную точку зрения на окружающую обстановку; каждый имеет свои конкретные ограничения; и каждый имеет свои преимущества.

Ограничения бывают разными. По мере понимания того, что сенсоры делают для беспилотного автомобиля, становятся заметными проблемы стоимости,

чувствительности к освещению, чувствительности к погоде и помехам (что означает чувствительность к замене сенсора, затрагивающей его точность), диапазону и разрешению. С другой стороны, они могут точно отслеживать скорость, позицию, высоту и дистанцию до объектов, а также обнаруживать объекты и классифицировать их.

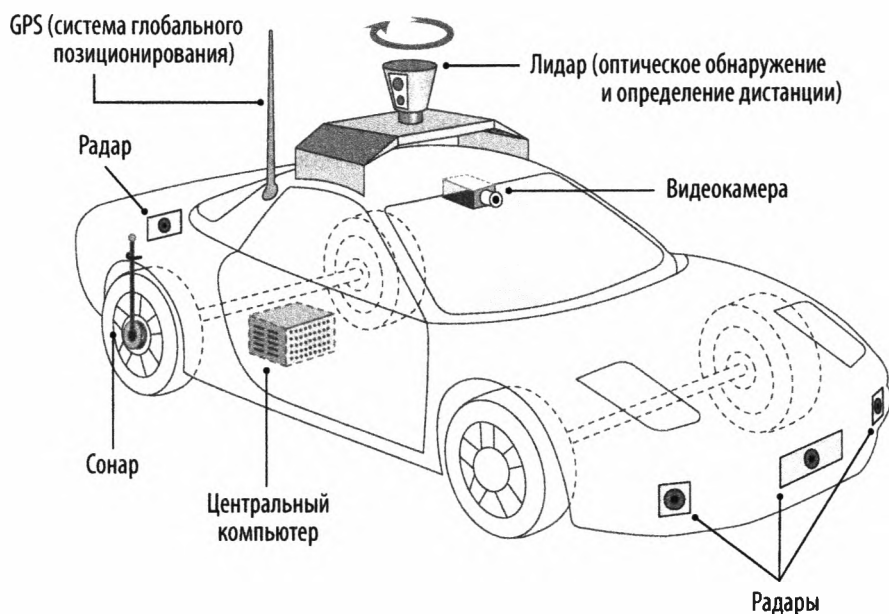


Рис. 14.2. Схематическое представление экстероцептивных сенсоров в беспилотном автомобиле

Камера

Камеры — это пассивные зрительные сенсоры. Они могут обеспечить монокулярное или бинокулярное зрение. С учетом их низкой цены вы можете оснастить автомобиль множеством камер: на ветровом стекле, на переднем бампере, на боковых зеркалах, на задней двери и на заднем стекле. Обычно стереокамеры имитируют человеческому зрению и предоставляют информацию о дороге и соседних транспортных средствах, тогда как монокулярные камеры обычно специализируются на обнаружении дорожных знаков и светофоров. Получаемые ими данные обрабатываются алгоритмами для обработки изображений или нейронными сетями глубокого обучения, чтобы обеспечить их обнаружение и классификацию (например, распознать красный свет светофора или знак ограничения дозволенной скорости). У камер может быть высокое разрешение (они могут распознавать маленькие детали), но они чувствительны к освещению и погоде (ночь, туман или визуальные помехи).

Лидар

Лидар использует инфракрасные лучи (длина волны — порядка 900 нм, они невидимы для человеческого глаза), позволяющие оценить расстояние до объектов. Они используют вращающуюся платформу, чтобы проецировать луч по кругу, а затем возвращают результат в виде облака из точек до препятствий, что позволяет оценить форму и размеры свободного пространства. В зависимости от цены (обычно чем дороже, тем лучше) у лидара вполне может быть более высокое разрешение, чем у радара. Но лидар непрочен и может быть загрязнен куда быстрее, чем радар, поскольку расположен вне автомобиля. (Лидар — это вращающееся устройство, которое вы видите наверху автомобиля Google в видео от CBS: https://www.youtube.com/watch?v=_qE5VzuYFPU.)

Радар

Это устройство работает по принципу отражения радиоволн от препятствий, а время следования радиоволн к препятствию и обратно позволяет установить дистанцию и скорость до него. Радар может быть установлен на переднем и заднем бамперах, а также по бокам автомобиля. Производители уже довольно давно используют их на автомобилях для адаптивной системы автоматического регулирования скорости, предупреждения о препятствиях в непросматриваемых зонах, предупреждения и предотвращения столкновений. В отличие от других сенсоров, нуждающихся в нескольких последовательных показаниях, радар может выяснить скорость объекта одним импульсом благодаря эффекту Доплера (см. <http://www.physicsclassroom.com/class/waves/Lesson-3/The-Doppler-Effect>). Радар может быть дальним или ближним, может создать схему окружающей среды, а может использоваться для локализации. На радар погодные условия воздействуют куда меньше, чем на другие средства обнаружения, в частности дождь или туман; он имеет угол зрения порядка 150 градусов и дальность 30–200 метров. Его наибольшая слабость — недостаток разрешения (радар не способен предоставлять подробности и обнаруживать статические объекты).

Сонары

Сонары подобны радару, но используют вместо микроволн высокочастотный звук (ультразвук, неслышимый людьми, но слышимый некоторыми животными). Главным недостатком сонаров, используемых изготовителями вместо весьма непрочных и более дорогих лидаров, является их малая дальность.

Объединение обнаруженного

Когда дело доходит до считывания обстановки вокруг беспилотного автомобиля, можно полагаться на различные показатели в зависимости от того, какие

сенсоры установлены на автомобиле. Но все же у каждого сенсора — собственные разрешение, дальность и чувствительность к помехам, что приводит к разным показателям для одной и той же ситуации. Другими словами, ни один из них не совершенен, и иногда их недостатки препятствует надлежащему обнаружению. Сигналы сонара и радара могут быть поглощены; лучи лидара не могут пройти сквозь твердые частицы. Кроме того, отражения и плохое освещение вполне могут вводить камеры в заблуждение, как описано в статье *MIT Technology Review* <https://www.technologyreview.com/s/608321/this-image-is-why-self-driving-cars-come-loaded-with-many-types-of-sensors/>.

Беспилотные автомобили призваны улучшить нашу мобильность, а значит, они должны быть безопасны и для пассажиров, и для окружающих. Нельзя разрешить беспилотному автомобилю оказаться не в состоянии обнаруживать пешехода, который внезапно появился перед ним. Из соображений безопасности производители сосредотачивают большие усилия на комбинировании и объединении данных от разных сенсоров, чтобы получить объединенный показатель, который куда лучше любого отдельного показателя. Объединение сенсоров — это зачастую результат использования вариантов фильтра Калмана (таких, как расширенный фильтр Калмана или даже более сложный сигма-точечный фильтр Калмана). Рудольф Калман был венгерским инженером-электриком и изобретателем, иммигрировавшим в Соединенные Штаты во время второй мировой войны. За свои изобретения, нашедшие весьма многочисленные применения в управлении, навигации и контроле транспортных средств, от автомобилей и самолетов до космических кораблей, Калман получил Национальную научную медаль США в 2009 году от президента США Барака Обамы.

Алгоритм фильтра Калмана осуществляет фильтрацию результатов нескольких разных измерений, полученных на протяжении некоторого времени, в единую последовательность показателей, которые обеспечивают реальную оценку (предыдущие показатели были не особенно точны). Это работает так: объект обнаруживается при первом проведении всех измерений и обработке их результатов (фаза прогноза состояния), что позволяет оценить текущую позицию объекта. Затем, когда поступают новые показатели, новые результаты используются для модификации предыдущих, чтобы получить более надежную оценку позиции и скорость объекта (фаза обновления показателей), как показано на рис. 14.3.

Таким образом, беспилотный автомобиль может передать алгоритму показатели сенсоров и использовать их для получения результирующей оценки положения окружающих объектов. Оценка объединяет преимущества показаний всех сенсоров и избегает их недостатков. Это возможно потому, что фильтр использует более сложную версию вероятностей и теоремы Байеса, рассматриваемых в главе 10.

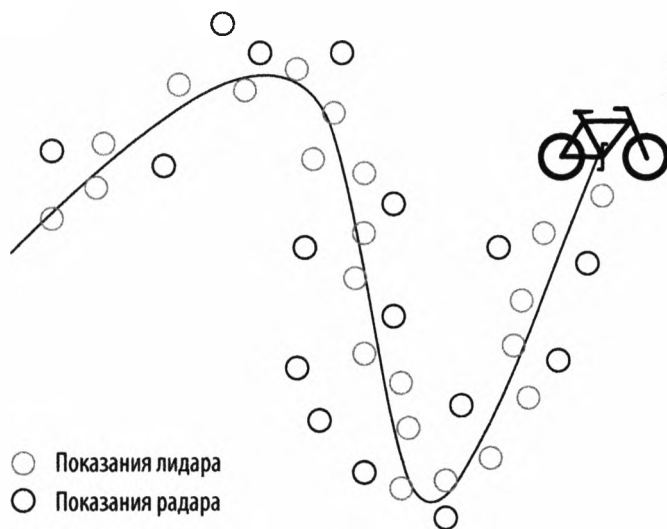


Рис. 14.3. Фильтр Калмана оценивает траекторию движения велосипеда, объединяя данные лидара и радара

5 Будущее искусственного интеллекта

В ЭТОЙ ЧАСТИ...

- » Определение, когда приложение не будет работать**
- » Использование искусственного интеллекта в космосе**
- » Создание новых профессий**



Глава 15

Причины неудач приложений

В ЭТОЙ ГЛАВЕ...

- » Сценарии применения искусственного интеллекта и последствия его отказа
- » Поиск решений несуществующих проблем

В предыдущих главах этой книги исследовалось, чем искусственный интеллект является и чем не является, а также задачи, некоторые из которых искусственный интеллект может решать хорошо, а некоторые — не может. Даже обладая всей этой информацией, вы можете легко распознать приложения, которые никогда не выйдут в свет, поскольку искусственный интеллект просто не может удовлетворить данную конкретную потребность. В этой главе рассматриваются приложения, обреченные на неудачу. Возможно, эту главу следовало бы назвать “Почему нам всем еще нужны люди”, но текущее название более понятное.

Часть этой главы посвящена попыткам создания неудачных приложений. Самым серьезным последствием таких попыток стала зима искусственного интеллекта. *Зима искусственного интеллекта* случается всякий раз, когда обязательства сторонников искусственного интеллекта превышают возможности их выполнения, что приводит к потере финансирования от предпринимателей.

Искусственный интеллект может также попасть в ловушку поиска решения проблемы, которой на самом деле не существует. Да, выдающиеся решения

действительно выглядят весьма странно, но если решение не относится к реальной потребности, никто его покупать не будет. Технологии процветают только тогда, когда удовлетворяют некие потребности и пользователи согласны потратить на это деньги. Эта глава завершается рассмотрением решения проблем, которых на самом деле не существуют.

Использование искусственного интеллекта там, где он не будет работать

В табл. 1.1 в главе 1 перечислены семь видов интеллекта. Общество в целом обладает всеми семью видами интеллекта, а люди по отдельности выделяются способностями в некоторых видах интеллекта.

Объединив усилия всех людей, можно использовать все семь видов интеллекта таким способом, который удовлетворяет потребностям общества.



ЗАПОМНИ

Из табл. 1.1 ясно, что искусственный интеллект не способен на два вида интеллекта вообще и демонстрирует весьма скромные способности еще в трех. Искусственный интеллект превосходит, когда дело касается математики, логики и кинестетического интеллекта; его способности ограничены решением множества видов задач, которые должно решать сообщество в целом. В следующих разделах описаны ситуации, в которых искусственный интеллект просто не сможет работать, поскольку это лишь технология, а не человек.

Определение границ искусственного интеллекта

Когда говоришь с Alexa, иногда забываешь, что это всего лишь машина. Машина понятия не имеет, о чем вы говорите; она не воспринимает вас как человека и не имеет ни малейшего желания общаться с вами; она просто действует по определенным алгоритмам, созданным для нее, и предоставленным ей данным. И даже в этом случае результаты удивительны. Довольно просто наделить искусственный интеллект человеческими качествами, не понимая его и рассматривая лишь как разновидность подобной человеку сущности. Однако искусственному интеллекту не хватает важнейших возможностей, рассматриваемых в следующих разделах.

Творчество

Можно найти бесконечное разнообразие статей, сайтов, музыки, картин, текстов и других примеров творчества искусственного интеллекта. Проблема

искусственного интеллекта в том, что он ничего не может создать. Думая о творческом потенциале, вы имеете в виду образ мыслей.

Например, у Бетховена было прекрасное понимание музыки. Вы можете узнать фрагмент классического произведения Бетховена, даже не будучи знакомыми со всеми его работами, поскольку у музыки есть специфический шаблон, сформированный тем способом, которым думал Бетховен.

Искусственный интеллект может создать новую пьесу Бетховена, отследив его мыслительный процесс математически, что искусственный интеллект и делает при обучении на примерах музыки Бетховена. В основе процесса создания новой пьесы Бетховена лежит математический механизм. Фактически благодаря математическим шаблонам вы можете услышать, как искусственный интеллект играет Бетховена в аранжировке других, включая Битлз <https://techcrunch.com/2016/04/29/paul-mccartificial-intelligence/>.



ЗАПОМНИ

Проблема составления математического уравнения творческого потенциала в том, что математика не является творческой. Быть творческим означает выработать новый шаблон мышления, которого не было прежде (см. <https://www.csun.edu/~vcpsy00h/creativity/define.htm>). Творческий потенциал — это не только мышление вне рамок; это действие по определению новых рамок.

Творчество подразумевает также разработку иной точки зрения, которая по существу определяет иной вид набора данных (если вы настаиваете на математической точке зрения). Искусственный интеллект ограничивается теми данными, которые ему были предоставлены. Он не может создать собственные данные; он может создать только варианты имеющихся данных — тех данных, на которых он обучался. Эта идея разъясняется во врезке “Понятие ориентации обучения” главы 13. Чтобы научить искусственный интеллект чему-то новому, совершенно другому и удивительному, человек должен обеспечить соответствующую ориентацию данных.

Воображение

Творчество означает определение чего-то реального, будь то музыка, живопись, литература, или любое другое действие, которое приводит к созданию чего-то, что другие могут увидеть, услышать, коснуться или ощутить другими способами. Воображение — это абстракция творчества, а потому оно еще дальше от диапазона возможностей искусственного интеллекта. Некто может вообразить такие вещи, которые не реальны и никогда не могут быть реальны. Воображение — это сознание, блуждающее в полях стремлений, играющего с тем, что могло бы быть, если бы правила не мешали. Истинный творческий потенциал, как правило, является результатом хорошего воображения.

С чисто человеческой точки зрения, все могут вообразить что угодно. Кроме того, воображение зачастую ставит нас мысленно в такие ситуации, которые не реальны вообще. Статья *Huffington Post* https://www.huffingtonpost.com/lamisha-serfwalls/5-reasons-imagination-is-_b_6096368.html повествует о пяти причинах критической важности воображения в преодолении пределов действительности.

Подобно тому как искусственный интеллект не может создать новый образ мышления или разработать новые данные, не используя существующие источники, он еще должен существовать в пределах границ действительности. Следовательно, маловероятно, что некто когда-нибудь разработает искусственный интеллект с воображением. Мало того, что воображение требует творческого интеллекта, оно также требует внутриличностного интеллекта, а искусственный интеллект такой формой интеллекта не обладает.



СОВЕТ

Воображение эмоционально, как и многие другие черты человека. Искусственный интеллект не имеет эмоций. Фактически, сравнивая то, на что способен искусственный интеллект и человек, имеет смысл задаться простым вопросом, требует ли эта задача эмоций.

Оригинальные идеи

Чтобы вообразить нечто или создать нечто реальное из возникшего в воображении, или использовать пример реальной ситуации для того, чего никогда не было в прошлом, должна возникнуть идея. Успешное создание идей — это способность человека с хорошим творческим, внутриличностным и межличностным интеллектом. Придумывание чего-то нового замечательно, если вы хотите помечтать о чем-то или развлечь себя. Но чтобы превратить это в идею, ею следует поделиться с другими, чтобы о ней могли узнать.

Неточности данных

В разделе “Пять недостоверностей данных” главы 2 рассматриваются проблемы с данными, которые должен преодолеть искусственный интеллект, чтобы решать задачи, для которых он предназначен. Единственная проблема заключается в том, что обычно искусственный интеллект не может с легкостью распознавать недостоверности в данных, если ему не предоставлено разнообразие примеров данных, в которых есть пропуски или недостоверности, что может быть куда труднее, чем вы думаете. Люди, напротив, способны определять недостоверности с относительной легкостью. Повидав куда больше примеров, чем любой искусственный интеллект когда-либо увидит, человек может выявить недостоверности, используя как воображение, так и творческий потенциал. Человек может вообразить такую недостоверность, которую

искусственный интеллект не сможет вообразить никак, поскольку в реальности искусственный интеллект ничего не понимает.



ЗАПОМНИ

Недостоверности попадают в данные столь многими способами, что их даже не перечислить.

Недостоверности зачастую вносят люди, даже непреднамеренно. Фактически избежать недостоверностей может быть невозможно, поскольку они зависят от точки зрения, а она может со временем меняться. Поскольку искусственный интеллект не может выявить все недостоверности, у используемых им для принятия решения данных всегда будет некая неточность. Повлияют ли эти неточности на возможность искусственного интеллекта сформировать полезный результат, зависит от вида и уровня неточности, а также от возможностей используемых алгоритмов.

Самый странный вид неточности данных, который следует учитывать, — это когда человек сам хочет повлиять на результат. Эта ситуация возникает чаще, чем думает большинство людей, и единственный способ преодолеть эту специфическую для людей проблему связан с наличием межличностного интеллекта, которого искусственный интеллект лишен. Например, некто покупает новый комплект одежды. Она выглядит отвратительно, по крайней мере для вас (ведь одежда бывает на удивление субъективной). Но если вы достаточно интеллектуальны, то скажете, что одежда выглядит непривычной. Человеку не нужно ваше объективное мнение, ему нужны ваши поддержка и одобрение. Вопрос ставится не как “Как эта одежда выглядит?” (искусственный интеллект на него не ответит), а как “Вы меня одобряете?” или “Поддерживаете ли вы мое решение купить эту одежду?” Частично эту проблему можно преодолеть предложением аксессуаров, которые дополняют одежду, или других средств, таких как замечание, что эту одежду они ведь не будут носить публично.

Есть также проблема с высказыванием горькой правды, которую искусственному интеллекту никогда не понять, поскольку у него нет эмоций. *Горькая правда* (hurtful truth) — это когда тот, кому ее говорят, не получает от нее ничего полезного, а только вред, эмоциональный, физический или интеллектуальный. Например, ребенок может не знать, что один из его родителей был неверен другому. Поскольку оба родителя пережили это, информация больше не является актуальной и лучше позволить ребенку оставаться в состоянии счастливого неведения. Подробное обсуждение супружеской неверности родителей наверняка нанесет вред психике ребенка.

Ребенок ничего не выиграет от разоблачения, но, определенно, получит вред. Искусственный интеллект может стать причиной вреда того же вида, разгласив семейную информацию способами, в которых ребенок не будет

учитываться. Например, обнаружив неточности в комбинации полицейских отчетов, гостиничных записях, магазинных чеках и других источниках, искусственный интеллект сообщает ребенку об изменах его родителей, нанося вред, оглашая правду. Однако в случае искусственного интеллекта правда разглашается из-за нехватки эмоционального интеллекта (сочувствия); искусственный интеллект не способен понять потребность ребенка в блаженном неведении о поведении родителей. К сожалению, даже когда набор данных для искусственного интеллекта содержит достаточно правильную и правдивую информацию, чтобы выдать пригодный для использования результат, результат может оказаться скорее вредным, чем полезным.

Неправильное применение искусственного интеллекта

Пределы искусственного интеллекта определяют область его возможного применения. Но даже в рамках этой области вы можете получить неожиданный или бесполезный вывод. Например, вы можете обеспечить искусственный интеллект различными исходными данными, а затем запросить вероятность определенных событий на основании этих данных. Когда доступно достаточное количество данных, искусственный интеллект может прийти к выводу, который соответствует математическому основанию исходных данных. Однако искусственный интеллект не может производить новые данные, чтобы создать решения на их основании, вообразить новые способы работы с этими данными или предложить идеи для реализации решения. Все эти роли принадлежат человеку. Все, чего вы можете ожидать, является вероятностным прогнозом.



ЗАПОМНИ

Большинство результатов искусственного интеллекта создается на основании вероятности или статистики. К сожалению, ни один из этих математических методов не применим к отдельным личностям; эти методы работают только с группами. Фактически использование статистики создает бесчисленные проблемы практически в любом деле, кроме конкретного вывода, такого как вождение автомобиля. В статье <https://public.wsu.edu/~taflinge/evistats.html> обсуждаются проблемы использования статистики. Когда ваше приложение искусственного интеллекта затрагивает личности, вам следует быть готовым к неожиданностям, включая полную неудачу в достижении любой из поставленных задач.

Другая проблема заключается в том, содержит ли набор данных некий вид мнения, которое распространено куда шире, чем можно было подумать. Мнение отличается от факта тем, что факт полностью доказуем и все соглашаются с тем, что факт — это правда (по крайней мере, люди без предубеждений).

Мнения вмешиваются тогда, когда в основе данных нет достаточных научных фактов. Кроме того, мнения вмешиваются тогда, когда задействованы эмоции. Даже встретившись с противоречащим мнению заключительным доказательством, некоторые люди продолжают полагаться на мнение, а не на факт. Мнение позволяет нам чувствовать себя уверенно, а факт — нет. Искусственный интеллект будет почти всегда терпеть неудачу, когда присутствует мнение. Даже наилучший алгоритм оставит кто-то недовольным выводом.

Мир нереалистичных ожиданий

В предыдущих разделах главы упоминается, что ожидания того, что искусственный интеллект выполнит определенные задачи или будет применен в конкретных ситуациях, вызывают проблемы. К сожалению, люди, похоже, не понимают, что многие задачи, которые предположительно будет решать искусственный интеллект, никогда не возникнут. Эти нереалистичные ожидания имеют много источников, включая следующие.

- » **Средства массовой информации.** Книги, фильмы и другие формы средств массовой информации, рассчитывающие на эмоциональный отклик от нас. Но этот эмоциональный отклик сам является источником нереалистичных ожиданий.

Мы предполагаем, что искусственный интеллект может сделать нечто, на что он в действительности не способен.

- » **Антропоморфизация.** Наряду с эмоциями, создаваемыми средствами массовой информации, люди склонны на все накладывать свои привязанности. Люди нередко дают имена своим автомобилям, говорят с ними и задаются вопросом, плохо ли они себя чувствуют, когда повреждаются. Искусственный интеллект не может чувствовать, не может понимать, не может общаться (на самом деле), не может делать ничего, кроме обработки чисел — больших и еще больших количеств чисел. Когда ожидание заключается в том, что у искусственного интеллекта внезапно появятся чувства и способность действовать, как человек, результат обречен на неудачу.
- » **Проблема неопределенности.** Искусственный интеллект вполне может справиться с определенными проблемами, но не с неопределенностью. Вы можете предоставить человеку набор потенциальных входных данных и ожидать, что человек создаст соответствующий вопрос на основании экстраполяции. Скажем, что набор тестов регулярно терпит неудачу, по большей части, но некоторые тесты действительно достигают желаемого результата. Искусственный интеллект мог бы попытаться улучшать результаты проверки за счет интерполяции, приближая характеристики новых тестов к тем,

которые закончились успехом. Но человек мог бы улучшить результаты тестов, задавшись вопросом, почему некоторые тесты преуспели найти причины и успеха, и неудачи (возможно, разнятся условия окружающей среды или объекты теста). Чтобы искусственный интеллект мог решить некую проблему, человек должен выразить ее таким способом, который понятен искусственному интеллекту. Проблемы неопределенности, решения которых нет в человеческой практике, просто не решаются и с использованием искусственного интеллекта.

- » **Несовершенная технология.** Во многих местах этой книги упоминается, что проблема не была решена в определенное время из-за несовершенства технологии. Нереально ожидать от искусственного интеллекта решения задачи, потому что его технология несовершенна. Например, отсутствие сенсоров и достаточной вычислительной мощности сделало бы создание беспилотного автомобиля в 1960-х годах невозможным, а современные достижения в технологиях сделали такие усилия вполне возможными.

Влияние зим искусственного интеллекта

Зимы искусственного интеллекта наступают, когда ученые анонсируют преимущества искусственного интеллекта, которые оказываются неосуществимыми за ожидаемый период времени, что приводит к замораживанию финансирования искусственного интеллекта, и его исследования продолжаются в темпе ледника. С 1956 года мир видел две зимы искусственного интеллекта. (Прямо сейчас мир переживает третье лето искусственного интеллекта.) В следующих разделах подробно обсуждаются причины, влияние и последствия зим искусственного интеллекта.

Понятие зимы искусственного интеллекта

Довольно трудно сказать точно, когда появился искусственный интеллект. В конце концов, даже древние греки мечтали о создании механических людей, как в мифах о Гефесте и Галатее Пигмалиона, и мы можем предположить, что у этих механических людей был бы своего рода интеллект. Следовательно, можно утверждать, что первая зима искусственного интеллекта фактически произошла между падением Римской империи и Средневековьем, когда люди начали мечтать об алхимическом способе создания разума — Таквин Джабир ибн Хайяна, Гомункул Парацельса и Голем раввина Йехуда Лёв бен Бецаделя из Праги.

Однако успех этих усилий не подтвержден историческими документами. Искусственный интеллект был создан позже, в 1956 году в Дартмутском колледже при правительственном финансировании исследований.



ЗАПОМНИ

Зима искусственного интеллекта начинается, когда заканчивается финансирование его исследований. Термин *зима* очень подходит, поскольку рост искусственного интеллекта, как и деревьев зимой, прекращается в целом.

Глядя на кольца дерева, вы видите, что дерево продолжает расти и зимой, только не очень быстро. Аналогично на протяжении зим искусственного интеллекта с 1974 по 1980 и с 1987 по 1993 год искусственный интеллект действительно продолжал расти, но с ледниковой скоростью.

Причины зим искусственного интеллекта

Причиной зим искусственного интеллекта вполне можно считать захватывающие обязательства, которые невозможно выполнить. В начале исследований Дартмутского колледжа в 1956 году лидеры разработки искусственного интеллекта прогнозировали, что компьютер станет столь же интеллектуальным, как человек, не более чем через поколение. Сейчас, шестьдесят с лишним лет спустя, интеллект компьютеров еще весьма далек от человеческого. Фактически, если вы читали предыдущие главы, то знаете, что компьютеры вряд ли когда-либо станут столь же интеллектуальными, как люди, по крайней мере не в каждом виде интеллекта (в настоящее время они превзошли человеческие возможности только в очень немногих видах интеллекта).



ЗАПОМНИ

Часть проблемы со сверхперспективными возможностями в том, что первые сторонники искусственного интеллекта полагали, что весь процесс человеческого мышления может быть формализован в алгоритмы. Фактически это идея китайских, индийских и греческих философов. Но как демонстрируется в табл. 1.1 из главы 1, формализованными могут быть лишь некоторые компоненты человеческого интеллекта. В наилучшем случае получится выяснить механизмы математического и логического мышления человека. Но даже в этом случае Давид Гильберт (David Hilbert) бросил в 1920- и 1930-х годах вызов математикам, попытавшись доказать, что все математические рассуждения могут быть формализованы. Ответ на этот вызов поступил из доказательства неполноты Курта Гёделя, машины, использованной Тьюрингом, и лямбда-исчислений Алонзо Чёрча. Появились два результата: формализация *всего* математического рассуждения невозможна; а в областях, в которых формализация возможна, вы можете механизировать рассуждение, которое является основанием для искусственного интеллекта.

Другая часть проблемы сверхобещаний — в чрезмерном оптимизме. На заре искусственного интеллекта компьютеры решали алгебраические задачи, доказывали геометрические теоремы и учились говорить на английском языке. Первые два вывода вполне разумны, когда вы полагаете, что компьютер просто анализирует ввод и приводит его в форму, которой может манипулировать компьютер. Проблема с третьим выводом. В действительности компьютер не умел говорить на английском; он преобразовывал текстовые данные в цифровые шаблоны, которые, в свою очередь, преобразовывались в аналоговые шаблоны, а результат весьма походил на речь, но это было не так. Компьютер ничего не понимал в английском языке, как, впрочем, и в любом другом языке. Да, ученые действительно слышали английскую речь, но компьютер видел просто нули и единицы в определенном шаблоне, который компьютер вообще не воспринимал как язык.



ВНИМАНИЕ!

Даже исследователи зачастую заблуждаются, полагая что компьютер делает больше, чем на самом деле. Например, ELIZA Джозефа Вейценбаума казалась слышащей сказанное и отвечавшей осмысленно. К сожалению, ответы были заданы заранее, и приложение не слышало ничего, не понимало и не высказывало. Но все же ELIZA была первым виртуальным собеседником и действительно шагом вперед, хотя и невероятно маленьким. Просто заблуждение было значительно более великим, чем фактическая технология — в этом и проблема, перед которой стоит сегодня искусственный интеллект. Люди чувствуют себя разочарованными, когда открывается обман, поэтому ученые и их покровители продолжают терпеть неудачи, демонстрируя блеск, а не реальную технологию. Первую зиму искусственного интеллекта вызвали следующие прогнозы ученых.

- » **Герберт Саймон** (H.A. Simon). “Через десять лет цифровой компьютер станет чемпионом по шахматам в мире” (1958) и “через двадцать лет машины смогут выполнять любые работы, на которые способен человек” (1965).
- » **Аллен Ньюэлл** (Allen Newell). “Через десять лет цифровой компьютер сформулирует и докажет новую важную математическую теорему” (1958).
- » **Марвин Ли Минский** (Marvin Minsky). “В пределах одного поколения... проблема создания искусственного интеллекта будет в основном решена” (1967) и “Через три-восемь лет у нас будет машина с общим интеллектом среднего человека” (1970).

Сегодня эти заявления выглядят забавно; вполне понятно, почему правительства закрыли их финансирование. В разделе “Аргумент китайской комнаты”

главы 5 описаны лишь некоторые из многих контрдоводов, выдвигаемых даже людьми из сообщества искусственного интеллекта против этих прогнозов.

Вторая зима искусственного интеллекта наступила в результате тех же самых проблем, которые вызвали первую — сверхобещания, перевозбуждение и чрезмерный оптимизм. В данном случае бум начался с экспертной системы, своего рода программы искусственного интеллекта, решавшей задачи, используя логические правила. Кроме того, вмешались японцы, выступив со своим проектом Fifth Generation Computer — компьютерной системой, продемонстрировавшей вычисления с массовым параллелизмом. Идея состояла в создании компьютера, способного выполнять множество задач параллельно, как человеческий мозг. И наконец, Джон Хопфилд (John Hopfield) и Дэвид Румельхарт (David Rumelhart) возродили коннекционизм, стратегию моделирования процессов мышления как объединенной сети простых модулей.

Конец наступил в виде финансового мыльного пузыря. Экспертные системы оказались хрупкими, даже работая на специализированных компьютерных системах. Специализированные компьютерные системы закончились с перекрытием финансовых шлюзов, более новые общедоступные компьютерные системы вполне могли бы заменить их за значительно меньшую стоимость. Фактически японский проект Fifth Generation Computer также провалился из-за экономических факторов. Его, оказалось, чрезвычайно дорого создать и поддерживать.

Пересмотр ожиданий и новые цели

Зима искусственного интеллекта необязательно оказывается разрушительной. Совсем наоборот.

Такие времена могут позволить отступить и подумать о различных проблемах, которые возникли во время порыва, а также разработать нечто удивительное. Из первой зимы искусственного интеллекта извлекли пользу два главных направления (при незначительном перевесе одного из них).

- » **Логическое программирование.** Это направление задействует набор предложений в логической форме (выполняемый как приложение), который выражает факты и правила специфической предметной области.

Примерами языков программирования, использующих эту конкретную парадигму, являются Prolog, Answer Set Programming (ASP) и Datalog. Это форма программирования на базе правил является основной технологией, используемой для экспертных систем.

- » **Рассуждение на основе здравого смысла.** Это направление использует метод моделирования человеческой способности

прогнозировать результат последовательности событий на основании свойств, цели, намерения и поведения конкретного объекта. Рассуждение на основе здравого смысла — это немаловажный компонент искусственного интеллекта, поскольку он затрагивает широкое разнообразие дисциплин, включая компьютерное зрение, робототехническую манипуляцию, таксономическое рассуждение, действие и изменение, временное рассуждение и качественное рассуждение.

Вторая зима искусственного интеллекта внесла дополнительные изменения, которые привлекли к искусственному интеллекту то внимание, которым он пользуется до сих пор. Эти изменения таковы.

- » **Использование обычных аппаратных средств.** Ранее использование экспертных систем и других средств искусственного интеллекта полагалось на специализированные аппаратные средства. Причина была в том, что обычные аппаратные средства не обеспечивали необходимой вычислительной мощности и памяти. Однако эти уникальные системы оказались слишком дорогими, сложными в обслуживании и программировании, а также чрезвычайно ненадежными, когда сталкивались с необычными ситуациями. Обычные аппаратные средства являются универсальными по природе и менее склонными к проблемам при поиске решения (см. раздел “Решения в поиске задачи” далее в главе).
- » **Осознание необходимости в обучении.** Экспертные системы и другие ранние формы искусственного интеллекта требовали, чтобы специальное программирование удовлетворяло всем потребностям, что делало их чрезвычайно негибкими. Стало очевидным, что компьютеры должны быть в состоянии обучаться на окружающей обстановке, используя сенсоры и предоставленные данные.
- » **Гибкость окружающей среды.** Системы, которые действительно выполняли полезную работу между первой и второй зимами искусственного интеллекта, работали жестко. Когда входные данные совсем не соответствовали ожиданиям, эти системы давали гротескные ошибки в выводе. Стало очевидно, что любые новые системы должны будут знать, как реагировать на реальные данные, которые полны ошибок, частично отсутствуют или неправильно оформлены.
- » **Доверие новым стратегиям.** Предположим, что вы работаете на правительство и пообещали всякого рода удивительные вещи от разрабатываемого искусственного интеллекта, но осуществить ни одну из них не удалось. Это причина второй зимы искусственного интеллекта: различные правительства опробовали различные

способы сделать обещания искусственного интеллекта реальностью. Когда текущие стратегии однозначно не сработали, эти же самые правительства начали искать другие способы продвижения вычислительной техники, некоторые из которых привели к интересным результатам, таким как достижения в робототехнике.

Дело в том, что зимы искусственного интеллекта не обязательно для него плохи. Фактически это шансы отстраниться и пересмотреть прогресс (или его отсутствие).

Довольно трудно задумываться обо всех этих важных моментах, когда все несется сломя голову к очередному успеху.



ЗАПОМНИ

При рассмотрении зим искусственного интеллекта и последующего возобновления разработок искусственного интеллекта с модифицированными идеями и целями имеет смысл вспомнить высказывание американского ученого и футуриста Роя Чарльза Амары (Roy Charles Amara) (известное также как закон Амары): “Мы склонны переоценивать эффект технологии в краткосрочной перспективе и недооценивать в долгосрочной”. После всех обманов и разочарования всегда есть время, когда люди не могут точно оценить долгосрочные последствия новой технологии и понять вызванные ею революции. Как технология искусственный интеллект изменит наш мир к лучшему или к худшему независимо от того, какое количество зим еще предстоит пережить.

Решения в поиске задачи

Два человека смотрят на грудку проводов, колес, кусков металла и разных странных предметов, выглядящую как куча мусора. Первый человек задается вопросом “Что это делает?”

Второй отвечает “А чего оно не делает?” Таким образом, изобретение, делающее все что угодно, на самом деле не делает ничего вообще. Средства массовой информации изобилуют примерами решений, ищущих свою задачу. Мы смеемся, поскольку все мы встречались с решением, для которого осталось только найти задачу. Эти решения заканчивают на свалке, даже если они действительно работают, поскольку они не являются ответом на насущную потребность. В следующих разделах подробнее обсуждаются решения на базе искусственного интеллекта, находящиеся в поисках задачи.

Определение примочки

Когда дело доходит до искусственного интеллекта, мир полон *примочек* (gizmo). Одни из них действительно полезны, другие — нет, а третьи попадают в промежуток между ними. Например, Alexa имеет много полезных возможностей, но обладает также солидным запасом примочек, которые заставляют хорошо почесать голову, если попытаться их использовать. Следующая статья Джона Дворака (John Dvorak) может показаться слишком пессимистичной, но она предоставляет пищу для размышлений о видах возможностей Alexa: <https://www.pcmag.com/commentary/354629/just-say-no-to-amazons-echo-show>.

Примочка искусственного интеллекта (AI gizmo) — это любое приложение, на первый взгляд кажущееся интересным, но в конечном счете оказывающееся неспособным ни на что. Вот некоторые из наиболее популярных аспектов, которые следует учитывать при определении, не является ли нечто примочкой. (Первые буквы каждого пункта списка складываются в аббревиатуру “CREEP”, намекая на то, что не стоит создавать жутких (сгееру) приложений искусственного интеллекта.)

- » **Доступная цена** (cost effective). Прежде чем некто решит купить приложение искусственного интеллекта, он должен убедиться, что оно не дороже существующего решения. Все ищут выгоду. Переплата за те же самые преимущества просто не привлечет внимания.
- » **Воспроизводимость** (reproducible). Результаты приложения искусственного интеллекта должны быть воспроизводимыми, даже если изменяются обстоятельства решения задачи. В отличие от процедурных решений задачи люди ожидают, что искусственный интеллект будет адаптироваться — учиться на собственном опыте, а значит, планка воспроизведения результатов поднимается выше.
- » **Экономность** (efficient). Когда решение для искусственного интеллекта внезапно расходует огромные объемы ресурсов любого вида, пользователь ищет в другом месте. Сейчас компании весьма заинтересованы решением задач с минимально возможной тратой ресурсов.
- » **Эффективность** (effective). Простого предоставления дешевого и экономного практического преимущества недостаточно; искусственный интеллект должен обеспечить полное решение задачи. Эффективные решения позволяют автоматизировать выполнение неких задач без необходимости регулярно перепроверять результаты или поддерживать автоматизацию.
- » **Практичность** (practical). Полезное приложение должно приносить практическую пользу. Создаваемое преимущество должно быть нужным конечному пользователю, например доступ к дорожной карте или напоминание о принятии лекарств.

Реклама

Шумиха среди потенциальных пользователей вашего приложения искусственного интеллекта является верным признаком того, что приложение потерпит неудачу. Достаточно странно, но легче всего добиваются успеха те приложения, цели и намерения которых очевидны с самого начала. Цели приложения распознавания речи очевидны: вы говорите, и компьютер делает нечто полезное. Вам не нужно никого убеждать в полезности программного обеспечения распознавания речи. В этой книге упоминается множество действительно полезных приложений, ни одно из которых не требует навязчивой рекламы. Если люди начинают спрашивать, что нечто делает, значит, пришло время заново продумать проект.

Когда люди справляются лучше

Даже при использовании искусственного интеллекта люди все равно остаются в деле. В этой книге нередко упоминается, что в некоторых задачах люди добиваются большего успеха, чем искусственный интеллект, а многие задачи искусственный интеллект не может решать вообще.

Все, что требует воображения, творческого потенциала, проницательности, выработки мнения или идеи, лучше оставить людям. Достаточно странно, но ограниченность искусственного интеллекта оставляет много места для людей, многие из которых сегодня не заняты, поскольку люди чрезвычайно заняты в повторяющихся, скучных задачах, которые легко мог бы выполнять искусственный интеллект.

Заглянем в будущее, где искусственный интеллект действует как помощник людей. Фактически со временем вы будете видеть случаи использования искусственного интеллекта все чаще. Наилучшими приложениями искусственного интеллекта будут те, которые смогут помогать людям, а не заменять их. Да, роботы заменят людей в опасных условиях, но решения должны будут принимать люди, чтобы не усугубить ситуацию, а значит, для работы роботу нужен человек в безопасном месте. Это взаимовыгодное сотрудничество между технологией и людьми.

Поиск простого решения

Принцип *не усложняй* (Keep It Simple, Stupid — KISS) является наилучшим, когда дело доходит до развития приложений искусственного интеллекта. Вы можете прочитать об этом принципе больше по адресу <https://www.techopedia.com/definition/20262/keep-it-simplestupid-principle-kiss-principle>, но главная идея в том, что наилучшим является самое простое решение.

Существуют прецеденты для применения всякого рода простых решений. Но, вероятно, самым известным из них является Бритва Оккама (<https://science.howstuffworks.com/innovation/scientific-experiments/occams-razor.htm>).

Конечно, возникает вопрос, почему простота настолько важна. Самый простой ответ — сложность ведет к отказам: чем больше важных деталей, тем вероятнее отказ. Этот принцип имеет свои корни в математике и прост в доказательстве.



ЗАПОМНИ

Когда дело доходит до приложений, вступают в действие и другие принципы. Для большинства людей приложение — это лишь средство для достижения цели. Люди заинтересованы в конечном результате, а не в приложении. Если приложение исчезнет из виду, пользователь будет только доволен, поскольку ему важен только конечный результат. Простые приложения просты в применении, не особо бросаются в глаза и не требуют каких-либо сложных инструкций. На самом деле лучшие приложения очевидны. Если ваше решение искусственного интеллекта должно полагаться на всевозможные сложные взаимодействия, подумайте, не пришло ли время вернуться к чертежной доске и придумать что-то получше.

УЧЕТ ПРОМЫШЛЕННОЙ РЕВОЛЮЦИИ

Сотрудничество человека и искусственного интеллекта не случится так сразу. Кроме того, новые виды работ, которые будут в состоянии выполнить только люди, не появятся немедленно.

Однако ожидания людей, просто сидящих без дела и ждущих, что машины будут их обслуживать, несбыточны и, очевидно, неприемлемы. Люди продолжают выполнять различные задачи.

Конечно, те же самые заявления о машинах, заменяющих людей, были во все времена, главным образом во времена промышленных революций (см. http://www.historydoctor.net/Advanced%20Placement%20World%20History/40.%20The_Industrial_revolution.htm). Но люди всегда будут делать некоторые вещи лучше, чем искусственный интеллект, и вы можете быть уверены, что продолжите делать свое дело и не потеряете места в обществе. Следует только иметь в виду, что этот переворот будет менее мощным, чем промышленная революция.



Глава 16

Искусственный интеллект в космосе

В ЭТОЙ ГЛАВЕ...

- » Исследование Вселенной
- » Добыча ресурсов в космосе
- » Поиск новых мест для исследований
- » Создание строений в космосе

Люди наблюдали за небом с незапамятных времен. Большинство имен созвездиям и звездам дали еще древние греки и другие древние народы (в зависимости от того, где вы живете). Большой ковш тоже имеет много названий и может быть известен даже как “Медведь”, когда группируется с другими звездами (см. <https://newsok.com/article/3035192/various-cultures-offer-legends-regarding-big-dipper?>). Люди всегда любили смотреть на звезды и думать о них, а потому многие культуры задумывались о природе и виде звезд. Когда люди стали летать в космос, Вселенная в целом приобрела новый смысл, как описано в этой главе. Искусственный интеллект позволяет людям видеть Вселенную более ясно и смотреть на нее новыми способами.

Впоследствии люди начали жить в космосе (например, на Международной космической станции, https://www.nasa.gov/mission_pages/station/main/index.html) и посещать другие планеты, такие как Луна. Люди также начали работать в космосе. Конечно, в космосе можно создать такие материалы,

которые невозможно получить в других условиях. Компания Made In Space (<http://madeinspace.us/>) фактически специализируется на этом. Кроме всего этого, использование роботов и специализированного искусственного интеллекта позволяет наладить промышленную добычу в космосе всякого рода материалов. Фактически Конгресс США принял в 2015 году закон, разрешающий такую деятельность (<https://www.space.com/31177-space-mining-commercial-spaceflight-congress.html>), предоставив компаниям право на продажу добытого. В этой главе рассматривается также роль искусственного интеллекта в создании космической добывающей промышленности.

Вселенная полна почти бесконечным количеством тайн. Одна из недавно обнаруженных тайн — существование экзопланет вне Солнечной системы (см. <https://www.nasa.gov/feature/jpl/20-intriguing-exoplanets>). Существование экзопланет означает, что люди могут в конечном счете найти жизнь на других планетах, но даже поиск экзопланет требует искусственного интеллекта. Пути, которыми искусственный интеллект делает все это возможным, действительно удивительны.

Жить и работать в космосе — это одно, а провести отпуск в космосе — совсем другое. Уже в 2011 году люди начали говорить о возможности создания гостиницы на околоземной орбите (<http://mashable.com/2011/08/17/commercial-space-station/>) или Луне. Хотя постройка орбитальной гостиницы кажется вполне выполнимой на настоящий момент (<https://www.newsweek.com/spacex-takes-space-hotel-module-orbit-445616>), лунная гостиница кажется следующим большим делом (<http://www.bbc.com/future/story/20120712-where-is-hiltons-lunar-hotel>). Дело в том, что искусственный интеллект позволит людям жить, работать и даже проводить каникулы в специализированных космических строениях, как описано в этой главе.

Наблюдение за Вселенной

Изобретение телескопа приписывают голландскому производителю линз Иоанну Липперсгею (Hans Lippershey) в 1600 годах. (Фактически тема первооткрывателя телескопа обсуждается до сих пор <https://www.space.com/21950-who-invented-the-telescope.html>.) Такие ученые, как итальянский астроном Галилео Галилей, немедленно начали осматривать небеса, используя нечто лучшее, чем их глаза. Таким образом, телескопы распространялись и улучшались, становясь больше и сложнее, а недавно даже разместились в космосе.



ЗАПОМНИ

Причина размещения телескопов в космосе в том, что атмосфера Земли не позволяет получить четкие изображения слишком дальних объектов. Телескоп Хаббл — один из первых и самый известный

из космических телескопов (см. <https://www.nasa.gov/audience/forstudents/5-8/features/nasa-knows/what-is-the-hubble-space-telescope-58.html>). Как описано в следующих разделах, использование современных телескопов требует искусственного интеллекта, например для планирования времени использования Хаббла (см. <http://ieeexplore.ieee.org/document/63800/?reload=true>).

Первый ясный взгляд

Одним из способов выхода за пределы атмосферы Земли является размещение телескопа в космосе. Но этот подход дорог, а обслуживание является кошмаром. Большинство людей, смотрящих на небеса, нуждается в альтернативе, такой как телескоп, способный корректировать рассеивающее действие атмосферы Земли за счет деформации зеркала телескопа (см. <https://www.space.com/8884-telescope-laser-vision-heavens-blurry.html>).



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Представьте необходимость вычислять рассеивающий эффект атмосферы Земли на основании света от, скажем, лазера за доли секунды. Единственный способ осуществлять такие массивные вычисления, а затем перемещать исполнительные механизмы зеркала исключительно правильным способом заключается в том, чтобы использовать искусственный интеллект, который имеет весьма большой опыт по части выполнения таких вычислений, необходимых, чтобы сделать адаптивную оптику возможной. В статье <https://www.spiedigitallibrary.org/conference-proceedings-of-spie/2201/1/Artificial-intelligence-system-and-optimized-modal-control-for-the-ADONIS/10.1117/12.176120.short?SSO=1> приводится только один из примеров использования искусственного интеллекта в адаптивной оптике. Сайт <https://www.helsinki.fi/en/news/data-science/neural-networks-and-temporal-control-in-adaptive-optics> предоставляет дополнительные ресурсы по использованию нейронных сетей в адаптивных оптических системах.

Для обеспечения еще лучших оптических характеристик будущие телескопы будут осуществлять трехмерную коррекцию эффекта рассеивания с использованием мультисопряженной адаптивной оптики (<http://eso-ao.indmath.uni-linz.ac.at/index.php/systems/multi-conjugate-adaptive-optics.html>). Эта новая технология исправит узость поля зрения, присущую современным телескопам, но это потребует куда более точного управления всеми уровнями исполнительных механизмов нескольких зеркал. Новые телескопы, такие как Giant Magellan Telescope, Thirty-Meter Telescope и европейский European Extremely Large Telescope (см. <https://www.space.com>).

com/8299-world-largest-telescope-built-chile.html), полагаются на эту технологию, чтобы оправдать усилия по их созданию ценой более одного миллиарда долларов.

Поиск новых мест

До XVIII века люди были прикованы к поверхности Земли, но все равно пристально глядели в небо и мечтали. Они проделывали всякого рода странные эксперименты, например прыгали с башен (см. <https://jkconnectors.com/news/the-history-of-aviation-part-1/>), но несмотря на баллоны с горячим воздухом любой вид истинного полета казался недостижимым. Мы все еще остались исследователями; люди продолжают исследования и сейчас, находя все новые и новые места для поиска.



ЗАПОМНИ

Идея наличия мест для исследования действительно не ставилась до первой посадки на Луну 20 июля 1969 года (см. https://www.nasa.gov/mission_pages/apollo/apollo11.html). Да, мы могли смотреть, но не могли коснуться. Даже в этом случае с тех пор люди смотрели на всякого рода места и не могли достичь их, например Марс (<https://www.space.com/33468-viking-1-first-mars-landing-pictures.html>) и комета Розетта¹ (см. <https://www.usnews.com/news/articles/2014/11/12/rosetta-comet-landing-is-space-game-changer>). Каждое из этих исследований стимулирует человеческое желание пойти в новые места. Но что важнее всего, ни одно из этих предприятий не осуществилось бы без сложной математики, в которой искусственный интеллект может весьма помочь.

Поиск обычно полагается на телескопы. Однако НАСА и другие организации все более полагаются на другие подходы, такие как использование искусственного интеллекта, как описано в статье <https://www.astrobio.net/also-in-news/artificial-intelligence-nasa-data-used-discover-eighth-planet-circling-distant-star/>. В данном случае машинное обучение позволяло найти восьмую планету, вращающуюся вокруг звезды Kepler-90. Конечно, проблема поиска новых мест практически определяет, можем ли мы фактически достичь некоторых из более экзотических мест. Спутник Voyager 1, продвинувшийся дальше всех от Земли, только недавно достиг межзвездного пространства (<https://www.space.com/26462-voyager-1-interstellar-space-confirmed.html>). Его механизмы ломаются, но он все еще пригоден для использования (<https://www.nasa.gov/feature/jpl/voyager-1-fires-up-thrusters-after-37>). Находясь в 13 миллиардах миль (20,9215 млрд. км),

¹ Чисто технически имеется в виду комета Чурюмова–Герасименко. — *Примеч. ред.*

Voyager прошел только 0,0022 светового года за 40 лет. Звезда Kepler-90 находится на расстоянии 2 545 световых лет, поэтому без новой технологии, которая, возможно, будет создана с помощью искусственного интеллекта в будущем, достичь ее будет невозможно.



СОВЕТ

К счастью, наша Солнечная система содержит все виды мест, которые могли бы быть достижимы. Например, *Британская энциклопедия* рекомендует посетить такие места, как Равнина Жары (Caloris Planitia) на Меркурии (см. <https://www.britannica.com/list/10-places-to-visit-in-the-solar-system>). Вы также могли бы посетить сайт TravelTips4Life (<http://www.traveltips4life.com/15-places-we-want-to-visit-in-outer-space/>), который рекомендует Международную космическую станцию в качестве первой остановки.

Эволюция Вселенной

Люди наблюдают Вселенную довольно давно, но все еще не имеют никакого реального представления о том, какова она точно, кроме знания, что мы в ней живем. Конечно, наблюдения продолжаются, но сущность Вселенной — все еще большой вопрос. Недавно ученые начали использовать искусственный интеллект, чтобы тщательно прорисовать движение различных частей Вселенной и попытаться выяснить, как она работает (см. <https://www.sciencedaily.com/releases/2012/09/120924080307.htm>). Использование модели Лямбда-CDM (Lambda Cold Dark Matter — LCDM) для космоса поможет понять, как работает Вселенная, немного лучше. Однако она, вероятно, даже не начнет отвечать на все наши вопросы.

Новые научные принципы

В конечном счете исследования, в ходе которых люди больше узнают о космосе, Солнечной системе, Галактике и Вселенной, должно принести некую прибыль. В противном случае никто не захочет продолжать их финансирование. Зимы искусственного интеллекта, обсуждавшиеся в главе 15, являются примером того, что случается с технологией независимо от того, насколько великие обещания она оказалась не в состоянии выполнить. Следовательно, с учетом длинной истории исследования космоса люди должны получать некую пользу. В большинстве случаев эта польза приходит в виде новых научных принципов — улучшения понимания того, как нечто работает. Применив урок, полученный из космических исследований и путешествий, люди смогут сделать жизнь на Земле лучше. Кроме того, космические технологии зачастую находят свою реализацию в товарах повседневного использования.

Рассмотрим одно из исследований: посадку Аполлон-11 на Луну. Люди все еще ощущают последствия взрыва технологий, который произошел во время работ над этой задачей. Например, необходимость покорения космоса потребовала от правительства потратить много денег на такие технологии, как интегральные схемы (IC), что сегодня мы считаем само собой разумеющимся (см. <https://www.computerworld.com/article/2525898/app-development/nasa-s-apollo-technology-has-changed-history.html>). По разным данным каждый доллар, инвестированный правительством в исследования НАСА, сегодня приносит американцам от 7 до 8 долларов в виде товаров или услуг.

Однако космическая гонка создала новые технологии и кроме самих кораблей и всего с ними связанного. Например, фильм *Скрытые фигуры* (Hidden Figures) (<https://www.amazon.com/exec/obidos/ASIN/B01LTI1RHQ/datacserver/vip0f-20/>) представляет НАСА с такой точки зрения, о которой большинство людей даже не думают: все, что нужно математике, — это большая вычислительная мощь. В фильме вы видите развитие математики НАСА от человеческих вычислений до электронно-вычислительных машин. Но если смотреть фильм внимательно, то можно заметить, что передовые компьютеры работают рядом с человеком так, как искусственный интеллект будет работать рядом с людьми в будущем.



ЗАПОМНИ

Сегодня данные о космосе поступают отовсюду. Они помогают выработать новые научные принципы о вещах, которых мы не можем даже видеть, таких как *dark space*² (область космоса, обладающая массой, но не видимым присутствием) и *темная энергия* (неизвестная и неопознанная форма энергии, которая противодействует силе притяжения между телами в пространстве). Поняв эту невидимую сущность, мы получим новое знание о том, какие силы воздействуют на нашу планету. Но исследователи настолько перегружены данными, что им нужен искусственный интеллект только для того, чтобы выделить смысл хотя бы из малой их части (см. <https://www.theverge.com/2017/11/15/16654352/ai-astronomy-space-exploration-data>). Дело в том, что будущее космоса и использование созданных нами для него технологий зависят от применения всех данных, которые мы собираем, а для этого на настоящий момент требуется искусственный интеллект.

² Вероятно, автор имел в виду не блэк-метал/эмбиент-группу, а темную материю. — Примеч. ред.

Добыча ресурсов в космосе

Космическая добывающая промышленность привлекла более чем пристальное внимание в средствах массовой информации и научном сообществе. Такие фильмы, как *Чужой* (<https://www.amazon.com/exec/obidos/ASIN/B001AQO-3QA/datacservip0f-20/>), дают представление о том, как могло бы выглядеть будущее судно космических шахтеров. (Если повезет, космическая добывающая промышленность не будет иметь дела с враждебными чужими.) Куда практичнее перспективы, представленные в таких статьях, как <https://www.outerplaces.com/science/item/17125-asteroid-miningspace-era>. Фактически такие компании, как Deep Space Mining (<http://deepspaceindustries.com/mining/>³), уже изучают возможности добычи ресурсов в космосе. Что удивительно, ищут эти шахтеры в основном такие вещи, как вода, которая фактически весьма распространена здесь на Земле, но относительно редко встречается в космосе. В следующих разделах рассматривается суть некоторых наиболее интересных аспектов космической добывающей промышленности.

УЧЕТ КРИТИКИ

Лишь некоторые люди ценят роль критика в сообществе — вы наверняка знаете человека, который видит только пятна на Солнце, выбоины на каждой дороге и обратную сторону каждой медали. Критик — как старая сварливая карга, которую большинство средств массовой информации представляют как самый худший вид зла. Однако у критика действительно есть важная роль в искусственном интеллекте для космического применения. Конструктивная критика может повлиять на перспективное планирование, но она вряд ли поступит от позитивно настроенных членов группы. В то время как все остальные сосредоточены на творческом решении существующих проблем, критик видит будущие проблемы, которые действительно будут иметь значение, когда дело дойдет до реального применения приложений искусственного интеллекта, например при добыче ресурсов.

У космически-ориентированного искусственного интеллекта должно быть больше независимости, чем у любого земного аналога. С учетом различных тестов, проведенных на настоящий момент, становится очевидно, что предвидение непредвиденного является требованием, а не пожеланием. Космически-ориентированный искусственный интеллект должен быть способен обучаться исходя из окружающей обстановки и находить решения задач, о которых его разработчики даже не думали, например непредвиденные гравитационные

³ Актуальность ссылки не гарантируется. — *Примеч. ред.*

эффекты, отказы оборудования, отсутствие подходящих запасных частей и т.д. Следует учитывать и некоторые другие трудности, которые перед космическим искусственным интеллектом в настоящее время еще стоят, таких как попытки хакеров перехватить доставку. Для критически настроенного ума это достаточная пища для размышлений, поэтому он становится основной частью любой группы.

Во врезке “Понятие ориентации обучения” главы 13 также описаны важные уроки для космического искусственного интеллекта. Один из этих уроков — в тщетности, т.е. в знании, что сценарий безнадежен. Космический искусственный интеллект должен предусмотреть контрмеры, чтобы предотвратить возможный ущерб, а не пытаться затем устранить последствия, которые могут уже быть непоправимыми. В космосе бесконечное количество неизвестностей, а значит, потребуется человеческое вмешательство, но это вмешательство может иметь место месяцы спустя. Космический искусственный интеллект должен знать, как поддерживать ситуацию, пока вмешательство не станет возможным. Обсуждение на <https://worldbuilding.stackexchange.com/questions/66698/what-issues-would-an-ai-asteroid-mining-stations-have-to-be-prepared-for> предоставляет лишь небольшой набор примеров невероятных проблем, с которыми может столкнуться космический искусственный интеллект.

Добыча воды

Вода покрывает примерно 71 процент поверхности Земли. Фактически на Земле так много воды, что нам зачастую трудно убрать ее из тех мест, где она не нужна. Однако Земля — исключение из правил. В космосе изобилия воды нет. Конечно, вы могли бы задаться вопросом, зачем нужна вода в космосе, кроме как для нужд астронавтов и, возможно, полива растений. Факт в том, что из воды получается отличное ракетное топливо. Разделение H_2O на составляющие компоненты дает водород и кислород, являющиеся компонентами современного ракетного топлива (см. https://www.nasa.gov/topics/technology/hydrogen/hydrogen_fuel_of_choice.html). Следовательно, большой шар грязного льда в небе может стать в будущем заправочной станцией.

Добыча редкоземельных и других металлов

Горнодобывающая промышленность всегда была делом грязным, но одни ее виды куда грязнее других, и добыча редкоземельных металлов относится именно к этой категории. Добыча редкоземельных металлов — действительно грязное дело (см. <https://earthjournalism.net//stories/the-dark-side-of-renewable-energy> и <https://tucson.com/business/local/big-pollution->

Вы не сможете определить, что содержит астероид, пока не подберетесь к нему действительно близко. Кроме того, количество требующих исследования астероидов, прежде чем на них удастся найти что-нибудь стоящее, весьма существенно — намного больше, чем смогут исследовать пилотируемые корабли. Кроме того, нахождение рядом с любым объектом, способным вращаться непредсказуемым способом и иметь странные характеристики, довольно опасно. По всем этим причинам большинство исследований астероидов в целях добычи ресурсов будет осуществляться с использованием автономных дронов различных видов. Эти дроны пойдут от астероида к астероиду, ища необходимые ресурсы. Когда дрон найдет необходимый материал, он предупредит центральную станцию, передав информацию о точном расположении астероида и других характеристиках.

Затем будет послан робот, чтобы сделать что-то с астероидом. Большинство людей полагает, что добыча ресурсов будет осуществляться на месте, но фактически это слишком опасно и дорого. Другая идея заключается в том, чтобы переместить астероид в более безопасное место, такое как орбита вокруг Луны, и выполнять добычу там. Фактически для этого проекта уже есть финансирование (см. <https://www.outerplaces.com/science/item/1332-nasa-to-receive-100m-budget-for-asteroid-capture>). Дело в том, что одни роботы могли бы осуществлять транспортировку, а другие — добычу ресурсов. Люди могли бы участвовать в ремонте роботов и, вероятно, в контроле действий дронов и роботов. Это куда менее опасная и более интересная работа, чем в горнодобывающей промышленности здесь на Земле.

risk-seen-in-rare-earth-mining/article_c604dd80-7a8d-5ab5-8342-0f9b-8dbb35fb.html), поскольку все их месторождения на территории США были закрыты как стратегический резерв для вооруженных сил, пока американское правительство не позволило вновь открыть месторождение Mountain Pass из-за китайцев (см. <https://www.popularmechanics.com/science/a12794/one-american-mine-versus-chinas-rare-earths-dominance-14977835/>). Одним из самых больших недостатков добычи редкоземельных металлов является загрязнение окружающей среды радиоактивным торием.

Из-за очень высокой цены, а также опасности для экологии и рабочих длительная добыча редкоземельных металлов в США на месторождении Mountain Pass вызывает большое сомнение (см. <https://www.bloomberg.com/news/articles/2017-06-22/the-distressed-debt-standoff-over-america-s-only-rare-earth-mine>). Фактически с китайцами ведется борьба, чтобы не позволить им купить единственное американское месторождение (см. <http://>

thehill.com/blogs/pundits-blog/economy-budget/339528-Rare-earth-rancor%3A-Feds-must-stop-Chinese-purchase-of-US-mine и <https://www.breitbart.com/economics/2017/08/30/exclusive-donald-trump-urged-to-nationalize-americas-only-rare-earth-mine/>).

Ваш мобильный телефон, iPad, автомобиль, телевизор, а также солнечная батарея или ветряная мельница, дающий электричество в ваш дом, — все используют чрезвычайно опасные редкоземельные материалы (см. <http://www.rareearthtechalliance.com/Applications/Electronics.html>). Большинство людей даже не знают, что эти материалы неустойчивы, из-за способа их современного использования (<http://www.pbs.org/wgbh/nova/next/physics/rare-earth-elements-in-cell-phones/>). С учетом длины списка подобных полезных ископаемых они являются наилучшей причиной для развития добычи на других планетах, где токсичные отходы не будут нам мешать. Фактически горная промышленность должна быть только первым этапом; все производство также должно покинуть планету (да, ее потенциал для загрязнения окружающей среды весьма велик).



ЗАПОМНИ

Искусственный интеллект очень важен для стараний по поиску лучших источников редкоземельных металлов, которые не будут загрязнять нашу планету. Как ни странно, одним из потенциальных источников редкоземельных металлов является Луна (см. https://www.washingtonpost.com/national/health-science/moon-draws-growing-interest-as-a-potential-source-of-rareminerals/2012/01/30/gIQAqHvUuQ_story.html?utm_term=.828c9cb19a34). Фактически многие политические деятели видят сейчас лунную горнодобывающую промышленность редкоземельных металлов как стратегическую задачу (см. <https://sservi.nasa.gov/articles/is-mining-rare-minerals-on-the-moon-vital-to-national-security/>). Проблема в том, что усилия по изучению точного строения Луны пока не были в целом успешны, поэтому важно знать, чего ожидать. Программа Moon Mineralogy Mapper (<https://www.jpl.nasa.gov/missions/moon-mineralogy-mapper-m3/>) является только реализацией одного из многих усилий по исследованию строения Луны. Кроме того, чтобы успешно обрабатывать редкоземельные металлы и превращать их в полезные товары, на Луне потребуется источник воды, которой на ней, очевидно, нет (см. <https://news.nationalgeographic.com/2017/07/water-moon-formed-volcanoes-glass-space-science/>). Зонды, роботы, анализ данных и все необходимое планирование потребуют использования искусственного интеллекта, поскольку проблемы куда сложнее, чем можно подумать.

Поиск новых элементов

Периодическая таблица содержит список всех доступных элементов и постоянно обновляется. Фактически в 2016 году в ней появились четыре новых элемента (см. <https://www.sciencenews.org/blog/science-ticker/four-newest-elements-periodic-table-get-names>). Однако поиск этих четырех новых элементов потребовал работы как минимум ста ученых, использовавших передовой искусственный интеллект (см. <https://www.wired.com/2016/01/smashing-new-elements-into-existence-gets-a-lot-harder-from-here/>), поскольку в лабораторных условиях они существуют лишь долю секунды. Достаточно интересно, что космос может быть тем местом, где эти новые элементы существуют в естественных условиях и не долю секунды, поскольку протоны в ядре отталкивают друг друга.



ЗАПОМНИ

Как демонстрируется в этой статье, мы все еще находим новые элементы для периодической таблицы и почти наверняка найдем еще больше. Сверхновые звезды и другие космические явления могут создавать элементы, которые ученые воспроизводят в ускорителях элементарных частиц и реакторах (<http://discovermagazine.com/2014/sept/3-ask-discover>). Фактически физики-ядерщики использовали искусственный интеллект в работе начиная с 1980-х годов (см. <http://www.sciencemag.org/news/2017/07/ai-changing-how-we-do-science-get-glimpse>). Вас это может удивить, но один из элементов, технеций, встречается только в космосе (<https://www.forbes.com/sites/ethansiegel/2015/08/01/a-periodic-table-surprise-the-one-element-in-stars-that-isnt-on-earth/#2bfc534edf74>).

Объединение элементов создает новые материалы. Искусственный интеллект непосредственно помогает химикам находить новые способы комбинирования элементов в новые интересные кристаллы (см. <https://www.sciencedaily.com/releases/2016/09/160921084705.htm>). Однажды ученые обнаружили 2 миллиона новых видов кристаллов, использующих только четыре элемента, но их поиски полагались на искусственный интеллект. Только представьте, что будет в будущем, когда ученые начнут использовать искусственный интеллект и глубокое обучение (которое позволит определить, будут ли полученные кристаллы фактически полезными).

Улучшение коммуникаций

Любое столь сложное космическое предприятие, как добыча ресурсов, требует использования передовых коммуникаций. Даже если зонды и роботы

используют для добычи возможности глубокого обучения, для преодоления больших и не очень инцидентов, которые могут произойти в процессе добычи ресурсов, людям все еще придется решить проблемы, с которыми искусственный интеллект не может справиться. Ожидание в течение многих часов, только чтобы обнаружить наличие проблемы, а затем еще больше часов на определение источника проблемы, является настоящим бедствием для космической горнодобывающей промышленности. Современные методы связи требуют модернизации, хотя, как ни странно, они также используют искусственный интеллект (см. <https://www.nasa.gov/feature/goddard/2017/nasa-explores-artificial-intelligence-for-space-communications>).



ЗАПОМНИ

Когнитивное радио (см. <http://ieeexplore.ieee.org/document/5783948/>) полагается на искусственный интеллект, чтобы автоматически принимать решения о необходимости улучшить эффективность радиосвязи различными способами. Человек-оператор не должен заботиться о том, как сигнал из одного места достигает другого; это просто делается самым эффективным способом. Во многих случаях для решения задачи когнитивное радио полагается на неиспользуемый или недогруженный спектр частот, но оно может полагаться и на другие методы. Другими словами, текущие методы контроля зондов, такие как представлены по адресу https://en.wikipedia.org/wiki/List_of_active_Solar_System_probes, не будут работать в будущем, когда будет необходимо общаться больше и быстрее.

Исследование новых мест

Космос обширен. Люди вряд ли когда-либо исследуют его весь. Любой, кто говорит вам, что все границы закончились, очевидно, не искал в небе. Даже авторы научно-фантастических романов уверены, что Вселенная еще долго будет содержать места для исследования людей. Конечно, если вам нравится теория множественности вселенных (<https://www.space.com/18811-multiple-universes-5-theories.html>), количество мест для исследования становится бесконечным. Проблема даже не в том, чтобы найти место, куда пойти, а в том, куда пойти сначала. Следующие разделы помогут понять роль искусственного интеллекта в путешествии людей с планеты Земля к другим планетам, а затем к звездам.

Сначала зонды

Люди уже начали запускать исследовательские зонды. На самом деле использование зондов началось куда раньше, чем думают многие. Еще в 1916 году д-р Роберт Х. Годдард (Robert H. Goddard), пионер американского ракетостроения, подсчитал, что на Луну можно отправить ракету со взрывчаткой, взрыв которой можно увидеть с Земли. Однако сам термин *зонд* (probe) дали миру Э. Берджесс (E. Burgess) и К.А. Кросс (C.A. Cross) в статье *The Martian Probe*, написанной в 1952 году. Большинство людей считают космический зонд средством, предназначенным для изучения других мест с Земли. Первый зонд, Луна-9, совершил мягкую посадку на Луну в 1966 году.

Сегодня зонды не просто пытаются достичь неких мест. Достигнув их, они решают сложные задачи, а затем сообщают результаты по радио ученым на Земле. Например, НАСА разработало зонд Mars Curiosity для поиска микробов на Марсе. Для решения этой задачи зонд Curiosity снабдили сложной системой, способной самостоятельно выполнять множество задач. Во многих случаях ожидание решения людей — просто не выход; некоторые проблемы требуют немедленного разрешения. Зонд Curiosity создает так много информации, что имеет собственный блог, подкасты и веб-сайт (https://www.nasa.gov/mission_pages/msl/index.html). О конструкции и возможностях зонда Curiosity можно прочитать по адресу <https://www.space.com/17963-mars-curiosity.html>.

Несложно вообразить огромный объем информации от отдельных зондов, таких как Curiosity. Анализ только его данных требует больших усилий от таких организаций, как Netflix и Goldman Sachs (см. <https://www.forbes.com/sites/bernardmarr/2016/04/14/amazing-big-data-at-nasa-real-time-analytics-150-million-miles-from-earth/#3b3ef365cc4f>). Различие лишь в том, что данные поступают с Марса, а не от местных пользователей, однако любой анализ данных требует времени, чтобы фактически получить информацию. Фактически сигнал от Земли до Марса идет целых 24 минуты. С учетом этого Curiosity и другие зонды должны думать сами за себя (<https://www.popsci.com/artificial-intelligence-curiosity-rover>), даже когда выполняют определенные виды анализа.

После того как данные возвращаются на Землю, ученые сохраняют их, а затем анализируют. Даже с помощью искусственного интеллекта процесс займет годы. Вполне очевидно, что полет к звездам потребует куда больше терпения и вычислительных мощностей, чем в настоящее время. С учетом сложности Вселенной использование зондов весьма важно, но зонды, возможно, должны стать более автономными для поиска правильных мест для исследований.

Ученые уже намечают вероятные места для колонизации людьми когда-нибудь в будущем. Колонизация станет важной по множеству причин, но рост населения Земли легко представить математически. Конечно, потенциальный перенос добычи и производства на другие планеты также является причиной. Плюс наличие другого места для жизни существенно улучшает наши шансы на случай, если в Землю случайно врежется астероид. С учетом этих соображений вот список наиболее вероятных целей колонизации (ваш список может отличаться).

- Луна
- Марс
- Европа
- Энцелад
- Церера
- Титан

Все эти потенциальные кандидаты предъявляют индивидуальные требования, справиться с которыми может помочь искусственный интеллект. Например, колонизация Луны потребует применения куполов. Кроме того, у колонистов должен быть достаточно большой источник воды, чтобы разделять ее на кислород для дыхания и водород для тепла. Таким образом, зонды предоставят некоторую информацию, но моделирование окружающей среды при колонизации потребует времени и больших вычислительных мощностей здесь на Земле, прежде чем люди смогут двинуться в некие другие места.

Роботизированные миссии

Вероятно, люди никогда не будут посещать другие планеты непосредственно, как их исследователи, вопреки научно-фантастическим книгам и фильмам. Куда больше смысла посылать на планеты роботы и предварительно выяснять, стоит ли вообще посылать туда людей, поскольку отправить роботы дешевле и проще. Люди фактически уже посылали роботы на многие планеты Солнечной системы и их спутники, но Марс кажется любимой целью по ряду причин.

- » Роботизированные миссии на Марс возможны каждые 26 месяцев.
- » Марс находится в пригодной для жизни зоне Солнечной системы, что делает его вероятной целью для колонизации.
- » Многие ученые полагают, что ранее на Марсе существовала жизнь.

Взаимоотношения людей с Марсом начались в октябре 1960, когда Советский Союз запустил межпланетные станции Марс-1960А и Марс-1960Б. К сожалению, зонды даже не вышли на орбиту Земли вследствие отказа ракет-носителей, не говоря уже о Марсе. Затем США предприняли попытки с космическим кораблем Mariner-3 в 1964 году и Mariner-4 в 1965 году. Станции Mariner-4 удалось отправить на Землю 12 фотографий Красной планеты. С этого времени люди послали на Марс бесчисленное количество зондов и роботов, которые начали раскрывать тайны Марса. (Коэффициент успеха марсианских миссий, однако, составляет меньше 50 процентов согласно <https://www.space.com/16777-curiosity-rover-many-mars-missions.html>). Помимо зондов, предназначенных для наблюдения Марса из космоса, посадку на Марс осуществляют роботы двух типов.

- » Lander. Стационарное роботизированное устройство, выполняющее относительно сложные задачи.
- » Rover. Мобильное роботизированное устройство, увеличивающее охват исследуемых площадей.

Вы можете найти список посадочных аппаратов и роверов, отправленных на Марс с 1971 года, по адресу <https://www.space.com/12404-mars-explored-landers-rovers-1971.html>. Несмотря на то, что большинство из них прибыло из Соединенных Штатов или Советского Союза, по крайней мере один аппарат — из Англии. Поскольку методы, необходимые для успешной посадки, становятся все более известными, другие страны тоже вскоре включатся в гонку к Марсу (даже если с помощью только дистанционного управления).



ЗАПОМНИ

По мере совершенствования зондов им потребуется улучшенный искусственный интеллект. Например, у Curiosity относительно сложный искусственный интеллект, позволяющий ему самостоятельно выбирать новые цели для исследования, как описано в статье <http://www.astronomy.com/news/2016/08/how-does-mars-rover-curiositys-new-ai-system-work>. Однако не думайте, что этот искусственный интеллект заменяет ученых на Земле. Свойства камней, которые находит искусственный интеллект, все еще определяют ученые. Кроме того, ученый может перенастроить искусственный интеллект и выбрать другую цель. Искусственный интеллект помогает ученым, а не заменяет их; это хороший пример того, как люди и искусственный интеллект будут сотрудничать в будущем.

Хотя все успешные роботизированные миссии на другие планеты полагались до сих пор на правительственное финансирование, добыча и другие коммерческие усилия в конечном счете потребуют коммерческих космических

экспедиций. Например, компания Google назначила премию Lunar XPRIZE (<https://lunar.xprize.org/>) за первое коммерческое предприятие на Луне, составляющую 20 миллионов долларов. Для победы коммерческое предприятие должно успешно посадить роботизированное устройство на Луну, устройство должно пройти не менее 500 метров и передать высококачественное видео на Землю. Конкурс важен, поскольку эта миссия будет осуществлена не только ради приза; он станет первым из многих других подобных мероприятий.

Добавление человеческого элемента

Люди хотят посетить другие места вне Земли. Конечно, единственное место, которое мы фактически посетили, — Луна. Первое такое посещение произошло 20 июля 1969 года в ходе миссии Аполлон-11. С тех пор люди высаживались на Луне шесть раз, вплоть до миссии Аполлон-17, выполненной 7 декабря 1972 года. У Китая, Индии и России есть планы относительно посадок на Луну. Россия планирует пилотируемый полет на Луну примерно в 2030 году. НАСА тоже планирует полет к Луне в будущем, но никаких дат не называет.

У НАСА действительно есть планы относительно Марса. Фактического посещения людьми Марса, вероятно, следует ожидать в 2030-х годах (<https://www.nasa.gov/topics/moon-to-mars/overview>). Как можно догадаться, наука о данных, искусственный интеллект, машинное обучение и глубокое обучение сыграют заметную роль в любом усилии по достижению Марса. Из-за расстояния и окружающей обстановки людям потребуется большая поддержка при полете на Марс. Кроме того, возвращение с Марса будет значительно более трудным, чем с Луны. Даже старт будет более трудным из-за наличия хотя и разреженной, но атмосферы, а также большей гравитации Марса.



ВНИМАНИЕ!

В 1968 году Артур Кларк выпустил книгу *Космическая одиссея 2001 года*. Книга, вероятно, вызвала большой отклик, поскольку по ней были сняты фильм и телесериал, не говоря уже о трех дополнительных книгах. В книге описан вымышленный компьютер с искусственным интеллектом HAL-9000, закончивший тем, что сошел с ума из-за конфликта в параметрах его задачи. Основная цель компьютера заключалась в помощи космическим путешественникам в их задании, но главная цель — не позволить им сойти с ума от одиночества.⁴ Безотносительно к имеющимся надеждам подобный HAL компьютер при любых космических полетах, вероятно, обречен на отказ. С одной стороны, любой запрограммированный для космоса искусственный интеллект, вероятно, не будет преднамеренно держать экипаж в

⁴ См. лучше https://ru.wikipedia.org/wiki/Космическая_одиссея_2001_года. — *Примеч. ред.*

неведении об их задании. Без сомнений, искусственный интеллект в космических полетах будет использован, но он будет иметь более практическую и обыденную функцию, чем HAL-9000.

Создание строений в космосе

В некий момент только посещения космоса станут недостаточно. Реальность космического полета в том, что все находятся в относительной тесноте, поэтому на пути к цели нужны промежуточные станции. Но даже с промежуточными станциями космический полет потребует серьезных усилий. Однако промежуточные станции важны уже сегодня. Предположим, что люди начнут фактическую добычу на Луне. На околоземной орбите потребуется склад из-за огромной стоимости доставки добываемого оборудования и других ресурсов с поверхности Земли. Конечно, должна будет осуществляться и обратная доставка добытых ресурсов или готовых изделий из космоса на Землю. Люди также захотят провести отпуск в космосе. Ученые уже используют различные космические строения для проведения своих исследований. В следующих разделах обсуждается использование различных строений, способных помочь людям на их пути от Земли к планетам и звездам.

Ваш первый отпуск в космосе

Сейчас различные компании обещают космические каникулы в ближайшем будущем. Компания Orbital Technologies сделала одно из первых таких предложений в 2011 году с первоначальной ожидаемой датой — 2016 год (см. <http://www.smh.com.au/technology/sci-tech/space-vacation-orbiting-hotel-ready-for-guests-by-2016-20110818-1j0w6.html>). Идея была в использовании российской ракеты “Союз” и проживании с шестью людьми на протяжении пяти дней. Хотя космические каникулы еще невозможны, видео по адресу <https://www.youtube.com/watch?v=2PEY0VV3ii0> знакомит с технологиями, необходимыми для того, чтобы сделать такие каникулы реальными. Большинство идей на таких сайтах вполне выполнимы, по крайней мере до некоторой степени, но сегодня еще не общедоступны. То, что вы видите, является *рекламной шумихой* (обещание товара, которого еще фактически нет, но привлечь внимание к которому уже можно), но так или иначе это интересно.



COBET

У компании Blue Origin, основанной Джеффом Безосом (Jeff Bezos), есть фактически работающая ракета (<https://www.csmonitor.com/Science/2017/0329/Blue-Origin-offers-window-into-what-a-space-vacation-might-look-like-literally>). На настоящий момент ракета совершила пять беспилотных полетов. Эти полеты

выведут людей не в космос, а скорее на низкую орбиту высотой 100 километров. У таких компаний, как Blue Origin (<https://www.blueorigin.com/>) и SpaceX (<http://www.spacex.com/>), есть хороший шанс прямо сейчас сделать космические каникулы реальностью. Компания SpaceX фактически рассматривает планы проведения каникул на Марсе (<http://www.spacex.com/mars>).

Безотносительно к последующим заявлениям люди в конечном счете окажутся в космосе по разным причинам, включая каникулы. Следует также учитывать стоимость, которая будет такой же астрономической, как и ваше расстояние от Земли. Космический полет не станет дешевым в обозримом будущем. В любом случае компании сейчас работают над космическими каникулами, хотя они еще и недоступны.

Проведение научных исследований

В космосе уже проводится множество научных исследований, и во всех них в настоящее время так или иначе помогает искусственный интеллект. Все, от международной космической станции до телескопа Hubbard⁵, весьма зависит от искусственного интеллекта (<https://spacenews.com/beyond-hal-how-artificial-intelligence-is-changing-space-systems/>). Заглядывая в будущее, вы вполне можете предположить реальность космических лабораторий и краткосрочных визитов в космос для проведения экспериментов. В настоящее время для экспериментов с имитацией невесомости используют *полет по параболической траектории* (<https://www.gozerog.com/>). Фактически полет происходит на самолете, который пикирует с большой высоты. Эта тенденция, вероятно, продолжится, причем на более высоких уровнях.

Промышленное пространство в космосе

У космических полетов могут быть разные цели. Люди уже пользуются значительными преимуществами технологий, разработанных для космических полетов и используемых сейчас для гражданских целей здесь, на Земле. (Важность космоса для жизни на Земле подчеркивается во многих статьях, одна из них — <https://www.nasa.gov/press-release/spinoff-2016-highlights-space-technologies-used-in-daily-life-on-earth>.) Но даже с передачей технологий космос все еще очень дорог, и лучшая окупаемость могла бы быть достигнута при использовании достижений другими способами, такими как создание космических фабрик (<https://www.popsci.com/factories-in-space>).

⁵ Вероятно, имелся в виду телескоп Хаббл. — *Примеч. ред.*

Фактически может оказаться, что космические фабрики являются единственным средством производства определенных материалов и продукции (см. <https://www.fastcodesign.com/3066988/mit-invented-the-material-we-need-to-build-in-space>). Условия невесомости влияют на то, как материалы реагируют и объединяются, а значит, на то, что невозможно здесь на Земле, становится вполне возможным в космосе. Кроме того, некоторые процессы легко выполнимы только в космосе, например совершенно круглый шарикоподшипник (<https://www.acorn-ind.co.uk/insight/The-Science-Experiment-Which-Took-Off-Like-A-Rocket---Creating-Space-Ball-Bearings/>).

Использование космоса для хранения

Люди в конечном счете будут хранить некоторые вещи в космосе, и это имеет смысл. По мере того как космические полеты будут становиться все более и более распространенными и люди начнут развивать промышленное производство в космосе, потребуется хранить такие вещи, как топливо и добытые ресурсы. Поскольку люди не будут знать, где будут использоваться добытые материалы (космические фабрики также будут требовать материалов), хранить их лучше будет в космосе, пока в них не возникнет потребность на Земле. Так будет намного дешевле, чем при их хранении на Земле. Космическая бензоколонка могла бы появиться куда раньше, чем вы думаете, поскольку она, возможно, понадобится при наших попытках посетить Марс (<https://futurism.com/a-gas-station-in-space-could-allow-us-to-reach-other-worlds/> и <https://www.smithsonianmag.com/innovation/nasa-sending-robotic-fueling-station-space-180963663/>).

Хотя никаких планов хранения опасных материалов в космосе пока нет, в будущем люди вполне смогут хранить загрязняющие планету отходы там. Конечно, на ум приходит вполне резонный вопрос “Зачем хранить опасные отходы, если их можно сжечь, отправив на Солнце?” В таком случае, возможно, было бы логичнее усомниться в потребности продолжать производить опасные отходы вообще. Но пока люди существуют, производство опасных отходов будет продолжаться. Их хранение в космосе дало бы нам шанс найти такие средства утилизации отходов, которые превратят их в нечто полезное.



Глава 17

Новые профессии

В ЭТОЙ ГЛАВЕ...

- » Работа в космосе
- » Строительство новых городов
- » Усиление человеческих возможностей
- » Ремонт нашей планеты

Когда люди слышат новости о роботах и других автоматах, созданных в результате технологических достижений, таких как искусственный интеллект, они обычно видят в них негатив, а не позитив. Например, в статье <https://www.theverge.com/2017/11/30/16719092/automation-robots-jobs-global-800-million-forecast> утверждается, что к 2030 году автоматизация освободит от 400 до 800 миллионов рабочих мест. Затем в статье рассматривается, как именно исчезнут эти рабочие места. Хотя в ней и признается, что некоторые технологические достижения создают новые рабочие места (например, персональный компьютер создал примерно 18,5 миллиона рабочих мест), основное внимание уделяется потере всех этих рабочих мест и возможности дальнейших потерь, которые могут стать постоянными (причем, вероятнее всего, в индустриальном секторе). Проблема в том, что большинство подобных статей вполне однозначны, когда дело доходит до потери работы, но весьма туманны, в лучшем случае, при обсуждении создания новых рабочих мест. Основная задача этой главы в том, чтобы развеять обман, прямую дезинформацию и страхи с помощью некоторых хороших новостей.

Здесь рассматриваются новые интересные профессии. Однако не рассчитывайте, что ваша профессия есть в этом списке. (Некоторые примеры профессий, неподвластных искусственному интеллекту, приведены в главе 18.) Если вы не задействованы на неких чрезвычайно простых монотонных работах, искусственному интеллекту вас не заменить. Куда вероятнее обратное: искусственный интеллект поможет вам и позволит получить больше удовольствия от вашей профессии. Даже в этом случае, после чтения данной главы, вы можете просто решить чуть повысить уровень своего образования и, получив квалификацию, освоить действительно новую и удивительную профессию.



ЗАПОМНИ

Некоторые из профессий, обсуждаемых в этой главе, являются к тому же опасными. Искусственный интеллект претендует также на выполнение множества повседневных дел в офисе или даже в вашем доме. Главное — не прекращать искать новую работу, если искусственный интеллект действительно сумеет занять вашу. Дело в том, что на протяжении своей истории люди попадали в такие ситуации многократно, особенно во времена промышленных революций, и вполне могли находить себе новые занятия. Даже если вы не извлечете ничего нового из этой главы, помните, что причиной всех страхов в мире является то, что кто-то пытается напугать вас и заставить поверить в истинность чего-то.

Жизнь и работа в космосе

Средства массовой информации заморочили людям головы идеей, что мы так или иначе исследуем всю Вселенную и выиграем все космические битвы с инопланетянами, которые прилетят, чтобы захватить нашу планету. Проблема в том, что большинству людей неизвестно, как осуществить любую из этих вещей. Тем не менее уже сегодня вы можете получить связанную с космосом работу в компании SpaceX (см. <http://www.spacex.com/careers>). Список потенциальных профессий огромен (<http://www.spacex.com/careers/list>), и большинство из них подразумевает интернатуру, чтобы вы могли получить лучшее представление о них прежде, чем начать серьезную карьеру. Конечно, вы могли бы ожидать, что все они будут сугубо техническими, но, заглянув в конец списка, можно найти буквально все профессии, включая бармена (на момент написания этой книги). Факт в том, что работа в космосе потребует практически всех профессий; в конечном счете у вас есть достаточно возможностей выбрать свой путь и найти нечто более интересное.



Компании, подобные SpaceX, заинтересованы в развитии собственных образовательных средств, поэтому они активно сотрудничают с университетами (<http://www.spacex.com/university>). Космос является относительно новым предприятием для людей, поэтому все начинают примерно на том же самом уровне, на котором все изучают что-то новое. Одна из самых волнующих частей освоения новых областей человеческих усилий состоит в том, что мы никогда не делали того, что делаем теперь, поэтому есть чему поучиться. У вас может появиться возможность сделать нечто действительно важное для человечества, но только если вы действительно пожелаете взять на себя труд исследований и все риски, связанные с выполнением чего-то нового.

Сегодня возможности фактически жить и работать в космосе ограничены, но со временем шансы увеличиваются. В главе 16 обсуждаются различные будущие занятия людей в космосе, такие как горнодобывающая промышленность или проведение исследований. Да, в конечном счете мы будем строить в космосе города после посещения других планет. Марс может стать следующей Землей. Многие люди описывают Марс как потенциально пригодный для проживания (см. <http://www.planetary.org/blogs/guest-blogs/2017/20170921-mars-isru-tech.html> и <https://www.nasa.gov/feature/goddard/2017/mars-mission-sheds-light-on-habitability-of-distant-planets>), необходимо только обновить его магнитосферу (<https://phys.org/news/2017-03-nasa-magnetic-shield-mars-atmosphere.html>).

Некоторые из обсуждаемых идей о жизни в космосе сегодня не кажутся выполнимыми, но люди весьма серьезно относятся к этим идеям, и теоретически они возможны. Например, после восстановления магнитосферы Марса станет возможным его терраформирование, чтобы сделать его пригодным для жизни. (По этой теме есть много статей; в статье <https://futurism.com/nasa-were-going-to-try-and-make-oxygen-from-the-atmosphere-on-mars/> обсуждается, как мы могли бы создать кислородную атмосферу.) Одни из этих изменений могли бы быть сделаны автоматами; другие потребовали бы вмешательства людей. Вообразите, на что могла бы походить часть группы терраформирования. Тем не менее, чтобы заставить все работать, людям придется в большой степени полагаться на искусственный интеллект. Люди и искусственные интеллекты будут сотрудничать, чтобы изменить такие места, как Марс, в соответствии с потребностями человека. Но что важнее всего, эти усилия потребуют огромного количества людей и здесь на Земле, и на Луне, и в космосе, и на Марсе. Координация действий также будет очень важна.

Строительство городов в опасной окружающей среде

На момент написания этой книги население Земли составляло 7,6 миллиарда человек (<http://www.worldometers.info/world-population/>), и это число постоянно увеличивается. Сегодня к населению Земли добавилось 153 030 человек. В 2030 году, когда НАСА планирует сделать первую попытку высадки на Марс, на Земле будет 8,5 миллиарда человек. Короче говоря, сегодня Землю населяет множество людей, а завтра будет еще больше. В конечном счете нам придется искать другие места для проживания. Кроме всего прочего, нам понадобится больше места для производства продуктов питания. Люди также захотят оставить некоторые из диких мест нетронутыми и резервировать территории в других целях. К счастью, искусственный интеллект может помочь нам найти подходящие места для построек, найти новые способы строительства зданий и помочь поддерживать состояние окружающей среды после того, как новое место станет доступным для использования.

По мере совершенствования искусственного интеллекта возможности людей возрастут, и некоторые из опасных ранее мест станут более доступными. Теоретически мы могли бы в конечном счете начать строить на вулкане, но к тому времени, конечно, будут и более подходящие места. В следующих разделах рассматриваются лишь некоторые из наиболее интересных мест, которые люди могли бы использовать для строительства городов. Эти новые места предоставляют преимущества, которых люди никогда не имели прежде, — возможность развить знание о способности жить в будущем в куда более неподходящих местах.

Строительство городов в океане

Есть несколько способов построить город в океане. Наиболее популярны две идеи: строительство плавающих городов и городов на дне океана. Фактически плавающий город находится в перспективном проектировании прямо сейчас, он предполагается недалеко от берега острова Таити (<http://www.dailymail.co.uk/sciencetech/article-4127954/Plans-world-s-floating-city-unveiled.html>). Задач у плавающих городов много, но здесь они вполне достижимы.

- » Защита от повышения уровня морей
- » Возможность опробовать новые агротехнические методы
- » Возможность опробовать новые методы выращивания рыбы
- » Выработка новых способов управления

Проживание в плавающих по океану городах — это *систейдинг* (seasteading), разновидность *хоместендинга* (homesteading), но в океане. Первые города будут существовать в относительно защищенных областях. Градостроение в открытом океане, определенно, возможно, нефтяные платформы уже полагаются на различные виды искусственного интеллекта, чтобы обеспечивать стабильность их работы и решать другие задачи (см. <https://www.techemergence.com/artificial-intelligence-in-oil-and-gas/>), но оно дорогостоящее.

Подводные города также вполне реальны, и в настоящее время существует множество подводных исследовательских лабораторий (<http://www.bbc.com/future/story/20130930-can-we-build-underwater-cities>). Ни одна из них не находится действительно глубоко, но глубина даже в 60 футов (18,288 метра) уже довольно велика. Согласно многим источникам существуют технологии создания больших городов на куда больших глубинах, но они требуют значительно лучшего мониторинга. Вот где искусственный интеллект, вероятно, будет играть ключевую роль. Искусственный интеллект может контролировать подводный город с поверхности и обеспечивать его безопасность.



ЗАПОМНИ

Важно понимать, что города в океане не могут выглядеть, как города на земле. Например, некоторые архитекторы хотят построить подводный город около Токио, который будет выглядеть, как гигантская спираль (<https://www.businessinsider.com/underwater-city-tokyo-japan-2017-1>). Эта спираль может стать домом для 5 тысяч человек. Этот город будет находиться на 16 400 футов (5 км) ниже уровня океана и полагаться на передовые технологии для обеспечения, например, электроснабжения. Это будет полнофункциональный город с лабораториями, ресторанами и школами.

Независимо от того, как люди освоят океан, это потребует широкого применения искусственного интеллекта. Частично этот искусственный интеллект уже находится на стадии разработки (<http://news.mit.edu/2017/unlocking-marine-mysteries-artificial-intelligence-1215>), равно как и подводные роботы, разрабатываемые студентами. Вполне понятно, что роботы будут частью любого строительства подводных городов, поскольку они способны выполнять различные виды работ, которые людям практически невозможно выполнить.

Создание космических поселений

Космическое поселение (space habitat) отличается от космической станции тем, что в космическом поселении живут постоянно. Причина создания космических поселений — в необходимости предоставить людям долгосрочное жилье. Предполагается, что космическое поселение будет обеспечивать

замкнутую (closed-loop) среду обитания, в которой люди смогут жить без пополнения необходимых запасов (или почти без пополнения). Следовательно, космическое поселение нуждалось бы в возобновлении воздуха и воды, методах выращивания продуктов питания и в средствах выполнения других задач, которые не нужны на космических станциях. Хотя все космические станции требуют, чтобы искусственный интеллект контролировал и поддерживал на них условия, искусственный интеллект для космического поселения был бы существенно мощнее и сложнее.

В главе 16 обсуждаются некоторые космические среды обитания (в разделе “Ваш первый отпуск в космосе”). Конечно, короткие посещения будут способом первого знакомства людей с космосом. Космические каникулы, конечно, были бы интересны! Однако место проведения каникул на близкой к Земле орбите отличается от долгосрочной среды обитания в открытом космосе, в которой будет нуждаться НАСА, если фактически преуспеет в полете на Марс. НАСА уже уполномочило шесть компаний начать изучение требований для создания среды обитания в открытом космосе (<https://www.nasa.gov/press-release/nasa-selects-six-companies-to-develop-prototypes-concepts-for-deep-space-habitats>). С некоторыми из прототипов, созданных этими компаниями, можно ознакомиться по адресу <https://www.nasa.gov/feature/nextstep-partnerships-develop-ground-prototypes>.

Для некоторых организаций космическая среда обитания является не средством проведения исследований, а скорее средством защиты цивилизации. Если сегодня гигантский астероид врежется в Землю, погибнет большая часть человечества. Но люди на Международной космической станции могли бы выжить, по крайней мере если астероид не заденет также и ее. Однако МКС не является долгосрочной стратегией выживания для людей, и количество людей на МКС ограничено. Такие организации, как Lifeboat Foundation (<https://lifeboat.com/ex/spacehabitats>), изучают космические поселения как средство выживания человечества. Их первая попытка космического поселения, Ark I (<https://lifeboat.com/ex/arki>), разрабатывается для 1 тысячи постоянных жителей и 500 гостей. Теоретически эта технология может сработать, но она потребует большого планирования.

Космические поселения могут использоваться как *корабль поколений* (generational ship), своего рода судно для исследования межзвездного пространства с использованием технологий, доступных уже сегодня. Люди просто жили бы на этом судне, пока оно двигалось бы к звездам. У них в космосе рождались бы дети, чтобы сделать столь продолжительные рейсы выполнимыми. Идея кораблей поколений не нова, они появились в книгах и фильмах много лет назад. Однако вы можете прочитать об усилиях по созданию реального корабля поколений в статье <http://www.icarusinterstellar.org/>

building-blocks-for-a-generation-ship. Проблема с кораблем поколений в том, что это судно потребовало бы соответствующего количества людей, желающих работать в различных областях, необходимых для поддержания движения судна. В этом случае люди будут взрослеть, зная, что у них есть важная задача, куда интереснее тех, с которыми они имеют дело сегодня.



Вместо того чтобы создавать компоненты космического поселения на Земле, а затем поднимать их в космос, нынешняя стратегия — добыть необходимые материалы на астероидах и использовать космические фабрики для сборки космических поселений на месте. Главный пояс астероидов Солнечной системы по современным оценкам содержит достаточно материала, чтобы построить поселения площадью, равной площади 3 тысяч планет Земля. Этого хватит для очень многих людей в космосе.

ПОСЕЛЕНИЯ ИЛИ ТЕРРАФОРМИРОВАНИЕ

Широкое использование искусственного интеллекта начнется независимо от того, как именно мы решим жить и работать в космосе. Пути создания искусственного интеллекта будут зависеть от того, куда мы пойдем и когда. Сейчас есть мнение, что мы могли бы жить на Марсе относительно короткий период времени. Однако материалы таких сайтов, как <https://phys.org/news/2017-03-future-space-colonization-terraforming-habitats.html>, наводят на мысль, что терраформирование Марса займет очень много времени. Только нагрев планеты (после того, как мы создадим технологию, способную обновить магнитосферу Марса) займет примерно сто лет. Следовательно, в действительности у нас нет выбора между поселениями и терраформированием; первыми будут поселения, и мы, вероятно, будем их широко использовать при выполнении любых работ на Марсе. Даже в этом случае искусственный интеллект для обоих проектов будет разным, с учетом типов проблем, которые будет помогать решить искусственный интеллект, но это должно быть захватывающе.

Строительство лунных баз

Вопрос не в том, вернемся ли мы на Луну; вопрос — когда. Большинство текущих стратегий колонизации космоса зависит от лунных ресурсов различных видов, включая усилия НАСА по высадке людей в конечном счете на Марсе. Мы также не страдаем от отсутствия проектов лунных баз. Некоторые из них можно увидеть по адресу <https://interestingengineering.com/8-interesting-moon-base-proposals-every-space-enthusiast-should-see>.



ЗАПОМНИ

Люди постоянно говорят о военных базах на Луне (<http://www.todayifoundout.com/index.php/2017/01/project-horizon/>), однако согласно Договору о космосе, подписанному 60 нациями, политика в космосе недопустима (<http://www.unoosa.org/oosa/en/ourwork/spacelaw/treaties/introouterspacetreaty.html>), что в значительной степени положило конец этой идее. Сначала для исследований, добычи ресурсов и фабричного производства, вероятно, будут использованы лунные базы, а затем появятся полнофункциональные города. Даже при том, что для этих проектов, вероятнее всего, будут использоваться роботы, широкий диапазон задач все еще потребует участия людей, включая ремонт роботов и управление ими. Строительство на Луне также потребует множества новых профессий, которые, вероятно, не будут фигурировать в сценариях создания поселений или исключительно космических работ. Например, кто-то должен будет ликвидировать последствия лунотрясений (см. https://science.nasa.gov/science-news/science-at-nasa/2006/15mar_moonquakes).

Построить жилой корпус на Луне можно, используя уже существующие лунные средства. Недавние исследования лунных структур, подходящих для использования при колонизации, сделали создание лунных баз еще проще. Например, вы можете прочитать об огромной пещере, вполне подходящей для колонизации, по адресу <http://time.com/4990676/moon-cave-base-lunar-colony-exploration/>. В данном случае Япония обнаружила то, что выглядит, как лавовая труба, способная защитить колонистов от множества экологических угроз.



ВНИМАНИЕ!

Конечно, о некоторых из этих структур (весьма вероятно, естественного происхождения) циркулирует множество фантастических слухов. Некоторые источники утверждают, что эти структуры на обратной стороне Луны построены инопланетянами (<http://www.dailymail.co.uk/sciencetech/article-4308270/UFO-hunters-claim-footage-aliens-moon.html>). Изображения на <https://www.youtube.com/watch?v=3caLr89zccw>¹ существенно четче. Помните: сенсацию можно раздуть на чем угодно. Структуры существуют; мы можем использовать их, чтобы проще строить здания лунных баз; а верить или не верить подобным источникам информации, решать вам.

¹ Видео удалено пользователем, который его добавил. — *Примеч. ред.*

Повышение эффективности людей

Искусственный интеллект может сделать человека более эффективным разными способами. Во многих главах этой книги приводятся примеры, когда человек полагается на искусственный интеллект, чтобы делать нечто более эффективно. Одной из самых интересных глав, тем не менее, является глава 7, в которой упоминается, как искусственный интеллект помогает в области медицины. Все эти случаи использования искусственного интеллекта предполагают, что ответственность несет человек, использующий искусственный интеллект, чтобы лучше выполнить свою задачу. Например, хирургическая система *da Vinci* не заменяет хирурга; она просто делает хирурга более способным выполнить свою задачу, причем с большей легкостью и меньшей вероятностью ошибок. Новые профессии, связанные с этим направлением, демонстрируют профессионалам, как использовать новые инструменты, обладающие искусственным интеллектом.



ЗАПОМНИ

В будущем следует ожидать появления консультантов, задачей которых будет поиск новых способов применения искусственного интеллекта в коммерческой деятельности, чтобы помочь людям стать более эффективными. До некоторой степени эта профессия уже существует, но в некоторый момент потребность в ней существенно возрастет, когда обобщенный искусственный интеллект с перестраиваемой конфигурацией станет общедоступным. Для многих компаний ключом доходности станет зависимость от поиска правильного искусственного интеллекта, который сможет повысить эффективность работы рабочих, чтобы они могли выполнять свои задачи как можно быстрее и без ошибок. Считайте этих людей частично поставщиками программного обеспечения, частично — продавцами и инструкторами в одном лице. Вы можете видеть пример мышления такого вида в статье <https://www.information-age.com/harness-ai-improve-workplace-efficiency-123469118/>.

Когда дело доходит до эффективности человека, следует подумать об областях, в которых искусственный интеллект может быть превосходен. Например, искусственный интеллект плох для решения творческих задач, поэтому их лучше оставить человеку. Но выполнение поиска искусственному интеллекту действительно удастся исключительно хорошо; таким образом, вы могли бы научить человека использовать искусственный интеллект для выполнения задач, связанных с поиском, пока он занят чем-то творческим. Вот некоторые из способов, которыми вы могли бы использовать искусственный интеллект, чтобы стать более эффективным в будущем.

- » **Наем на работу.** В настоящее время человек, нанимающий людей в организацию, не может проверить реальность сертификатов всех кандидатов и их историю. Искусственный интеллект может проверить всех кандидатов перед интервью, чтобы у нанимающего человека была вся подробная информация на момент интервью. Кроме того, поскольку искусственный интеллект использовал бы одну и ту же методологию поиска для каждого кандидата, организация может быть уверена, что все кандидаты прошли одинаковую проверку и получили справедливую оценку. Статья <https://www.forbes.com/sites/georgenehuang/2017/09/27/why-ai-doesnt-mean-taking-the-human-out-of-human-resources/#28db1e901ea6> предоставляет дополнительные детали о данной конкретной задаче. Компания товаров народного потребления, Unilever, также использует подобную технологию, что описано в статье <https://www.businessinsider.com/unilever-artificial-intelligence-hiring-process-2017-6>.
- » **Планирование.** Сегодня бизнес постоянно находится в опасности, поскольку кто-то не подумал о необходимости планировать задачу. Фактически у людей, возможно, просто не было времени, чтобы даже подумать о необходимости определения первоочередных задач. Обычно расписание контролируют секретари и ассистенты, но в нынешней упрощенной иерархии эти помощники исчезли, а обычные сотрудники выполняют собственные плановые задачи. Таким образом, перегруженные сотрудники зачастую упускают возможность помочь делу, поскольку слишком заняты контролем расписания. Сотрудничество с искусственным интеллектом освободит человека от фактического планирования. Вместо этого человек сможет оглянуться вокруг и увидеть, что именно нужно планировать. Это вопрос фокусировки внимания: человек, сосредоточенный на том, в чем он хорош, дает бизнесу куда больше. Искусственный интеллект позволит человеку сосредоточиться на том, в чем его возможности непревзойденны.
- » **Поиск скрытой информации.** Сегодня чаще, чем когда-либо, компании проигрывают в конкуренции из-за скрытой информации. В корне проблемы лежат информационная перегрузка, постоянное развитие науки, технологий, сложность бизнеса и общества. Возможно, существует новый способ упаковки товаров, который значительно снижает себестоимость, или можно изменить структуру бизнеса в результате внутренней политики. Знание того, что доступно и что было всегда, является единственным средством, которое может действительно помочь успеху компании, но эту задачу выполнить непросто. Если человек будет уделять все время приобретению знаний обо всем, чего требует конкретная работа, у него

не останется времени на фактическое выполнение этой работы. Однако искусственный интеллект исключительно хорош при поиске. Союз машинного обучения и человека позволит научить искусственный интеллект искать именно нужные проблемы и требования, чтобы сохранить бизнес на плаву, не тратя впустую ценное время на самостоятельные исследования.

- » **Адаптивная помощь.** Любой пользователь приложений сегодня должен признать, что необходимость помнить, как время от времени выполнять определенную задачу, является невероятно раздражающей, особенно если для повторного выполнения задачи приходится обратиться к справке приложения. Вы уже можете видеть, как искусственный интеллект оказывает адаптивную помощь, когда дело доходит до ввода определенных видов информации в формы. Однако искусственный интеллект может пойти значительно дальше. При использовании методов машинного обучения для поиска шаблона искусственный интеллект может в конечном счете предоставить адаптивную помощь, позволяющую пользователям выполнять те действия приложения, которые трудно запомнить. Поскольку каждый пользователь индивидуален, аппаратное решение без адаптивной помощи никогда не будет работать. Машинное обучение позволяет людям настроить систему помощи в соответствии с каждым конкретным пользователем.
- » **Адаптивное обучение.** Сегодня вы можете сдать адаптивный экзамен, показывающий пробелы в ваших знаниях. Адаптивный экзамен либо обнаруживает то, что вы действительно знаете, либо задает достаточно много дополнительных вопросов, чтобы выяснить, чему стоит поучиться дополнительно. В конечном счете приложения будут в состоянии учитывать, как вы их используете, а затем обеспечивать автоматизированное обучение, чтобы улучшить ваши навыки. Например, приложение может обнаружить, что, выполняя задачу, вы можете сделать на пять щелчков меньше, и оно может подсказать вам, как выполнить задачу, используя этот подход. Постоянно обучая людей использованию наиболее эффективных подходов при взаимодействии с компьютерами или выполнении других задач, человек становится куда эффективнее, но необходимость в человеке в этой конкретной роли остается.

Решение проблем в планетарном масштабе

Независимо от того, верите ли вы в глобальное потепление или в то, что загрязнение среды — это проблема, или вы обеспокоены перенаселенностью, факт остается фактом: планета Земля у нас только одна и на ней есть проблемы.

Погода, определенно, становится весьма странной; большие области становятся бесплодными из-за загрязнения; а некоторые области мира на самом деле слишком перенаселены. Неконтролируемые шторма и лесные пожары могут не занимать ваши мысли, но результат всегда один: сокращение областей обитания людей. Попытки поместить слишком много людей в относительно небольшое пространство обычно приводят к болезням, росту преступности и другим проблемам. Эти проблемы не являются политическими или связанными с определенными верованиями. Проблемы реальны, и искусственный интеллект может помочь решить их, предоставив компетентным сотрудникам правильные шаблоны. В следующих разделах обсуждаются планетарные проблемы с точки зрения использования искусственного интеллекта, чтобы выявить, понять причины и, возможно, устранить их. Мы не заявляем и не подразумеваем никаких политических или иных идей.

Как устроен мир

Сегодня сенсоры контролируют каждый глобальный аспект. Фактически существует так много информации, что просто удивительно, что некто может собрать всю ее в одном месте, не говоря уже о том, чтобы сделать с ней что-нибудь. Кроме того, из-за взаимодействий между различными областями Земли вы не можете в действительности знать, у каких фактов есть причинно-следственное влияние на некоторые другие части окружающей среды. Например, довольно трудно определить точно, сколько ветров влияет на потепление моря, что, в свою очередь, влияет на течения, которые потенциально могут вызывать штормы. Если бы люди понимали все эти взаимодействия, прогноз погоды был бы куда точнее. К сожалению, сейчас прогноз погоды выглядит правильно, если скосить глаза вправо и держать рот определенным образом. Тот факт, что мы принимаем этот уровень эффективности от прогнозирующих погоду людей, свидетельствует о том, что мы осознаем сложность задачи.

За последние годы прогноз погоды стал намного точнее. Частично причиной этого являются все эти сенсоры. Служба погоды также создала лучшие погодные модели и накопила намного больше данных, используемых при составлении прогнозов. Но основной причиной того, что прогнозы погоды стали куда точнее, является использование искусственного интеллекта для обработки огромных массивов данных и поиска узнаваемых шаблонов в получаемых данных (см. <https://www.techemergence.com/ai-for-weather-forecasting/>).

Фактически погода — это один из наиболее хорошо изученных земных процессов. Рассмотрим трудности в прогнозе землетрясений. Использование машинного обучения повысило вероятность правильного предсказания землетрясений (<https://www.express.co.uk/news/science/871022/earthquake-artificial-intelligence-AI-cambridge-university>), но только время

покажет, полезна ли эта новая информация. Когда-то люди думали, что на землетрясения может влиять погода, но это не так. С другой стороны, землетрясения могут влиять на погоду, изменяя условия окружающей среды. Кроме того, землетрясения и погода могут объединиться, чтобы сделать ситуацию еще хуже (<https://www.usatoday.com/story/news/nation/2015/05/02/kostigen-earthquake-weather/26649071/>).

Еще сложнее прогнозировать извержения вулканов. НАСА по крайней мере может обнаружить и получить изображения вулканических извержений с высокой точностью (<https://www.livescience.com/58423-nasa-artificial-intelligence-captures-volcano-eruption.html>). Извержения вулканов зачастую вызывают землетрясения, поэтому знание об одном помогают прогнозировать другое (<http://volcano.oregonstate.edu/how-are-volcanoes-and-earthquakes-related>). Конечно, вулканы также влияют на погоду (<http://volcano.oregonstate.edu/how-do-volcanoes-affect-atmosphere-and-climate>).

Природные события, рассмотренные в этом разделе, являются только верхушкой айсберга. Если вы пришли к выводу, что Земля настолько сложна, что ни один человек ничего не понимает, то вы правы. Именно поэтому необходимо создать и обучить искусственный интеллект для помощи людям в лучшем понимании устройства мира. Выработка знаний этого вида позволит избегать в будущем природных катастроф, а также, возможно, ограничить последствия определенных искусственных проблем.



ВНИМАНИЕ!

Независимо от того, что вы читали, в настоящее время нет способа предотвратить плохую погоду, землетрясения или извержения вулканов. Самое большее, на что могут рассчитывать люди сегодня, — это прогнозирование подобных событий и принятие мер по ограничению их последствий. Но даже возможность снижения влияния природных явлений является огромным шагом вперед. До появления искусственного интеллекта люди находились во власти любых происходящих событий, поскольку прогноз был невозможен, пока не становилось слишком поздно, чтобы предпринимать предварительные меры по ограничению последствий стихийного бедствия.

Аналогично, хотя предотвращение всех искусственных бедствий может казаться возможным, это зачастую не так. Никакое планирование не предотвратит происшествий. Однако большинство вызванных человеком событий контролируемо и потенциально предотвратимо при правильном понимании их сути, что может быть обеспечено в соответствии с шаблоном, который может предоставить искусственный интеллект.

Выявление потенциальных источников проблем

Сегодня, когда все глаза устремлены в небо, многие полагают, что спутники могут предоставить абсолютный источник данных для прогнозирования любых проблем на Земле. Однако у этой точки зрения есть много недостатков.

- » Земля настолько огромна, что обнаружение определенного события означает обработку миллионов изображений каждую секунду каждый день.
- » Изображения должны быть в правильном разрешении, чтобы фактически обнаружить событие.
- » Очень важно использовать подходящий светофильтр, поскольку некоторые события становятся видимыми только при правильном свете.
- » Погодные условия могут помешать получению изображений определенных типов.

Однако при всем этом ученые и другие заинтересованные лица используют искусственный интеллект для обработки получаемых каждый день снимков в поисках потенциальных проблем (<https://www.cnet.com/news/descartes-labs-satellite-imagery-artificial-intelligence-geovisual-search/>). Искусственный интеллект может подсказать возможные проблемные области и выполнить анализ, только если изображения представлены в правильном виде. Тем не менее только человек может определить, реальна ли проблема и как с ней справиться. Например, сильный шторм в середине Тихого океана далеко от транспортных маршрутов или суши, вероятно, не будет считаться первоочередной проблемой. А тот же самый шторм над сушей — уже причина для беспокойства. Конечно, когда дело доходит до штормов, их лучше обнаруживать до того, как это станет проблемой.



СОВЕТ

Помимо просмотра изображений в поисках потенциальных проблем, искусственный интеллект может увеличивать изображения. В статье <https://www.wired.com/story/how-ai-could-really-enhance-images-from-space/> говорится о том, как искусственный интеллект может увеличить разрешение полученных из космоса изображений, а также удобство и простоту их использования. Улучшая изображения, искусственный интеллект может упростить определение некоторых видов событий на основании их шаблонов. Конечно, если искусственный интеллект не встречал конкретный шаблон прежде, он не сможет сделать прогноз. Люди всегда должны будут проверять искусственный интеллект, чтобы удостовериться в действительности заподозренного им события.

Поиск возможных решений

Решение планетарных проблем зависит от самих проблем. Например, предотвращение таких событий, как шторм, землетрясение или извержение вулкана, даже не рассматривается. Наибольший успех, на который люди могут рассчитывать сегодня, — это выяснить область эвакуации и предоставить людям безопасное место для нахождения. Но узнав о таком событии как можно раньше, люди могут предпринять упреждающие действия, а не реагировать на событие уже после того, как наступил полный хаос.

Другие события не требуют обязательной эвакуации. Например, при нынешних технологиях и некоторой удаче люди вполне могут справиться с лесным пожаром. Фактически некоторые профессиональные пожарные уже используют искусственный интеллект, чтобы прогнозировать лесные пожары прежде, чем они произойдут (<https://www.ctvnews.ca/sci-tech/artificial-intelligence-can-better-predict-forest-fires-says-alberta-researcher-1.3542249>). Использование искусственного интеллекта для обнаружения проблемы с последующей выработкой ее решения на основании исторических данных вполне выполнимо, поскольку люди записали огромный объем информации о таких событиях в прошлом.

Использовать исторические данные при решении планетарных проблем очень важно. Наличие только одного возможного решения обычно является плохой идеей. Наилучшие планы решения проблем включают несколько решений, и искусственный интеллект может помочь ранжировать потенциальные решения на основании исторических результатов. Конечно, человек может увидеть в возможных решениях нечто такое, что сделает один выбор предпочтительнее других. Например, некое решение может быть неприемлемым из-за недоступности ресурсов или отсутствия у задействованных людей достаточной квалификации.

Контроль результатов решений

Отслеживание результатов конкретного решения означает запись данных в режиме реального времени, как можно более быстрый их анализ и последующее отображение результатов способом, понятным людям. Искусственный интеллект может собрать данные, проанализировать их и обеспечить представление результатов намного быстрее, чем любой человек. Люди все еще устанавливают критерии для выполнения всех этих задач и принимают окончательное решение, а искусственный интеллект действует просто как инструмент, позволяющий человеку действовать быстрее.



СОВЕТ

В будущем некоторые люди могли бы специализироваться на взаимодействии с искусственными интеллектами, чтобы улучшить их совместную работу с данными. Получение правильных результатов зачастую означает знание вопроса, который нужно задать, и способа, которым следует его задать. Сегодня люди нередко наблюдают негативные последствия применения искусственного интеллекта, поскольку они недостаточно хорошо знают, как работает искусственный интеллект, чтобы задавать ему корректные вопросы.

Люди, предполагающие, что искусственный интеллект думает подобно человеку, обречены на неудачу в получении хороших результатов от его использования. Конечно, это именно то, что наше общество опробует сегодня. Рекламные ролики с Siri и Alexa демонстрируют искусственный интеллект, кажущийся человеческим, но это, конечно, не так. В чрезвычайной ситуации, даже если искусственный интеллект доступен людям, пытающимся с ней справиться, людям понадобится знать, как задать соответствующие вопросы и как добиться получения необходимых результатов. Вы можете не увидеть результат решения, если не знаете, чего ожидать от искусственного интеллекта.

Повторная попытка

Земля — сложное место. Одни факторы взаимодействуют с другими факторами такими способами, которые никто не мог бы даже предвидеть. Следовательно, выработанное решение может фактически не решить проблему. Если вы читаете новости достаточно часто, то вы знаете, что многие решения не решают ничего вообще. Метод проб и ошибок позволяет людям понять, что работает, а что — нет. Однако при использовании искусственного интеллекта, способного распознавать шаблоны неудачи (решения, которые не сработали по неким причинам), вы можете ограничить количество подлежащих опробованию решений при поиске того, которое сработает. Кроме того, искусственный интеллект может искать сценарии, в которых подобные решения сработали в прошлом, что иногда экономит время и силы на попытках поиска новых вариантов решений. Искусственный интеллект — не волшебная палочка, которой можно взмахнуть и получить работоспособное решение в любой момент. Причина, по которой люди всегда будут оставаться в деле, — только люди могут видеть и оценивать результаты.



ЗАПОМНИ

Сегодня искусственный интеллект всегда программируется на победу. Во врезке “Понятие ориентации обучения” главы 13 обсуждается ситуация, когда искусственному интеллекту следует уяснить концепцию безнадежного сценария. Однако такого искусственного

интеллекта в настоящее время не существует и не может существовать. Люди действительно понимают возможность безнадёжного сценария, а поэтому нередко могут создавать не столь оптимальные решения, которые работают достаточно хорошо. При оценке, почему решение не работает, немаловажно учитывать безнадёжный сценарий, поскольку искусственный интеллект никогда его вам не представит.

Используемый при создании решений искусственный интеллект в конечном счете исчерпает идеи и в некоторый момент станет в основном бесполезным. Поэтому он не является творческим. Шаблоны, с которыми работает искусственный интеллект, уже существуют, но они не могут решить текущую задачу, а значит, необходимы новые шаблоны. Люди имеют большой опыт по части создания новых шаблонов применительно к проблемам. Следовательно, повторная попытка становится важной как средство создания новых шаблонов, к которым искусственный интеллект может затем обратиться и которые может использовать, чтобы помочь человеку вспомнить нечто, что работало в прошлом. Короче говоря, люди — основная часть цикла решения проблемы.

6

**Великолепные
десятки**

В ЭТОЙ ЧАСТИ...

- » Профессии, недоступные искусственному интеллекту**
- » Как искусственный интеллект помогает обществу**
- » Почему искусственный интеллект терпит неудачу в некоторых ситуациях**



Глава 18

Десять профессий, недоступных искусственному интеллекту

В ЭТОЙ ГЛАВЕ...

- » Взаимодействие с людьми
- » Творчество
- » Использование интуиции

В этой книге достаточно много внимания уделено различиям между искусственным интеллектом и людьми, а также уверениям в том, что людям абсолютно не о чем волноваться. Да, некоторые профессии исчезнут, но, как описано в главе 17, при использовании искусственного интеллекта фактически будет создано множество новых профессий, большинство из которых будут намного интереснее работы на сборочной линии. Новые профессии будут полагаться на такие типы интеллекта, которыми искусственный интеллект просто не может овладеть (как описано в главе 1). Фактически неспособность искусственного интеллекта овладеть очень многими областями процесса мышления человека сохранит большинству людей их нынешние профессии, что и является темой данной главы.



ЗАПОМНИ

Вы можете найти, что ваша профессия недоступна для искусственного интеллекта, если она относится к определенным категориям: общение с людьми, творчество или использование интуиции. Но в этой главе затрагивается только верхушка айсберга. Бойтесь спекуляций определенных личностей (см. <https://www.theinquirer.net/inquirer/news/3013919/elon-musk-spews-more-ai-fear-mongering-is-desperate-bid-for-more-media-attention>), как бы заботящихся о людях, профессии которых вскоре исчезнут. Бойтесь спекуляций, препятствующих людям использовать полный потенциал искусственного интеллекта, чтобы сделать жизнь проще (см. <https://www.cnbc.com/2017/09/21/head-of-google-a-i-slams-fear-mongering-about-the-future-of-a-i.html>). Общее послание данной главы — не бойтесь. Искусственный интеллект — это инструмент, который, как любой другой инструмент, предназначен для того, чтобы сделать жизнь проще и лучше.

Общение с людьми

Роботы уже выполняют небольшой объем работ по общению с людьми, а в будущем, вероятно, будут выполнять еще больше. Но если взглянуть на подобные приложения внимательнее, то окажется, что, по существу, они делают вещи, смехотворно примитивные: служат указателем в магазине, подсказывая людям, куда пойти; используются в качестве будильника, напоминающего пожилым людям о приеме лекарства; и т.д. Большинство случаев общения с людьми не так просты. В следующих разделах рассматриваются некоторые из наиболее интерактивных действий, требующих таких видов человеческого общения, на которые искусственный интеллект вообще не способен.

Обучение детей

Уделим время начальной школе и посмотрим, как учителя пасут детей. Вы будете поражены. Так или иначе учителям удастся перемещать всех детей из пункта А в пункт В с минимумом суеты, вероятно, одной только силой воли. Но даже в этом случае один ребенок будет нуждаться в одном уровне внимания, а другой ребенок — в другом. Когда что-то идет не так, учителю приходится справляться с несколькими проблемами одновременно. Любая из этих ситуаций сокрушит современный искусственный интеллект, поскольку он полагается на сотрудничество с человеком при общении. Задумайтесь на минуту о реакции Alexa или Siri на упрямство ребенка (или попытайтесь смоделировать

подобную реакцию на собственном устройстве). Это просто не сработает. Искусственный интеллект может только помочь учителю в следующих случаях.

- » Сортировка бумаг
- » Использование адаптивного образовательного программного обеспечения
- » Улучшение учебных курсов исходя из шаблонов учеников
- » Снабжение учеников обучающими программами
- » Обучение учеников поиску информации
- » Создание более безопасной обстановки для эмпирического обучения
- » Помощь ученикам в принятии решений о выборе курсов и занятий после школы на основании их навыков
- » Помощь ученикам с домашними заданиями

Медицинский уход

Робот может поднять пациента, заменив медсестру. Однако искусственный интеллект не может принять решение о том, когда, где и как поднять пациента, поскольку не может правильно оценить все необходимые невербальные данные о пациенте или понять его психологию, например склонность врать (см. раздел “Пять недостоверностей данных” главы 2). Искусственный интеллект может задавать вопросы пациенту, но, вероятно, не тем способом, который лучше всего подходит для получения полезных ответов. Робот может помыть, но маловероятно, что он сделает это способом, сохраняющим достоинство пациента, и позволит ему почувствовать заботу. Короче говоря, робот — хороший молоток: прекрасный для выполнения сложных и грубых задач, но не особенно нежный и заботливый.



ЗАПОМНИ

Использование искусственного интеллекта в медицине, несомненно, расширится, но области его применения будут чрезвычайно специфическими и ограниченными. Глава 7 дает хорошее представление о том, где искусственный интеллект может помочь в области медицины. Немногие из этих действий имеют какое-либо отношение к общению с человеком. Они скорее будут действовать по линии усиления человеческих возможностей и сбора медицинских данных.

Удовлетворение личных нужд

Можно подумать, что ваш искусственный интеллект — это совершенный компаньон. В конце концов, он никогда не возражает, всегда внимателен и

никогда не оставит вас ради кого-то еще. Вы можете поведать ему свои самые сокровенные мысли, и он не будет смеяться. Фактически такой искусственный интеллект, как Alexa или Siri, может быть великолепным компаньоном, как в фильме *Она* (<https://www.amazon.com/exec/obidos/ASIN/B00H9HZGQ0/data-cservip0f-20/>). Единственная проблема в том, что искусственный интеллект не является компаньоном вообще. Но что он действительно делает, так это обеспечивает приложения голосом. Наделение искусственного интеллекта человеческими качествами не делает их реальностью.

Проблема с удовлетворением личных нужд искусственным интеллектом заключается в том, что он понятия не имеет о личных нуждах. Искусственный интеллект может найти радиостанцию, статью в новостях, сделать покупки, запланировать встречу, напомнить о времени приема лекарств и даже включить или выключить свет в вашем доме. Однако он не может подсказать вам, что некая идея плоха, и избавить вас от последующих проблем. Чтобы получить полезный совет в ситуациях, не предусмотренных правилами, нужно поговорить с человеком, обладающим практическим опытом выхода из подобной ситуации в реальной жизни, поэтому вам действительно нужен человек. Именно поэтому такие люди, как адвокаты, врачи, медсестры и даже собеседники в кафе, необходимы. Одни из этих людей оказывают платные услуги, а другие готовы просто выслушать вас, когда вы нуждаетесь в помощи. При удовлетворении личных нужд, которые действительно являются личными, человеческое общение всегда обязательно.

Решение проблем, связанных с развитием

Людям со специальными потребностями требуется человеческая забота. Нередко специальная потребность оказывается специальным даром, но только тогда, когда обладатель признает ее как таковую. Некто со специальной потребностью может быть полностью функциональным во всем, кроме одного — творческого потенциала и воображения, чтобы находить способы преодоления затруднений. Поиск способа использования специальной потребности в мире, который не принимает специальные потребности, весьма труден. Например, большинство людей не рассматривает дальтонизм (который фактически является смещением восприятия цветов) как преимущество при создании произведений искусства. Но кто-то пришел, и превратил это в преимущество (<https://www.artsy.net/article/artsy-editorial-the-advantages-of-being-a-colorblind-artist>).

Искусственный интеллект мог бы помочь людям со специальными потребностями определенными способами. Например, робот может помочь в выполнении реабилитационной или физической терапии, чтобы сделать человека более подвижным. Абсолютное терпение робота гарантирует человеку

получение той же беспристрастной помощи каждый день. Но нужен человек, чтобы понять, что реабилитационная или физическая терапия не работает и требует изменений.



ВНИМАНИЕ!

Помощь с проблемами развития — это одна из областей, в которых искусственный интеллект, независимо от того, насколько хорошо он запрограммирован и обучен, может оказаться фактически вредным. Человек может заметить, когда кто-то переусердствовал, даже если кажется, что он достиг успехов в решении различных задач. Для достижения успеха поможет комплекс невербальных признаков, но это также вопрос опыта и интуиции — качеств, которыми искусственный интеллект не обладает в изобилии, поскольку потребуется множество ситуаций, чтобы искусственный интеллект *экстраполировал* (расширил свое знание до неизвестной ситуации), а не *интерполировал* (использовал знания между двумя известными пунктами). Коротко говоря, люди не только будут контролировать человека, которому помогают они и искусственный интеллект; им придется контролировать еще и искусственный интеллект, чтобы он работал так, как ожидалось.

Создание нового

Как упоминалось в табл. 1.1, роботы не способны к творчеству. Этот немаловажный фактор следует учитывать при выработке нового образа мыслей. Хорошее приложение глубокого обучения может анализировать существующие шаблоны мышления, а искусственный интеллект может превращать эти шаблоны в новые версии вещей, которые уже существовали прежде, что создает впечатление оригинального мышления, но никакого творчества здесь нет. То, что вы видите, является результатом математического и логического анализа существующих работ, а не определение того, что могло бы быть. С учетом этого ограничения искусственного интеллекта в следующих разделах описано создание новых вещей — область, в которой люди всегда будут вне конкуренции.

Изобретения

Когда люди говорят об изобретателях, они вспоминают Томаса Эдисона, имевшего 2 332 патента на изобретения во всем мире (1093 только в Соединенных Штатах) (<http://www.businessinsider.com/thomas-edisons-inventions-2014-2>). Вы все еще можете использовать одно из его изобретений, лампочку, но большинство из них, таких как фонограф, изменили мир. Эдисон не

единственный. Есть такие люди, как Бетт Несмит Грэм (Bette Nesmith Graham) (<http://www.women-inventors.com/Bette-Nesmith-Graham.asp>), которая изобрела корректирующую жидкость Whiteout (известную также как Liquid Paper) в 1956 году. Ее изобретение находилось на столе каждой машинистки в мире как средство для исправления опечаток. Оба эти человека сделали нечто, на что не способен искусственный интеллект: создание нового шаблона мышления в форме физической сущности.



ЗАПОМНИ

Да, каждый из этих людей черпал вдохновение из других источников, но идея была действительно их собственной. Дело в том, что люди постоянно что-то изобретают. В Интернете можно найти миллионы и миллионы идей, созданных людьми, которые просто видят вещи не так, как другие. Если уж на то пошло, люди станут еще более изобретательными, поскольку у них появится на это время. Искусственный интеллект может освободить людей от обыденного, чтобы они могли делать то, на что они способны лучше всех: изобретать новые вещи.

Искусство

Стиль и видение делают Пикассо (<https://www.pablopicasso.org/>) отличным от Моне (<https://www.claudemonetgallery.org/>). Люди способны различать их потому, что мы видим шаблоны в методах этих художников: все, от выбора холста и красок до стиля представления и тем изображений. Искусственный интеллект тоже может видеть эти различия. Фактически с учетом высокой точности выполнения искусственным интеллектом анализа и более широкого набора сенсоров в его распоряжении (как правило) искусственный интеллект, вероятно, способен описать шаблоны художников куда лучше, чем человек, и подражать этим шаблонам способом, который художник никогда не использовал. Однако на этом преимущества искусственного интеллекта заканчиваются.



СОВЕТ

Искусственный интеллект будет придерживаться того, что он знает, а люди экспериментируют. Вы можете найти 59 примеров человеческих экспериментов с различными материалами по адресу <https://www.pinterest.com/aydeeyai/art-made-with-non-traditional-materials-or-methods/>. Только человек может додуматься создавать произведения искусства из проволоочной сетки (<https://www.pinterest.com/pin/491947959277129127/>) или листьев (<https://www.pinterest.com/pin/451697037596827773/>). Если есть материал, найдется и кто-то, кто создаст из него произведение искусства, на что искусственный интеллект не способен.

Воображение нереального

Люди постоянно расширяют рамки реального, делая возможным нереальное. Когда-то никто не думал, что люди будут летать, но все придумывали воздушные машины, которые тяжелее воздуха. Фактически эксперименты подтверждали теорию о том, что даже попытка летать была глупой. Затем появились братья Райт (<http://www.history.com/topics/inventions/wright-brothers>). Их полет на Kitty Hawk изменил мир¹. Однако важно понимать, что братья Райт просто сделали нереальные мысли многих людей (включая собственные) реальными. Искусственный интеллект никогда не имел бы нереальных мыслей, не говоря уже об их воплощении в действительность. Только люди способны на это.

Интуитивное принятие решений

Интуиция — прямое восприятие факта независимого от любого процесса рассуждения. Поскольку этот факт не является результатом логики, его невероятно трудно анализировать. Имеется большой опыт по части интуиции, и у людей с развитой интуицией обычно есть существенное преимущество перед теми, у кого ее нет. У искусственного интеллекта, в основе которого заложены логические и математические принципы, интуиция полностью отсутствует. Следовательно, искусственному интеллекту придется продраться через все возможные логические решения и в конечном счете прийти к заключению, что никакого решения проблемы не существует, даже если человек находит его с относительной легкостью. Человеческая интуиция и способность проникновения в суть зачастую играют огромную роль в некоторых профессиях, как описано в следующих разделах.

Расследование преступлений

Любители детективных сериалов знают, что следователь обычно находит один небольшой факт, который раскрывает все дело. В реальности расследование преступления осуществляется иначе. Для выполнения своей задачи детективы полагаются на точное знание, но иногда преступники упрощают их работу. Процедуры и правила, многочасовые копания в фактах и поиск всех доказательств играют важную роль в расследовании преступления. Но иногда у человека наступает это нелогичное просветление, которое внезапно складывает вместе все, казалось бы, никак не связанные части.

¹ Чисто технически Китти Хоук — это городок, недалеко от которого проходили испытания, а аппарат братьев Райт назывался “Флайер-1”. См. https://ru.wikipedia.org/wiki/Wright_Flyer. — *Примеч. ред.*

Работа детектива подразумевает решение широкого диапазона проблем. Фактически некоторые из этих проблем даже не подразумевают незаконных действий. Например, детектив может просто искать того, кто кажется пропавшим без вести. Возможно, у человека есть серьезное основание для того, чтобы не желать быть найденным. Дело в том, что большинство подобных поисков подразумевает взгляд на факты такими способами, до которых искусственный интеллект никогда не додумается, поскольку это требует озарения — расширения интеллекта. На ум приходит фраза *нестандартное мышление*.

Контроль ситуации в режиме реального времени

Искусственный интеллект будет контролировать ситуации, используя имеющиеся данные в качестве основания для будущих решений. Другими словами, искусственный интеллект использует для прогнозов шаблоны. В большинстве ситуаций использование шаблонов работает прекрасно, а значит, искусственный интеллект вполне может фактически прогнозировать с высокой степенью точности то, что произойдет при определенном сценарии развития событий. Но иногда возникают ситуации, в которых соответствия шаблону нет и данные, казалось бы, не обеспечивают заключения. Возможно, для ситуации в настоящее время не хватает сопутствующих данных, такое иногда случается. В таких случаях человеческая интуиция — единственный выход в аварийных условиях. Надежда только на искусственный интеллект в чрезвычайной ситуации — это явно плохая идея. Хотя искусственный интеллект действительно предоставит проверенное решение, человек может думать нестандартно и придумать лучшую альтернативную идею.

Распознавание фактов и вымысла

У искусственного интеллекта никогда не будет интуиции. Интуиция находится в противоречии со всеми правилами, по которым в настоящее время создается искусственный интеллект. Поэтому некоторые люди решили создать *искусственную интуицию* (Artificial Intuition — AN) (см. <http://www.artificial-intuition.com/>). Читая эти материалы, быстро приходишь к мнению, что здесь имеет место некая магия (т.е. изобретатели пытаются выдать желаемое за действительное), поскольку теория просто не соответствует предложенной реализации.



ЗАПОМНИ

Некоторые немаловажные проблемы искусственной интуиции связаны, в первую очередь, с тем, что все программы, даже с поддержкой искусственного интеллекта, выполняются на процессорах, единственной возможностью которых является выполнение только

самых простых математических и логических функций. Данный искусственный интеллект работает на уровне доступных в настоящее время аппаратных средств, в чем нет ничего удивительного.

Вторая проблема заключается в том, что искусственный интеллект и все компьютерные программы, по существу, для решения задач полагаются на математику. Искусственный интеллект ничего не понимает. В разделе “Аргумент китайской комнаты” главы 5 затрагивается только одна из многих проблем на пути искусственного интеллекта к пониманию. Дело в том, что интуиция нелогична, а значит, люди даже не понимают ее основ. Без понимания люди не могут создать систему, которая могла бы подражать интуиции любым осмысленным способом.



Глава 19

Десять достижений искусственного интеллекта, наиболее полезных для общества

В ЭТОЙ ГЛАВЕ...

- » Сотрудничество с людьми
- » Решение производственных задач
- » Разработка новых технологий
- » Решение задач в космосе

Эта книга помогает понять историю искусственного интеллекта, его настоящее и возможное будущее. Однако технология полезна только тогда, когда она вносит существенный вклад в жизнь общества. Кроме того, вклад должен сопровождаться мощным финансовым стимулом, иначе инвесторы не будут ему способствовать. Хотя правительство и может поддержать технологию, которую считает полезной для вооруженных сил или других целей,

эта поддержка будет краткосрочной; долгосрочное благополучие технологии полагается на поддержку инвесторов. Данная глава сосредоточена на тех компонентах искусственного интеллекта, которые полезны уже сегодня, а значит, приносят существенную пользу обществу прямо сейчас.



ЗАПОМНИ

Некоторые люди говорят, что нынешние обещания сверхпреимуществ искусственного интеллекта могут вызвать следующую зиму искусственного интеллекта (<https://codeahoy.com/2017/07/27/ai-winter-is-coming/>). Кроме того, спекуляции одних людей вселяют страхи в других и заставляют заново осмыслить значение искусственного интеллекта (<https://www.theinquirer.net/inquirer/news/3013919/elon-musk-spews-more-ai-fear-mongering-is-desperate-bid-for-more-media-attention>). Им противостоят люди, полагающие, что вероятность зимы искусственного интеллекта невелика (<https://www.technologyreview.com/s/603062/ai-winter-isnt-coming/>) и что опасения неуместны (<https://www.cnn.com/2017/09/21/head-of-google-a-i-slams-fear-mongering-about-the-future-of-a-i.html>). Обсуждение важно при оценке любой технологии, но инвесторам нужны не слова; их интересуют результаты. Эта глава — о результатах, о том, что искусственный интеллект уже достаточно глубоко интегрирован в общество, чтобы следующая зима искусственного интеллекта действительно стала вероятной. Конечно, было бы плюсом отсутствие обмана, чтобы люди могли понять, что искусственный интеллект может сделать для них в настоящий момент.

Учет взаимодействий, специфических для человека

Продажей товаров управляют люди. Кроме того, люди решают, что скажет большинство, что вызовет шум, что, в свою очередь, повысит продажи. Хотя о технологиях, обсуждаемых в следующих разделах, вы, вероятно, не слышали по радио, их уровень влияния на людей потрясающий. Активный протез ноги (первый случай) действительно помогает людям ходить и использовать протезы с такой же легкостью, как и собственную ногу. Несмотря на то что нуждающаяся в этом группе людей относительно невелика, результат его применения становится известным весьма широко. Второй и третий случаи могут влиять на миллионы, а может, и миллиарды людей. Это товары повседневного пользования, но зачастую именно повседневность стимулирует

повторные продажи. Во всех трех случаях технологии не будут работать без искусственного интеллекта, а это значит, что прекращение исследований, разработок и продаж искусственного интеллекта, вероятно, повредит людям, использующим эти технологии.

Изобретение активного протеза ноги

Эндопротезы стоят больших денег. Их изготовление стоит состояния, но они являются необходимым элементом для любого потерявшего конечность, желающего иметь приличное качество жизни. Многие эндопротезы полагаются на пассивную технологию, а значит, не обеспечивают реакции и не корректируют автоматически свои функциональные возможности, чтобы приспособиться к личным нуждам. Все, что изменилось за годы, когда такие ученые, как Хью Херр (<https://www.smithsonianmag.com/innovation/future-robotic-legs-180953040/>), создали активный эндопротез, — так это способность имитировать действия реальных конечностей и автоматически приспосабливаться к использующему их человеку. Хотя основные аплодисменты сорвал Хью Херр, сегодня вы можете найти активную технологию в эндопротезах всякого рода, включая протезы колен, рук и кистей.



ЗАПОМНИ

Вы можете задаться вопросом о потенциальном значении использования активных и пассивных эндопротезов. Поставщики медицинского оборудования уже проводят исследования (некоторые из результатов содержатся в отчете https://www.rand.org/pubs/research_reports/RR2096.html). Кроме того, гарантия правильного взаимодействия с пользователем такого устройства, как основанный на микропроцессоре эндопротез, полагающийся на искусственный интеллект, уже является огромной победой. Люди, использующие технологии активного эндопротеза, не только дольше живут активной жизнью, но и экономят на прямых и косвенных медицинских затратах. Например, человек, использующий технологию активного протеза, имеет меньше шансов споткнуться. Даже при том, что начальная стоимость активного протеза выше, со временем она окупается.

Постоянный мониторинг

В главе 7 обсуждается множество используемых в медицине устройств мониторинга, а также устройств, гарантирующих прием лекарств в нужное время и в правильной дозировке. Кроме того, медицинский мониторинг может помочь пациентам быстрее получить медицинскую помощь после приступа и даже спрогнозировать следующий приступ, например сердечный. Большинство подобных устройств, особенно прогнозирующих, полагаются в работе на искусственный

интеллект некого вида. Однако остается вопрос, обеспечивают ли эти устройства финансовый стимул для людей, создающих и использующих их.

Исследования идут трудно, но результаты (<https://academic.oup.com/europace/article-abstract/19/9/1493/3605206>) показывают, что дистанционный мониторинг кардиологических пациентов обеспечивает значительную экономию (помимо помощи пациенту жить долго и счастливо). Фактически согласно Файнэншл Таймс (<https://www.ft.com/content/837f904e-9fd4-11e4-9a74-00144feab7de>) использование дистанционного мониторинга даже для здоровых людей оказывает существенное влияние на медицинские расходы (для чтения статьи нужна подписка). Экономия настолько существенна, что дистанционный мониторинг может изменить работу всей медицины.

Прием лекарств

Пациенты, забывающие принимать свои лекарства, обходятся медицинским учреждениям в огромные суммы денег. Согласно этой статье CNBC.com (<https://www.cnbc.com/2016/08/03/patients-skipping-meds-cost-290-billion-per-year-can-smart-pills-help.html>) в одних только Соединенных Штатах эта сумма составляет 290 миллиардов долларов в год. При объединении такой технологии, как Near Field Communication (NFC) (<https://www.nfcworld.com/2015/11/18/339766/nxp-launches-nfcb blister-packs-and-pill-bottles-for-medication-tracking/>), с приложениями, полагающимися на искусственный интеллект, вы можете отслеживать, как люди принимают свои лекарства и когда. Кроме того, искусственный интеллект может помочь людям запомнить, когда принимать лекарства, какие и сколько. Совместно с мониторингом даже люди, нуждающиеся в специальном контроле, могут получить правильную дозу своих лекарств (<https://clinicaltrials.gov/ct2/show/NCT02243670>).

Выработка промышленных решений

Люди совершают тысячи небольших продаж. Но когда думаешь о покупательной способности человека, она бледнеет по сравнению с тем, что может потратить одна организация. Все дело в количестве. Однако инвесторов интересуют оба вида продаж, поскольку оба делают деньги, причем довольно много. Индустриальные решения затрагивают организации. Обычно они весьма дороги, но все же промышленность использует их для повышения продуктивности, эффективности и, главное, доходности. В следующих разделах рассматривается, как искусственный интеллект затрагивает организации, которые используют полученные решения.

Искусственный интеллект и трехмерная печать

Трехмерная печать начинала как игрушечная технология, впечатлившая некоторых, но не особенно ценная по результатам. Но это было прежде, чем НАСА использовало трехмерную печать на Международной космической станции для создания инструментов (<https://www.nasa.gov/content/international-space-station-s-3-d-printer>). Большинство людей полагают, что все необходимые на МКС инструменты можно было взять с собой с Земли. К сожалению, инструменты теряются и ломаются. Кроме того, на МКС просто нет достаточного пространства для хранения абсолютно всех необходимых инструментов. Трехмерная печать может также создавать запасные части, а МКС, конечно, не может нести все возможные запасные части. При микрогравитации трехмерные принтеры работают точно так же, как и на Земле (https://www.nasa.gov/mision_pages/station/research/experiments/1115.html). Таким образом, трехмерная печать — это технология, которую ученые могут использовать тем же самым способом в обоих местах.

Тем временем промышленность использует трехмерную печать для удовлетворения самых разнообразных запросов. Добавление в этот комплекс искусственного интеллекта позволит устройству получать результаты, видеть созданное и учиться на своих ошибках (<https://www.digitaltrends.com/cool-tech/ai-build-wants-to-change-the-way-we-build-the-future/>). Это значит, что промышленность в конечном счете будет в состоянии создавать роботы, которые будут исправлять свои ошибки по крайней мере до такой степени, которая снизит количество ошибок и увеличит прибыль. Искусственный интеллект позволяет также снизить риск, связанный с трехмерной печатью за счет применения таких продуктов, как Business Case (<https://www.sculpteo.com/blog/2017/08/10/the-artificial-intelligence-for-your-3d-printing-projects-business-case/>).

Передовые роботизированные технологии

В этой книге содержится множество примеров использования роботов, от домашних до медицинских и промышленных. Здесь также упоминается о роботах в автомобилях, в космосе и под водой. Если вы полагаете, что роботы — это основная движущая сила искусственного интеллекта, вы правы. Роботы становятся надежной, доступной и известной технологией с видимым присутствием и хорошим послужным списком, а потому очень многие организации вкладывают капитал в их дальнейшее совершенствование.

Сегодня на роботы полагается куда больше традиционных компаний, чем можно подумать. Например, нефтедобывающая промышленность полностью полагается на роботы при поиске новых нефтяных месторождений, инспекции и

обслуживании трубопроводов. В некоторых случаях роботы осуществляют также ремонт в таких местах, к которым людям нелегко добраться, например в трубопроводах (<http://insights.globalspec.com/article/2772/the-growing-role-of-artificial-intelligence-in-oil-and-gas>). Согласно Oil & Gas Monitor искусственный интеллект обеспечивает интерполяцию между разведочными скважинами, снижая стоимость буровых работ и создавая модели, демонстрирующие потенциальные проблемы бурения (<http://www.oilgas-monitor.com/artificial-intelligence-upstream-oil-gas/>). Использование искусственного интеллекта позволяет инженерам снизить общий риск, а значит, добыча нефти окажет потенциально меньшее воздействие на окружающую среду из-за снижения количества разливов.



СОВЕТ

Снижение цены на нефть стало именно тем, что подтолкнуло нефтедобывающую промышленность к принятию искусственного интеллекта согласно Engineering 360 (<http://insights.globalspec.com/article/2772/the-growing-role-of-artificial-intelligence-in-oil-and-gas>). Поскольку нефтедобывающая промышленность — дело рискованное, использование ею искусственного интеллекта дает хороший повод для его проверки и демонстрации другим компаниям его преимуществ. При чтении статей о нефтедобывающей промышленности понимаешь, что она ожидала успехов в здравоохранении, финансах и обрабатывающей промышленности, прежде чем делать собственные инвестиции. Вы можете ожидать очередного роста от применения искусственного интеллекта после его успеха в других отраслях.



ЗАПОМНИ

В этой книге рассматриваются самые разные робототехнические решения, некоторые из них мобильны, некоторые — нет. В части 4 вообще рассматриваются летающие роботы (дроны) и самоходные или беспилотные автомобили. Обычно роботы способны приносить прибыль, когда выполняют конкретный вид задач, такой как чистка вашего пола (Roomba) или сборка вашего автомобиля. Аналогично дроны теперь являются источником доходов для оборонных подрядчиков и в конечном счете станут таковыми и для гражданских предприятий. Многие люди предсказывают, что беспилотный автомобиль не только будет зарабатывать деньги, но и станет чрезвычайно популярным (<https://www.forbes.com/sites/oliviergarret/2017/03/03/10-million-self-driving-cars-will-hit-the-road-by-2020-heres-how-to-profit/>).

Создание новых технологических сред

Обычно все ищут новые вещи, чтобы купить, а значит, компании должны их придумывать, чтобы продать. Искусственный интеллект помогает людям находить шаблоны во всякого рода вещах. Шаблоны зачастую показывают присутствие чего-то нового, такого как новый элемент или новый процесс для создания чего-то еще. В области разработки изделий задача искусственного интеллекта заключается в том, чтобы помочь обнаружить новый продукт (в отличие от продажи существующего продукта, находящегося в фокусе). Сокращая время, необходимое для поиска нового товара, чтобы продавать, искусственный интеллект помогает бизнесу повышать прибыль и снижать стоимость исследований, связанных с поиском новых товаров. Более подробно эти проблемы обсуждаются в следующих разделах.

Разработка новых редких ресурсов

Как можно заметить повсюду в этой книге, искусственный интеллект имеет особенно большой опыт по части поиска шаблонов, а шаблоны могут указать на всякого рода вещи, включая новые элементы (этому аспекту искусственного интеллекта посвящен раздел “Поиск новых элементов” главы 16). Новые элементы означают новые товары, а значит, и новые продажи. У организации, способной придумать новый материал, будет существенное преимущество в конкуренции. Статья <https://virulentwordofmouse.wordpress.com/2010/11/30/an-economic-perspective-on-revolutionary-us-inventions/> повествует о влиянии на экономику некоторых наиболее интересных изобретений. Большинство из них полагается на новый процесс или материал, который искусственный интеллект с легкостью может помочь найти.

Увидеть невидимое

Человек видит в довольно ограниченном спектре света, который фактически существует в природе. И даже с усилием люди пытаются мыслить в очень малых масштабах или в очень больших. Пристрастия не позволяют людям видеть неожиданное. Иногда у казалось бы случайного шаблона есть структура, но люди ее не видят. Искусственный интеллект может видеть то, чего не могут видеть люди, а затем воздействовать на это. Например, при поиске напряжений в металле искусственный интеллект может заметить усталость и воздействовать на это. Применение для радиопередачи таких элементов, как волноводы, может дать огромную экономию (<https://www.ncbi.nlm.nih.gov/pmc/articles/PMC4481976/>).

Сотрудничество с искусственным интеллектом в космосе

В главе 16 приведен обзор того, на что искусственный интеллект может быть потенциально способен в космосе. Даже при том, что планы относительно выполнения этих задач находятся на чертежной доске, большинство из них финансируется правительством, а значит, для них есть возможность необязательно быть прибыльными. В главе 16 описаны также некоторые научно-исследовательские работы, связанные с бизнесом. В данном случае бизнес фактически надеется получать прибыль, но сегодня не может этого сделать. В следующих разделах космос рассмотрен в другом отношении и указано на то, что происходит сегодня. В настоящее время искусственный интеллект позволяет компаниям зарабатывать деньги в космосе, что дает им стимул продолжать вкладывать капитал и в искусственный интеллект, связанный с космическими проектами.

Доставка товаров на космические станции

Возможно, самым успешным коммерческим применением искусственного интеллекта в космосе является пополнение запасов на МКС такими компаниями, как SpaceX и Orbital ATK (https://www.nasa.gov/mission_pages/station/structure/launch/overview.html). Конечно, эти организации делают деньги на каждом рейсе, но и НАСА получает доход. Фактически все Соединенные Штаты извлекают пользу из этого предприятия.

- » Возможность избежать использования транспортных средств других стран, а следовательно, снижение цены доставляемых на МКС материалов
- » Повышение роли таких американских организаций, как Kennedy Space Center, что означает долговременную амортизацию вложенных в них средств
- » Создание дополнительных пусковых центров для космических полетов в будущем
- » Более доступная коммерческая доставка спутников и других элементов

Компании SpaceX и Orbital ATK сотрудничают с множеством других компаний. Следовательно, хотя и может показаться, что только эти две компании могли бы извлечь выгоду из этой ситуации, на самом деле пользу получают и многие их партнеры. Использование искусственного интеллекта делает все это возможным, и может быть прямо сейчас. Компании зарабатывают деньги в

космосе и сегодня, не дожидаясь завтра, как вы могли бы подумать, посмотрев новости. Однако то, что доход поступает от того, что, по сути, является обычной службой доставки, не имеет никакого значения.



ЗАПОМНИ

Космическая доставка чрезвычайно нова. Многие основанные на Интернете компании были убыточными на протяжении многих лет, прежде чем стали прибыльными. Однако SpaceX по крайней мере ищет возможности заработать деньги после нескольких первых провалов (<https://www.fool.com/investing/2017/02/05/how-profitable-is-spacex-really.aspx>). Космически-ориентированным компаниям потребуется время, чтобы достичь того же финансового влияния, которым обладают сегодня сугубо земные компании.

Добыча внепланетных ресурсов

Такие компании, как Planetary Resources (<https://www.planetaryresources.com/>), готовы начать добычу на астероидах и других планетарных телах. Есть вероятность заработать огромные суммы (<http://theweek.com/articles/462830/how-asteroid-mining-could-add-trillions-world-economy>). Мы включили этот раздел в главу потому, что на Земле буквально исчерпаны ресурсы, а большинство остающихся месторождений требует исключительно новых методов добычи. Данный конкретный бизнес взлетит скорее раньше, чем позже.



ВНИМАНИЕ

Сегодня в этом специфическом бизнесе много спекуляций, включая добычу на астероиде (16) Психея (<https://www.usatoday.com/story/tech/nation-now/2017/01/18/nasa-planning-mission-asteroid-worth-10000-quadrillion/96709250/>). Но даже в этом случае людям в конечном счете придется создать уникальную программу утилизации отходов, которая кажется маловероятной, или найти ресурсы где-то в другом месте (весьма вероятно, в космосе). Люди, делающие деньги на данном конкретном проекте уже сегодня, поставляют инструментальные средства, как правило, на базе искусственного интеллекта, позволяющие определить наилучший способ решения задачи.

Исследование других планет

Кажется вполне вероятным, что люди рано или поздно исследуют и даже колонизируют другие планеты. Марс, вероятно, является первым кандидатом. Фактически 78 тысяч человек уже подписались на такой рейс (см. <http://>

newsfeed.time.com/2013/05/09/78000-people-apply-for-one-way-trip-to-mars/). Многие полагают, что после того, как люди достигнут других миров, включая Луну, единственным способом делать деньги будет продажа интеллектуальной собственности или, возможно, создание материалов, специфических только для данного мира (<https://www.forbes.com/sites/quora/2016/09/26/is-there-a-fortune-to-be-made-on-mars/#68d630ab6e28>).



ВНИМАНИЕ!

К сожалению, хотя некоторые люди делают деньги сегодня и на этом проекте, мы, вероятно, не увидим фактической прибыли от этих усилий в ближайшее время. Однако некоторые компании получают прибыль уже сегодня, предоставляя различные инструментальные средства, необходимые для разработки проекта. Исследования действительно инвестируют в экономику.



Глава 20

Десять областей, в которых искусственный интеллект терпит неудачу

В ЭТОЙ ГЛАВЕ...

- » Понимание мира
- » Генерация новых идей
- » Понимание состояния человека

В любой книге по искусственному интеллекту должны рассматриваться пути, на которых искусственный интеллект потерпел неудачу и обманул ожидания. Частично эта проблема обсуждалась при рассмотрении исторических причин наступления зим искусственного интеллекта. Но даже после этих обсуждений у вас могло сложиться впечатление, что искусственный интеллект терпел неудачи только из-за того, что не оправдал ожиданий своих чрезмерно восторженных сторонников; но он терпел неудачи, даже удовлетворяя специфические потребности и соответствуя основным требованиям.

Эта глава — о неудачах, не позволивших искусственному интеллекту показать себя и выполнить задачи, необходимые для достижения полного успеха, описанного в других главах. В настоящее время искусственный интеллект — это развивающаяся технология, которая успешна лишь частично (в лучшем случае).



ЗАПОМНИ

Одна из главных проблем, преследующих искусственный интеллект сегодня, состоит в том, что люди продолжают наделять его человеческими качествами и считать тем, чем он не является. Искусственный интеллект получает на входе очищенные данные, анализирует их, находит шаблоны и предоставляет требуемый вывод. Как описано в этой главе, искусственный интеллект ничего не понимает, он не может ни создать, ни обнаружить ничего нового, у него нет никакого внутриличностного знания. Поэтому он никому и ни в чем не может сочувствовать. Критически важная информация этой главы — искусственный интеллект ведет себя так, как заложено разработавшим его человеком, и его интеллект — это только объединение сложного программного обеспечения и обширных объемов данных, анализируемых определенным способом. Лучшее представление об этих и других проблемах дает статья “Asking the Right Questions About AI” по адресу <https://medium.com/@yonatanzunger/asking-the-right-questions-about-ai-7ed2d9820c48>.

Однако куда важнее то, что некоторые люди утверждают, что искусственный интеллект в конечном счете приведет к глобальному отказу, что невозможно с учетом современной технологии. Искусственный интеллект не может внезапно обрести самосознание, поскольку у него нет никаких средств выражения эмоций, обязательных для обладания самосознанием. Как демонстрируется в табл. 1.1 в главе 1, сегодня у искусственного интеллекта отсутствуют некоторые из важнейших видов интеллекта, обязательных для обладания самосознанием. Однако простого обладания этими уровнями интеллекта также не будет достаточно. В людях есть искра — нечто, непонятное ученым. Не поняв, что это за искра, наука не сможет воспроизвести ее как часть искусственного интеллекта.

Понимание

У людей способность постигать является врожденной, а у искусственного интеллекта она полностью отсутствует. Глядя на яблоко, человек видит больше, чем просто набор свойств, связанных с изображением объекта. Люди по-

нимают яблоки с помощью чувств, таких как цвет, вкус и осязание. Мы понимаем, что яблоко съедобно и содержит определенные питательные вещества. Мы имеем представление о яблоках; возможно, мы любим их и чувствуем, что они — наилучшие фрукты. Искусственный интеллект видит объект, с которым связаны некие свойства, значения которых он не понимает; он только манипулирует ими. В следующих разделах описано, как отсутствие понимания заставляет искусственный интеллект в целом терпеть неудачу.

Интерпретация, а не анализ

Как уже не раз упоминалось в этой книге, искусственный интеллект использует алгоритмы, чтобы манипулировать исходными данными и создать вывод. Акцент делается на выполнении анализа данных. Но направление этого анализа контролирует человек, и он же впоследствии должен интерпретировать результаты. Например, искусственный интеллект может выполнить анализ рентгеновского снимка, представляющего потенциальную опухоль. В полученном результате он может подчеркнуть часть снимка, вероятно, содержащую опухоль, чтобы врач мог ее заметить. В противном случае врач мог бы не заметить ее; таким образом, искусственный интеллект, несомненно, оказывает важную услугу. Но даже в этом случае врач все еще должен сделать обзор результата и определить, действительно ли на рентгеновском снимке опухоль. Как описано в нескольких разделах книги, особенно о беспилотных автомобилях в главе 14, время от времени искусственный интеллект легко вводится в заблуждение, когда даже небольшой элемент присутствует в неправильном месте. Следовательно, даже при том, что искусственный интеллект невероятно полезен, позволяя врачу рассмотреть нечто малозаметное для человеческого глаза, он не заслуживает абсолютного доверия, чтобы принимать решения любого вида.

УЧЕТ ЧЕЛОВЕЧЕСКОГО ПОВЕДЕНИЯ

В этой главе нередко фигурирует отсутствие понимания человеческого поведения. Но даже понимания поведения недостаточно, чтобы реплицировать или моделировать поведение. Чтобы это стало доступным для искусственного интеллекта, необходимо формальное математическое понимание поведения. С учетом разнообразия человеческих поведений возникает подозрение в крайне малой вероятности того, что хоть какая-нибудь их формальная математическая модель будет когда-либо создана. Без таких моделей искусственный интеллект не может думать способом, подобным человеческому, или достичь чего-то, напоминающего чувства.

Интерпретация подразумевает также способность видеть поверх данных. Это не способность создавать новые данные, а понимание того, что данные могут указывать и не только на то, что и так очевидно. Например, люди нередко говорят, что некие данные неверны или сфальсифицированы, даже при том что сами данные не представляют тому доказательств. Искусственный интеллект принимает как истинные все данные, а человек знает, что они могут и не быть такими. Точная формализация того, как люди решают эту задачу, в настоящее время невозможна, поскольку люди фактически этого не понимают.

Выход за пределы чистых чисел

Вопреки любому внешнему виду искусственный интеллект работает только с числами. Он не может понимать смысл слов, когда вы, например, с ним говорите; он просто ищет соответствие шаблону после преобразования вашей речи в числовую форму. Теряется сущность сказанного. Даже если бы искусственный интеллект был в состоянии понять слова, он не мог бы сделать так потому, что слова исчезли в процессе лексического анализа. Невозможность искусственного интеллекта понять нечто столь простое, как слова, означает, что перевод искусственным интеллектом с одного языка на другой всегда будет иметь недостатки, и это бесспорно; кто-то должен перевести чувства, кроющиеся за словами, а также сами слова. Слова выражают чувства, а искусственный интеллект этого сделать не может.

Тот же самый процесс преобразования происходит с каждой мыслью, которой обладают люди. Компьютер преобразует взгляд, звук, запах, вкус и касания в числовые представления, а затем находит соответствие шаблону, чтобы создать набор данных, который моделирует реальную практику. Дальше дело куда сложнее, поскольку разные люди обычно переживают одни и те же вещи по-разному. Например, каждый человек воспринимает цвет индивидуально (<https://www.livescience.com/21275-color-red-blue-scientists.html>). Для искусственного интеллекта каждый компьютер видит цвет совершенно одинаково, а значит, искусственный интеллект не может воспринимать цвета индивидуально. Однако из-за преобразований искусственный интеллект фактически не ощущает цвет вообще.

Учет последствий

Искусственный интеллект может анализировать данные, но не принимать моральные или этические решения. Если вы попросите искусственный интеллект сделать выбор, то он всегда выберет вариант с самой высокой вероятностью успеха, если вы не включите также своего рода функцию случайности.

Искусственный интеллект делает этот выбор независимо от возможного результата. В разделе “Беспилотные автомобили и проблема вагонетки” главы 14 эта проблема выражается весьма ясно. Сталкиваясь с выбором сохранить жизнь либо пассажирам автомобиля, либо пешеходам, принимающий решение искусственный интеллект должен иметь инструкции от человека. Искусственный интеллект не способен учитывать последствия, а потому не может участвовать в процессе принятия таких решений.



ВНИМАНИЕ!

Во многих ситуациях недооценка способности искусственного интеллекта решать задачу просто неудобна. В некоторых случаях вам, вероятно, придется решать задачу во второй или третий раз вручную, поскольку искусственному интеллекту не до задачи. Однако, когда дело доходит до последствий, перед вами может встать реальная проблема в дополнение к моральным и этическим проблемам, если вы доверите искусственному интеллекту решать неподходящие для него задачи. Например, разрешение движения беспилотных автомобилей в неподходящем для этого месте, вероятно, будет незаконным, и вы столкнетесь с юридическими проблемами в дополнение к возмещению ущерба и медицинским расходам, причиной которых может стать беспилотный автомобиль. Короче говоря, прежде чем доверять искусственному интеллекту нечто, подразумевающее потенциальные последствия, имеет смысл ознакомиться с юридическими требованиями.

Открытия

Искусственный интеллект может интерполировать существующее знание, но не может его экстраполировать, чтобы создать новое знание. Столкнувшись с новой ситуацией, он обычно пытается решить ее исходя из существующих знаний, вместо того чтобы придумать что-то новое. Фактически у искусственного интеллекта нет никакого способа создавать что-то новое или видеть нечто как-то уникально. Только человеческие чувства помогают нам открывать новые вещи, работать с ними, изобретать методы для взаимодействия с ними и создавать новые методы для того, чтобы использовать их для выполнения новых задач или лучшего решения прежних задач. В следующих разделах описано, как неспособность искусственного интеллекта делать открытия не позволяет ему оправдывать ожидания, которые люди на него возлагают.

Получение новых данных из старых

Одной из наиболее распространенных задач, решаемых людьми, является *экстраполяция* данных; например, если дано А, то каково В? Люди используют существующее знание, чтобы создать новое знание иного вида. Имея одну часть знания, человек может совершить скачок к новой части знания вне области первоначального знания, причем с высокой вероятностью успеха. Люди делают подобные скачки настолько часто, что они стали их второй натурой, с интуицией в виде крайнего случая. Даже дети могут делать такие прогнозы с высокой вероятностью успеха.



ЗАПОМНИ

Наилучшее, на что когда-либо будет способен искусственный интеллект, — это *интерполяция* данных. Например, даны А, В и С, что в промежутке? Возможность успешно интерполировать данные означает, что искусственный интеллект может дополнить шаблон, но не может создать новые данные. Однако, используя хитрые методики программирования, разработчики иногда могут ввести людей в заблуждение, заставив их думать, что данные новы. Присутствие С выглядит новым, тогда как в действительности это не так. Нехватка новых данных может создать условия, при которых будет казаться, что искусственный интеллект решает задачу, но и это не так. Проблема требует нового решения, а не интерполяции существующих решений.

Видение вне шаблонов

В настоящее время искусственный интеллект способен заметить в данных те шаблоны, которые не очевидны для людей. Именно возможность видеть эти шаблоны делает искусственный интеллект настолько ценным. Манипулирование данными и их анализ являются трудоемким, сложным и монотонным делом, но искусственный интеллект легко может выполнить эту задачу. Однако шаблоны данных — это просто вывод и необязательно решение. Люди полагаются на пять органов чувств, сочувствие, творческий потенциал и интуицию, чтобы видеть вне шаблонов и находить потенциальное решение вне рамок, обусловленных данными. Более подробно эта сфера человеческих возможностей обсуждается в главе 18.



СОВЕТ

Проще всего понять человеческую способность видеть вне шаблонов — это посмотреть на небо. В облачный день люди могут видеть шаблоны в облаках, но искусственный интеллект видит облака и только облака. Кроме того, два человека могут видеть различные вещи в одном и том же наборе облаков. При творческом взгляде на шаблоны в облаках один человек может видеть овцу, а другой —

фонтан. То же самое относится к звездам и другим видам шаблонов. Искусственный интеллект представляет шаблон как результат, но не понимает его, поскольку отсутствие творческого потенциала не позволяет ему сделать с шаблоном ничего, кроме как указать в отчете, что шаблон существует.

Реализация новых чувств

Впоследствии люди узнали о таких вариациях в человеческих органах чувств, которые фактически нереально корректно передать искусственному интеллекту, поскольку репликация этих чувств на аппаратных средствах сейчас действительно невозможна. Например, способность использовать разные органы чувств для обследования одного и того же ввода (синестезия; см. <https://www.mnn.com/health/fitness-well-being/stories/what-is-synesthesia-and-whats-it-like-to-have-it>) находится вне возможностей искусственного интеллекта.



ТЕХНИЧЕСКИЕ
ПОДРОБНОСТИ

Хотя общеизвестно, что у людей есть пять органов чувств, некоторые источники сейчас утверждают, что фактически люди имеют намного больше пяти обычных чувств (<http://www.todayifoundout.com/index.php/2010/07/humans-have-a-lot-more-than-five-senses/>). Некоторые из этих дополнительных чувств не очень хорошо поняты и едва доказуемы, такие как *магниторецепция* (способность чувствовать магнитные поля, такие как магнитное поле Земли). Это чувство позволяет людям указывать направление, подобно птицам, но в меньшей степени. Поскольку нет никакого способа даже определения этого чувства, его репликация как элемента искусственного интеллекта невозможна.

Сочувствие

Компьютеры ничего не чувствуют. Это необязательно плохо, но в данной главе это рассматривается как недостаток. Без способности чувствовать компьютер не может видеть вещи с точки зрения человека. Не понимая, что значит быть счастливым или грустным, нельзя реагировать на эти эмоции, если программа не обладает функцией анализа выражения лица и других индикаторов, позволяющих ей реагировать соответственно. Даже в этом случае реакция будет предопределенной и склонной к ошибкам. Подумайте, сколько решений вы принимаете на основании эмоций, а не фактов. В следующих разделах рассматривается, как отсутствие сочувствия со стороны искусственного интеллекта во многих случаях препятствует его корректному общению с людьми.

Войти в чужое положение

Идея *войти в чужое положение* означает взгляд на вещи с точки зрения другого человека, когда он испытает чувства, подобные испытываемым другим человеком. Никто действительно не чувствует то же самое, что и кто-то еще, но благодаря сочувствию люди могут себе это представить. Эта форма сочувствия требует сильного внутриличностного интеллекта в качестве отправной точки, чего никогда не будет иметь искусственный интеллект, если он не выработает самоощущение (*сингулярность*, как описано в статье <https://www.technologyreview.com/s/425733/paul-allen-the-singularity-isnt-near/>). Кроме того, искусственный интеллект должен быть в состоянии чувствовать нечто, что в настоящее время для него невозможно, и быть открытым для того, чтобы разделять чувства с некой другой сущностью (обычно — с человеком), что также невозможно. Текущее состояние технологии искусственного интеллекта не позволяет ему чувствовать или понимать любой вид эмоции, что делает сочувствие невозможным.



ЗАПОМНИ

Конечно, возникает вопрос, почему сочувствие столь важно. Отсутствие способности чувствовать то же самое, что чувствует кто-то еще, не позволяет искусственному интеллекту выработать побуждение выполнять определенные действия. Вы можете приказывать искусственному интеллекту выполнить некую задачу, но никакого самостоятельного побуждения у него не возникнет. Следовательно, искусственный интеллект никогда не выполнит определенные задачи, даже при том что эффективность в таких задачах является требованием для выработки навыков и знаний, обязательных для достижения интеллекта, подобного человеческому.

Выработка истинных отношений

Искусственный интеллект строит ваш образ исходя из данных, которые он собрал. Затем он создает шаблон из этих данных и, используя определенные алгоритмы, разрабатывает вывод, который позволит создать впечатление, что вы по крайней мере знакомы. Но поскольку искусственный интеллект не чувствует, он не может ценить вас как человека. Он может служить вам, если вы прикажете ему и если поставленная задача находится в списке его функций, но у него не может быть никаких чувств к вам.

Когда речь заходит о взаимоотношениях, люди должны учитывать как интеллектуальную привязанность, так и чувства. Источником интеллектуальной привязанности зачастую является общая выгода для двух сущностей. К сожалению, у искусственного интеллекта и человека (или любого другого субъекта,

если на то пошло) нет никакой общей выгоды. Искусственный интеллект просто обрабатывает данные, используя определенный алгоритм. Нечто не может претендовать на то, чтобы любить нечто другое, даже если ему это приказано. Эмоциональная привязанность несет с собой риск отказа, что подразумевает самооценку.

Смена точки зрения

Иногда люди могут менять свое мнение на основании чего-то, отличного от фактов. Несмотря на то что все шансы — за конкретный разумный ход действий, эмоции делают предпочтительным другой курс действий. У искусственного интеллекта нет никаких предпочтений, поэтому он не может выбрать другой курс действий по любой причине, отличной от изменения вероятностей, *ограничений* (правил, заставляющих его внести изменения) или требования предоставлять случайный вывод.

Вера в недоказуемое

Вера — это уверенность в истинности чего-то без доказывающих фактов. Как правило, вера имеет форму *доверия*, т.е. верой в искренность другого человека безо всяких доказательств тому, что этот человек заслуживает доверия. Искусственный интеллект не может продемонстрировать ни веры, ни доверия, что является частью причины невозможности экстраполировать знания. Само действие экстраполяции зачастую полагается на догадку, в основе которой лежит вера в истинность чего-то, несмотря на отсутствие любого вида данных, поддерживающих догадку. Поскольку эта способность у искусственного интеллекта отсутствует, он не может продемонстрировать способность проникновения в суть — требование, необходимое для шаблонов мышления, подобных человеческим.



СОВЕТ

За примерами веры далеко ходить не нужно, на ее основе были сделаны многие изобретения, создавшие нечто новое. Одним из самых заметных изобретателей был Эдисон. Например, он предпринял тысячу (а возможно, и больше) попыток создать лампочку. Искусственный интеллект, вероятно, сдался бы после определенного количества попыток в связи с заданными ограничениями. Вы можете просмотреть список людей, совершавших удивительные вещи исходя из веры, по адресу <https://www.uky.edu/~eushe2/Pajares/OnFailingG.html>. Каждый из них является примером того, на что искусственный интеллект не способен, поскольку ему недостает способности мыслить поверх конкретных данных, предоставленных ему в качестве ввода.

Предметный указатель

A

Activation function, 207
AI, 25
 effect, 66
 winter, 37
AI-complete, 68
Algorithm, 66
Alpha-beta pruning, 73
AN, 348
Android, 235
A posteriori probability, 188
Application Specific Integrated
 Circuit, 93
A priori probability, 188
Architecture, 85
Artificial
 Intelligence, 25
 Intuition, 348
ASIC, 93
Automata, 234

B

Backpropagation, 210
Batch learning, 215
BFS, 71
Big data, 44
Breadth-First Search, 71

C

Chatbot, 223
Classification problem, 176
CNN, 214; 217
Conditional probability, 188
Connectionism, 204

ConvNet, 217
Convolutional Neural Network, 214;
 217
Coordinate descent, 75

D

Data
 analysis, 163
 record, 55
Dataset, 55
Decision tree, 196
Deduction, 196
Deep learning, 38
Depth-First Search, 71
DFS, 71
Display adapter, 90
Drone, 239; 249

E

End-to-end learning, 217

F

Feed-forward input, 208
Field, 55
First-order logic, 79
FPGA, 94
Frame-of-reference mistruth, 61
Fuzzy logic, 80

G

GAN, 222; 226
Generative Adversarial Network, 222;
 226
Generative-based model, 225
Geo-fencing, 262

GNMT, 150
Google Neural Machine
Translation, 150
GPU, 90
Graph, 69
Graphic Processing Unit, 90

H

Harvard Architecture, 87
Heuristic, 74
Hill-climbing optimization, 75
Hurtful truth, 289
Hypothesis space, 170

I

IA, 158
ICE, 122
Induction, 196
Industrial Communication Engine, 122
Inference engine, 79
Intelligence Augmentation, 158
Internet of Things, 48
IoT, 48

K

Knowledge base, 79

L

Local search, 73

M

Machine learning, 37; 81
Manicuring, 54
Master algorithm, 40
Min-max approximation, 72
Mistruth, 57
 of bias, 61
 of commission, 58
 of omission, 59
 of perspective, 59
Multithreading, 90

N

Naïve Bayes, 189
Neural network architecture, 208
Neuron, 206
NLP, 224
NP-полная задача, 67

O

Objective function, 75
Online learning, 215
Outright error, 112
Overfitting, 174

P

Pattern, 213
Perception, 247
Perceptron, 206

R

Recurrent Neural Network, 222
Regression problem, 176
Reinforcement learning, 82; 177
Reliable, 51
RNN, 222
Robot Process Automation, 120
RPA, 120

S

SD car, 263
Seasteading, 325
Sensor, 246
Simulated annealing, 75
Singularity, 40
State space, 68
State-space search, 68
Supervised learning, 176
Symbolism, 196

T

Tabu search, 75
Target function, 169

Tensor Processing Unit, 93

TPU, 93

Transfer learning, 215

Trolley problem, 271

Twiddle, 75

U

Unsupervised learning, 177

V

Vanishing gradient, 211

Von Neumann bottleneck, 84

W

Weight, 209

A

Автомат, 234

Автоматизация роботизированных процессов, 120

Автоматизированный сбор данных, 53

Автономность, 258

Алгоритм, 49; 66
возврата, 75

Альфа-бета-отсечение, 73

Анализ данных, 66; 163

Андроид, 235

Апостериорная вероятность, 188

Априорная вероятность, 188

Архитектура, 85

нейронной сети, 208
фон Неймана, 84

Б

База знаний, 79

Байесовская сеть, 184; 194

Беспилотный автомобиль, 263

Большие данные, 44

В

Ввод с прямой подачей, 208

Вера, 369

Верховный алгоритм, 40

Видеоадаптер, 90

Восприятие, 247

Г

Гарвардская архитектура, 87

Генеративная модель, 225

Генеративно-состязательная сеть, 222; 226

Геозонирование, 262

Глубокое обучение, 38

Горькая правда, 289

Граф, 69

Графический процессор, 90

Групповое обучение, 215

Гуманоидный робот, 239

Д

Дедукция, 196

Дерево, 68

Дерево решений, 184; 196

Дистанционное обучение, 215

Дрон, 239; 249

З

Задача

классификации, 176

регрессии, 176

Закон

Амары, 297

Гордона Мура, 45

Запись данных, 55

Зима искусственного интеллекта, 37; 285

И

Имитация отжига, 75

Индукция, 196

Индустриальный механизм
связи, 122

Интеллект

сильный, 34

слабый, 34

Интерлингва, 150

Интернет вещей, 48

Интерполяция, 366

Интуиция, 347

Искусственная интуиция, 348

Искусственный интеллект, 25

Исправление, 107

Источник данных, 51

Исчезающий градиент, 211

К

Китайская комната, 104

Коннекционизм, 204

Контролируемое обучение, 176

Координация, 258

Корабль поколений, 326

Космическое поселение, 325

Коэффициент, 209

Критерий отбора, 75

Л

Логика первого порядка, 79

Локальный поиск, 73

М

Маникюр данных, 54

Манипулятор, 238

Машинное обучение, 37; 81; 168

Механизм логического вывода, 79

Многопоточность, 90

Мобильный манипулятор, 239

Моделирование, 164

Н

Набор данных, 55

Надежные данные, 51

Наивный байесовский
классификатор, 183; 189

Недостоверность, 57

недопонимания, 61

предубежденности, 61

точки зрения, 59

умолчания, 59

усердия, 58

Нейрон, 206

Нейронный машинный перевод

Google, 150

Неконтролируемое обучение, 177

Нечеткая логика, 80

О

Обратное распространение

ошибки, 210

Обучение с подкреплением, 82; 177

Оптимизация восхождением к

вершине, 75

П

Паноптикум, 270

Парадокс Моравека, 277

Параллелизм, 214

Перенос обучения, 215

Переобучение, 174

Перцептрон, 206

Подход приближения к мини-

максу, 72

Поиск в

глубину, 71

пространстве состояний, 68

ширину, 71

Поиск с запретами, 75

Покоординатный спуск, 75

Поле, 55

Полный искусственный
интеллект, 68

Преобразование, 163

Проблема вагонетки, 271

Проверка, 163

Пространство

гипотез, 170

состояний, 68

Прямая ошибка, 112

Р

Революция данных, 44

Рекомендация, 110

Рекуррентная нейронная сеть, 222

Робот, 234

С

Сверточная нейронная сеть, 214;
217

Сенсор, 246; 278

Символизм, 196

Символическая логика, 35

Симпатия, 131

Систейдинг, 325

Сквозное обучение, 217

Случайный выбор, 75

Сочувствие, 131

Специализированная интегральная
микросхема, 93

Т

Тензорный процессор, 93

Теорема Байеса, 190

Тест Тьюринга, 224

Технологическая сингулярность, 40

У

Узкое место Фон Неймана, 84

Усиление интеллекта, 158

Условная вероятность, 188

Ф

Функция активизации, 207

Ц

Целевая функция, 169

Ч

Чатбот, 223

Чистка, 163

Ш

Шаблон, 213

Э

Эвристика, 74

Экспертная система, 36

Экстраполяция, 366

Эмодзи, 149

Эмотикон, 149

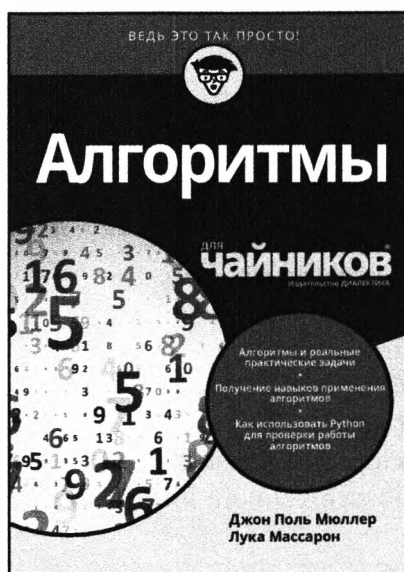
Эффект

зловещей долины, 240

искусственного интеллекта, 66

АЛГОРИТМЫ ДЛЯ ЧАЙНИКОВ

**Джон Поль Мюллер
Лука Массарон**



www.dialektika.com

Эта книга — действительно книга для чайников, поскольку основная ее задача не научить программировать реализации тех или иных давно известных алгоритмов, а познакомить вас с тем, что же такое алгоритмы, как они влияют на нашу повседневную жизнь, и каково состояние дел в этой области человеческих знаний сегодня. В книге рассматривается крайне широкий спектр вопросов, связанных с алгоритмами — это и стандартные сортировка и поиск, и работа с графами (но с уклоном не в стандартные базовые алгоритмы, а в приложения их к таким явлениям сегодняшнего дня, как, например, социальные сети), работа с большими данными и вопросы искусственного интеллекта. При этом материал книги — не просто отвлеченный рассказ о том или ином аспекте современных алгоритмов, но и демонстрация реализаций алгоритмов с конкретными примерами на языке программирования Python. Книга будет полезна всем, кто интересуется современным состоянием дел в области программирования и алгоритмов.

ISBN 978-5-9909446-2-6

в продаже

СОЗДАЕМ НЕЙРОННУЮ СЕТЬ

Тарик Рашид



www.williamspublishing.com

Эта книга представляет собой введение в теорию и практику создания нейронных сетей. Она предназначена для тех, кто хочет узнать, что такое нейронные сети, где они применяются и как самому создать такую сеть, не имея опыта работы в данной области. Изложение материала сопровождается подробным описанием процедуры поэтапного создания полностью функционального кода, который реализует нейронную сеть на языке Python и способен выполняться даже на таком миниатюрном компьютере, как Raspberry Pi Zero.

Основные темы книги:

- нейронные сети и системы искусственного интеллекта;
- структура нейронных сетей;
- сглаживание сигналов, распространяющихся по нейронной сети, с помощью функции активации;
- тренировка и тестирование нейронных сетей;
- интерактивная среда программирования IPython;
- распознавание образов с помощью нейронных сетей.

ISBN 978-5-9909445-7-2 **в продаже**

Искусственный интеллект для чайников®

Шпаргалка

Искусственный интеллект (AI) является технологией, которой уделяется повышенное внимание в фильмах, книгах, играх и т.д. Производители всего этого зачастую приравнивают искусственный интеллект к разуму: вы покупаете умное устройство, но получаете устройство с искусственным интеллектом, весь ум которого иногда заключается только в возможности общения и даже не содержит настоящего искусственного интеллекта. Многие товары рекламируются как содержащие искусственный интеллект, хотя на самом деле его в них нет. Конечно, некоторые производители очень хотят попасть в заголовки газет и сайтов невзирая на достоверность или правдивость сообщений. Эта шпаргалка не развеет все мифы, но даст представление о некоторых интересных возможностях, чтобы вы могли понять, почему искусственный интеллект так часто используется для решения рутинных задач. Да, его можно использовать для некоторых удивительных вещей, но производители зачастую настолько все искажают, что никто не может понять, насколько все действительно реально, а насколько это всего лишь результат чьего-то хорошего воображения.

7 видов интеллекта

Люди демонстрируют семь видов интеллекта. Эти виды интеллекта позволяют отличить людей от других мыслящих сущностей, таких как искусственный интеллект. Кроме того, понимание этих видов интеллекта помогает осознать, в чем люди всегда будут превосходить искусственный интеллект. Многие боятся, что устройства с искусственным интеллектом захватят мир и в конечном счете заменят людей. Да, искусственный интеллект может стать весьма умным в определенной области интеллекта, но не как человек; он никогда не сможет продемонстрировать определенные виды интеллекта, поскольку на самом деле мы не понимаем их сами.

Тип	Возможность симуляции	Человеческие инструменты	Описани
Визуально-пространственный	Средняя	Модели, графика, диаграммы, фотографии, рисунки, трехмерные модели, видео, телевидение и средства массовой информации	Интеллект физической среды, используется, например, моряками и архитекторами (и многими другими). Чтобы двигаться вообще, людям нужно ориентироваться в своем физическом пространстве, учитывая его размеры и особенности. Эта способность нужна интеллекту каждого робота или ноутбука, но ее зачастую трудно смоделировать (как с самоуправляемыми автомобилями) или сделать достаточно точной (как с пылесосом, который, обходя препятствия, создает впечатление осмысленного движения)



TM

BESTSELLING
BOOK
SERIES

Искусственный интеллект для чайников®



TM

СЕРИЯ
ПЕЧАТНЫХ
КНИГ ОТ
АИИТЕКНИКИ*Шпаргалка*

Тип	Возможность симуляции	Человеческие инструменты	Описание
Телесно-кинестетический	От средней до высокой	Специализированное оборудование и реальные объекты	Движения тела, как у хирурга или балерины, требуют точности и понимания своего тела. Роботы обычно используют этот вид интеллекта для выполнения повторяющихся задач зачастую куда точнее, чем люди, но иногда с меньшим изяществом. Следует делать различие между средством усиления возможностей человека, таким как хирургическое устройство, улучшающее физические способности хирурга, и истинным независимым движением. Вышесказанное — это просто демонстрация математической возможности, в которой все зависит от действий хирурга
Творческий	Нет	Художественная деятельность, новый образ мыслей, изобретения, новые виды музыкальных произведений	Творчество — это действие по выработке нового образа мыслей, приводящее к уникальному результату в виде произведения искусства, музыки или литературы. Действительно, новое произведение — это результат творчества. Искусственный интеллект способен моделировать существующий образ мыслей и даже комбинировать их несколько, чтобы создать уникальное представление о том, что уже существует, но в действительности является только математически ориентированной версией существующего образа. Для творчества искусственный интеллект должен обладать самосознанием, которое требует внутриличностного интеллекта
Межличностный	От низкой до средней	Телефон, звуковая конференц-связь, видео конференц-связь, записи, компьютерная конференц-связь, электронная почта	Взаимодействие с другими происходит на нескольких уровнях. Задача этого вида интеллекта — получать и предоставлять информацию, а также обмениваться и манипулировать ею на основании опыта других. Компьютеры способны отвечать на простые вопросы на основании введенных ключевых слов, а не потому, что понимают вопрос. Интеллект действует, получая информацию, находя подходящие ключевые слова и предоставляя затем информацию на основании этих ключевых слов. Поиск терминов в таблице с перекрестными ссылками и последующими действиями согласно предоставляемым таблицей инструкциям демонстрирует логический интеллект, а не межличностный

Искусственный интеллект для чайников®

Шпаргалка

Тип	Возможность симуляции	Человеческие инструменты	Описание
Внутри-личностный	Нет	Книги, творческие материалы, дневники, уединение и время	Взгляд внутрь себя, чтобы понять собственные интересы, и последующая выработка цели на основании этих интересов являются в настоящее время прерогативой только человеческого интеллекта. Как у машин у компьютеров нет никаких нужд, желаний, интересов или творческих способностей. Искусственный интеллект обрабатывает введенные числа, используя набор алгоритмов, и предоставляет вывод, но не знает ни о том, что он делает, ни о том, зачем он это делает
Лингвистический	Низкая	Игры, средства массовой информации, книги, устройства звукозаписи и произносимые слова	Работа со словами — это немаловажный инструмент общения, поскольку разговор и обмен письменной информацией происходят намного быстрее, чем в любой другой форме. Этот вид интеллекта подразумевает понимание речевого и письменного ввода, его обработку для формирования ответа и предоставление результата в качестве понятного ответа. Во многих случаях компьютеры могут лишь анализировать ввод по ключевым словам, но вопрос фактически не могут понять вообще
Логико-математический	Высокая	Логические игры, исследования, тайны и головоломки	Вычисление результата, осуществление сравнений, исследование шаблонов и выявление взаимозависимостей — все это области, в которых компьютеры сейчас превосходны. Победа компьютера над человеком на телевикторине — это практически единственный общеизвестный вид интеллекта из семи. Да, вы могли бы замечать небольшие фрагменты других видов интеллекта, но в фокусе находится этот. Оценка различий интеллекта человека и компьютера на базе только одной области — это отнюдь не лучшая идея

Искусственный интеллект для чайников®

Шпаргалка

Наиболее популярные реальные случаи использования искусственного интеллекта

При фактическом использовании искусственного интеллекта возникает два типа заблуждений. Первый тип имеет отношение к интеллектуальным устройствам, которые просто обеспечивают подключение к серверному приложению, создавая впечатление применения искусственного интеллекта. Например, интеллектуальный термометр может подключаться к вашему смартфону, но никак не полагается на искусственный интеллект, чтобы что-нибудь сделать. Однако термометр, программирующий сам себя на основании устанавливаемой вами температуры в доме, действительно полагается на искусственный интеллект, чтобы обеспечить дополнительные функциональные возможности.

Второй тип заблуждений имеет отношение к устройствам, которые используют искусственный интеллект, но не тем способом, которым он, вероятно, мог бы работать. Например, интеллектуальный помощник, предназначенный для поиска наилучших решений, обречен на неудачу, поскольку принятие решений находится вне области возможностей искусственного интеллекта. С другой стороны, интеллектуальный помощник, который помогает находить ресторан, включать свет и хранить список встреч (гарантируя отсутствие в них накладок), вероятно, будет работать, если в приложении нет ошибок и вы предоставите соответствующие исходные данные.

В следующей таблице приведены продукты, которые в настоящее время доступны, относительно автономны, достаточно недороги для многих людей и фактически работают. Все они в некотором роде полагаются на искусственный интеллект.

Продукт	URL	Описание
Arterys	https://arterys.com/	Движения тела, как у хирурга или балерины, требуют точности и понимания своего тела. Выполняет сканирование сердца за 6–10 мин, а не за час, как обычно. Пациентам не приходится также задерживать дыхание. Удивительно, но эта система получает несколько размерностей данных (анатомия сердца, скорость кровотока и направление кровотока) за очень короткое время
Clocky	https://nandahome.com/	Действует как будильник для тех, кому нелегко встать утром. Устройство дает вам один шанс подняться, а затем перемещается в случайном направлении, чтобы вынудить вас встать с кровати, чтобы его выключить
Enlitic	https://www.enlitic.com/	Анализирует рентгеновские снимки за миллисекунды — в 10 тысяч раз быстрее, чем рентгенолог. Кроме того, система на 50% лучше в классификации опухолей и имеет более низкий коэффициент ошибочно негативных диагнозов (0 против 7% у людей)

Искусственный интеллект для чайников®

Шпаргалка

Продукт	URL	Описание
Hom-Bo	http://www.lg.com/us/vacuum-cleaners/lg-CR5765GD	Пылесосит ковры и полы. У этого робота превосходный искусственный интеллект наряду со многими интеллектуальными сенсорами, поэтому обычно ему удается избежать столкновений с предметами. Его можно также запрограммировать на использование различных стратегий чистки (для гарантии, что он ничего не пропустит, следуя все время по одному и тому же шаблону чистки)
K'Watch	http://www.pkvitality.com/ktrack-glucose/	Обеспечивает постоянный контроль глюкозы, а также предоставляет приложение, которое люди могут использовать для получения полезной информации о состоянии своей диабетической болезни
Moov	https://welcome.moov.cc/	Контролирует ритм сердца и его трехмерное движение. Искусственный интеллект этого устройства отслеживает статистику и дает советы относительно лучшей разминки. Фактически вы получаете советы о том, как лучше ставить ногу во время бега и стоит ли увеличить шаг. Задача этого устройства — гарантировать наилучшую разминку, которая улучшит здоровье, без риска получить травму
QardioCore	https://www.getqardio.com/	Снимает кардиограмму без использования проводов, допускается применение персоналом без высокой медицинской квалификации. Как и многие подобные устройства, это полагается на ваш смартфон, чтобы выполнить анализ и связаться при необходимости с внешними источниками
Robomow	https://www.robomow.com/	Косит траву
Roomba	http://www.irobot.com/	Пылесосит ковры и полы. Имеет привычку врезаться в предметы; не видит их и не может избежать столкновений, поэтому его искусственный интеллект чрезвычайно прост. Его аналог, Braava, моет полы шваброй, а Miraа чистит бассейн. Если пол нужно и пропылесосить, и вымыть, то вместо них можно использовать Scooba
Sentrian	http://sentrian.com/	Контролирует уровень сахара в крови и другую статистику хронических болезней, позволяя использовать собранные данные для прогноза болезни прежде, чем она обострится. Вносит изменения в список лекарств и поведение пациента прежде, чем событие сможет произойти; Sentrian снижает количество неизбежных госпитализаций, улучшая качество жизни пациентов и сокращая медицинские затраты

Искусственный интеллект для чайников®

Журнал

Основные поставщики искусственного интеллекта

Невозможно перечислить всех производителей, имеющих отношение к искусственному интеллекту. Их довольно много, и небольшие компании нередко быстро закрываются, так как исследования довольно дороги. Вот список основных производителей искусственного интеллекта, на которых имеет смысл обратить внимание.

- Amazon
- Apple
- Baidu
- Cylance
- Deloitte
- Electronic Arts
- Facebook
- Google
- IBM
- Intel
- LinkedIn
- Lockheed Martin
- Microsoft
- MITRE
- NASA
- NVidia
- Sizmek, первоначально — Rocket Fuel
- Sentient Corporation
- Tesla
- Uber

Искусственный интеллект для чайников®

Шпаргалка

Основные производители искусственного интеллекта

Искусственный интеллект используется не во всех отраслях промышленности. Некоторые отрасли занимают выжидательную позицию, когда дело доходит до искусственного интеллекта, поскольку он все еще не доказал своей ценности, а владельцы этих сфер помнят прошлые зимы искусственного интеллекта. Кроме того, исследования в области искусственного интеллекта сосредоточены лишь на определенных отраслях промышленности, в которых искусственный интеллект фактически работает. Искусственный интеллект требует большого количества исходных данных и полагается на алгоритмы для обработки этих данных, чтобы затем предоставить вывод, который в наилучшем случае соответствует требованиям. Некоторые отрасли промышленности даже не могут соответствовать основным требованиям, и, что гораздо важнее, осталось сделать еще очень многое, чтобы искусственный интеллект стал полностью пригодным для использования. Вот основные отрасли промышленности, которые инвестируют в исследования искусственного интеллекта и используют его.

Область инвестирования	Типы приложений	Процент
Отрасли, связанные с применением компьютеров	Аппаратные средства, программное обеспечение и IT	33,33
Связь	Телекоммуникации, Интернет и сетевые средства массовой информации	14,30
Бизнес-услуги	Маркетинг, реклама, консалтинг по менеджменту, финансовые услуги и банковское дело	8,48
Промышленные товары	Производство автомобилей, электрического, электронного и полупроводникового оборудования	5,72
Потребительские товары	Розничная продажа, бытовая электроника и развлечения	5,21
Исследования	Поиск в Сети	3,37
Здравоохранение	Больницы, организация здравоохранения, страхование от болезней	2,86
Правительство	Администрирование	2,04
Управление ресурсами	Нефть и электроэнергия	1,33

Что такое искусственный интеллект?

Искусственный интеллект является захватывающим и немного жутковатым. Он вокруг нас. Искусственный интеллект помогает защитить от мошенничества, контролировать расписание медицинских процедур, он способен работать в клиентской службе и даже помогает вам в выборе телешоу и приборке вашего дома. Хотите узнать больше? Эта книга восполняет пробелы, знакомя вас с тем, что представляет собой искусственный интеллект и чем он не является, рассматриваются также этические вопросы использования искусственного интеллекта, его современное применение и некоторые из удивительных вещей, на которые он, вероятно, будет способен завтра. Независимо от уровня образования вы будете очарованы тем, что узнаете!

В книге...

- История искусственного интеллекта
- Роль данных
- Как искусственный интеллект используется в компьютерных приложениях, медицине, космосе и машинном обучении
- Мифы, окружающие искусственный интеллект
- Роботы и дроны

Джон Мюллер — автор более 108 книг и 600 статей на темы от искусственного интеллекта и работы с сетями до управления базами данных. Он также является техническим редактором и консультантом.

Лука Массарон — аналитик данных и директор по маркетинговому исследованию, специализирующийся на многомерном статистическом анализе, машинном обучении и изучении ожиданий потребителей.

 **ДИАЛЕКТИКА**

www.dialektika.com

Изображение на обложке:
@Depositphotos.com/27247835
Автор: VLADGRIN

Искусственный интеллект

ISBN 978-5-907114-57-9



9 785907 114579


для
Чайников