

经 典 原 版 书 库

Spark数据分析

基于Python语言

[澳] 杰夫瑞·艾文 (Jeffrey Aven) 著

(英文版)

ADDISON
WESLEY
DATA &
ANALYTICS
SERIES

Data Analytics
with **SPARK**
Using **PYTHON**



JEFFREY AVEN



机械工业出版社
China Machine Press

经

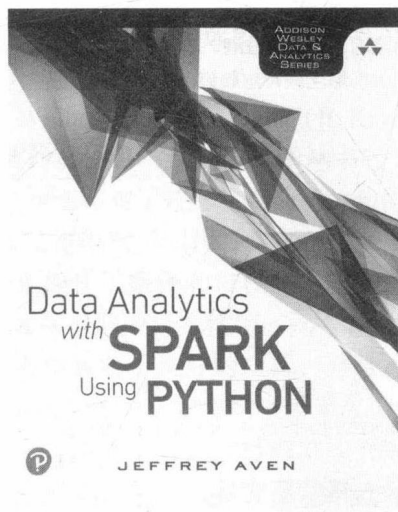
库

Spark数据分析

基于Python语言

(英文版)

*Data Analytics with
Spark Using Python*



[澳] 杰夫瑞·艾文 (Jeffrey Aven) 著



机械工业出版社
China Machine Press

图书在版编目 (CIP) 数据

Spark 数据分析: 基于 Python 语言 (英文版) / (澳) 杰夫瑞·艾文 (Jeffrey Aven) 著. —北京: 机械工业出版社, 2019.4

(经典原版书库)

书名原文: Data Analytics with Spark Using Python

ISBN 978-7-111-62003-7

I. S… II. 杰… III. 数据处理软件-英文 IV. TP274

中国版本图书馆 CIP 数据核字 (2019) 第 031718 号

本书版权登记号: 图字 01-2018-7390

Authorized translation from the English language edition, entitled Data Analytics with Spark Using Python, ISBN: 9780134846019, by Jeffrey Aven, published by Pearson Education, Inc., Copyright © 2018 Pearson Education, Inc.

All rights reserved. No part of this book may be reproduced or transmitted in any form or by any means, electronic or mechanical, including photocopying, recording or by any information storage retrieval system, without permission from Pearson Education, Inc.

English language edition published by China Machine Press, Copyright © 2019.

本书英文影印版由 Pearson Education Inc. 授权机械工业出版社独家出版。未经出版者书面许可, 不得以任何方式复制或抄袭本书内容。

此影印版仅限于中华人民共和国境内 (不包括香港、澳门特别行政区及台湾地区) 销售发行。

本书封面贴有 Pearson Education (培生教育出版集团) 激光防伪标签, 无标签者不得销售。

出版发行: 机械工业出版社 (北京市西城区百万庄大街 22 号 邮政编码: 100037)

责任编辑: 张梦玲

责任校对: 李秋荣

印 刷: 北京市荣盛彩色印刷有限公司

版 次: 2019 年 3 月第 1 版第 1 次印刷

开 本: 186mm×240mm 1/16

印 张: 18.25

书 号: ISBN 978-7-111-62003-7

定 价: 79.00 元

凡购本书, 如有缺页、倒页、脱页, 由本社发行部调换

客服热线: (010) 88378991 88361066

投稿热线: (010) 88379604

购书热线: (010) 68326294 88379649 68995259

读者信箱: hzjsj@hzbook.com

版权所有·侵权必究

封底无防伪标均为盗版

本书法律顾问: 北京大成律师事务所 韩光/邹晓东

出版者的话

文艺复兴以来，源远流长的科学精神和逐步形成的学术规范，使西方国家在自然科学的各个领域取得了垄断性的优势；也正是这样的优势，使美国在信息技术发展的六十多年间名家辈出、独领风骚。在商业化的进程中，美国的产业界与教育界越来越紧密地结合，计算机学科中的许多泰山北斗同时身处科研和教学的最前线，由此而产生的经典科学著作，不仅擘划了研究的范畴，还揭示了学术的源变，既遵循学术规范，又自有学者个性，其价值并不会因年月的流逝而减退。

近年，在全球信息化大潮的推动下，我国的计算机产业发展迅猛，对专业人才的需求日益迫切。这对计算机教育界和出版界都既是机遇，也是挑战；而专业教材的建设在教育战略上显得举足轻重。在我国信息技术发展时间较短的现状下，美国等发达国家在其计算机科学发展的几十年间积淀和发展的经典教材仍有许多值得借鉴之处。因此，引进一批国外优秀计算机教材将对我国计算机教育事业的发展起到积极的推动作用，也是与世界接轨、建设真正的世界一流大学的必由之路。

机械工业出版社华章公司较早意识到“出版要为教育服务”。自1998年开始，我们就将工作重点放在了遴选、移译国外优秀教材上。经过多年的不懈努力，我们与 Pearson、McGraw-Hill、Elsevier、MIT、John Wiley & Sons、Cengage 等世界著名出版公司建立了良好的合作关系，从它们现有的数百种教材中甄选出 Andrew S. Tanenbaum、Bjarne Stroustrup、Brian W. Kernighan、Dennis Ritchie、Jim Gray、Afred V. Aho、John E. Hopcroft、Jeffrey D. Ullman、Abraham Silberschatz、William Stallings、Donald E. Knuth、John L. Hennessy、Larry L. Peterson 等大师名家的一批经典作品，以“计算机科学丛书”为总称出版，供读者学习、研究及珍藏。大理石纹理的封面，也正体现了这套丛书的品位和格调。

“计算机科学丛书”的出版工作得到了国内外学者的鼎力相助，国内的专家不仅提供了中肯的选题指导，还不辞劳苦地担任了翻译和审校的工作；而原书的作者也相当关注其作品在中国的传播，有的还专门为其书的中译本作序。迄今，“计算机科学丛书”已经出版了近500个品种，这些书籍在读者中树立了良好的口碑，并被许多高校采用为正式教材和参考书籍。其影印版“经典原版书库”作为姊妹篇也被越来越多实施双语教学的学校所采用。

权威的作者、经典的教材、一流的译者、严格的审校、精细的编辑，这些因素使我们的图书有了质量的保证。随着计算机科学与技术专业学科建设的不断完善和教材改革的逐渐深化，教育界对国外计算机教材的需求和应用都将步入一个新的阶段，我们的目标是尽善尽美，而反馈的意见正是我们达到这一终极目标的重要帮助。华章公司欢迎老师和读者对我们的工作提出建议或给予指正，我们的联系方式如下：

华章网站：www.hzbook.com

电子邮件：hzsj@hzbook.com

联系电话：(010) 88379604

联系地址：北京市西城区百万庄南街1号

邮政编码：100037



华章科技图书出版中心

前 言

Spark 在这场由大数据与开源软件掀起的颠覆性革命中处于核心位置。不论是尝试 Spark 的意向还是实际用例的数量都在以几何级数增长，而且毫无衰退的迹象。本书将手把手引导你在大数据分析领域中收获事业上的成功。

本书重点

本书重点关注 Spark 项目的基本知识，从 Spark 核心开始，然后拓展到各种 Spark 扩展、Spark 相关项目及子项目，以及 Spark 所处的丰富的生态系统里各种别的开源技术，比如 Hadoop、Kafka、Cassandra 等。

尽管对于本书所介绍的 Spark 基本概念（包括运行环境、集群架构、应用架构等）的理解是与编程语言无关且透明的，本书中大多数示例程序和练习是用 Python 实现的。Spark 的 Python API (PySpark) 为数据分析师、数据工程师、数据科学家等提供了易用的编程环境，让开发者能在获得 Python 语言的灵活性和可扩展性的同时，获得 Spark 的分布式处理能力和伸缩性。

本书所涉及的范围非常广泛，涵盖了从基本的 Spark 核心编程到 Spark SQL、Spark Streaming、机器学习等方方面面的内容。本书对于每个主题都给出了良好的介绍和概览，足以让你以 Spark 项目为基础构建出针对任何特定领域或学科的平台。

目标读者

本书是为有志进入大数据领域或是已经入门想要进一步巩固大数据领域知识的数据分析师和工程师而写的。当前市场对于具备大数据技能、懂得大数据领域优秀处理框架 Spark 的工程师的需求特别大。本书的目标是针对这一不断增长的雇佣市场需求培训读者，使得读者获得雇主亟需的技能。

对于阅读本书来说，有 Python 使用经验是有帮助的，没有的话也没关系，毕竟 Python 对于任何有编程经验的人来说都非常直观易懂。读者最好对数据分析和数据处理有一定了解。这本书尤其适合有兴趣进入大数据领域的数据仓库技术人员阅读。

如何使用本书

本书分为两大部分共 8 章。第一部分“Spark 基础”包括 4 章，会让读者深刻理解 Spark 是什么，如何部署 Spark，如何使用 Spark 进行基本的数据处理操作。

第 1 章提供了大数据生态圈的概览，包括 Spark 项目的起源和演进过程。本章讨论了 Spark 项目的关键属性，包括 Spark 是什么，用起来如何，以及 Spark 与 Hadoop 项目之间的关系。

第 2 章展示了如何部署一个 Spark 集群，包括 Spark 集群的各种部署模式，以及调用

Spark 的各种方法。

第 3 章讨论 Spark 集群和应用是如何运作的，让读者深刻理解 Spark 是如何工作的。

第 4 章关注使用弹性分布式数据集 (RDD) 进行 Spark 初级编程的基础。

第二部分“基础拓展”包括后四章的内容，扩展到 Spark 的核心模块以外，包括 SQL 和 NoSQL 系统、流处理应用、数据科学与机器学习中 Spark 的使用：

第 5 章讲解了用来扩展、加速和优化常规 Spark 例程的高级元件，包括各种共享变量和 RDD 存储，以及分区的概念及其实现。

第 6 章讨论 Spark 与广袤的 SQL 土壤的整合，还有 Spark 与非关系型存储的整合。

第 7 章介绍了 Spark 的 Streaming 子项目，以及 Streaming 中最基本的 DStream 对象。本章还涵盖了 Spark 对于 Apache Kafka 这样的常用消息系统的使用。

第 8 章介绍了使用 R 语言使用 Spark 建立预测模型，以及 Spark 中用来实现机器学习的子项目 MLlib。

本书代码

本书中各个练习的示例数据和源代码可以从 <http://sparkusingpython.com> 下载。你也可以从 https://github.com/sparktraining/spark_using_python 查看或者下载本书的 GitHub 代码仓库。

引 言

Spark 是用于大数据的一流的数据处理平台和编程接口，与大数据技术浪潮密不可分。在本书撰写时，Spark 是 Apache 软件基金会（ASF）框架下最活跃的开源项目之一，也是有史以来最活跃的开源大数据项目之一。

数据分析、数据处理以及数据科学社区都对 Spark 有着浓厚兴趣，因而理解 Spark 是什么，Spark 是用来干什么的，Spark 的优势在哪里，以及如何利用 Spark 进行大数据分析都是很重要的。这本书涵盖了上述内容的方方面面。

与其他专门针对 Spark 的出版物不同，后者几乎仅使用 Scala API，而本书专注于 Spark 的 Python API，也就是 PySpark。选择 Python 作为本书基础语言是因为 Python 是一门直观易懂的解释型语言，广为人知而且对于新手也极易上手。更何况 Python 对于数据科学家而言是一门非常常用的编程语言，而数据科学家在 Spark 社区也是相当大的一个群体。

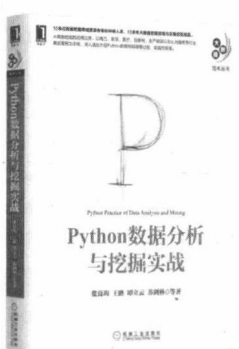
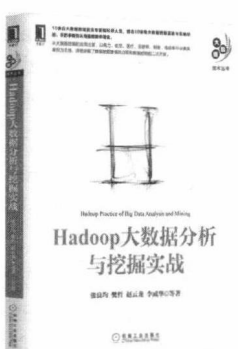
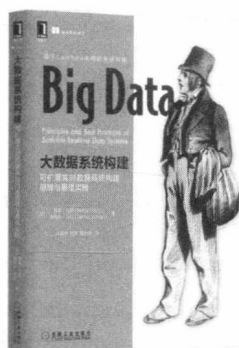
本书会从零开始讲大数据和 Spark，因此无论你是从未接触 Spark 和 Hadoop，还是已经有所接触，正在寻求全面了解 Spark 的运作方式和如何充分利用 Spark 丰富的功能，这本书都是合适的。

在阅读本书的过程中，你会了解一些相近的或者相互协作的平台、项目以及技术，比如 Hadoop、HBase、Kafka 等，并且看到它们如何与 Spark 交互与集成。

过去的几年中，我的工作都集中在这个领域，包括针对大数据分析进行授课和向客户提供咨询服务。我经历了 Spark 和大数据乃至整个开源运动的出现和成熟，并且参与到了开源软件融入企业使用的进程中。我尽量把个人的学习历程总结到了本书中。

希望这本书能助你开启大数据和 Spark 从业人员的旅程。

推荐阅读



目 录

第一部分 Spark 基础

第 1 章 大数据、Hadoop、Spark 介绍

1.1 大数据、分布式计算、Hadoop 简介	3
1.1.1 大数据与 Hadoop 简史	4
1.1.2 Hadoop 详解	5
1.2 Apache Spark 简介	11
1.2.1 Apache Spark 背景	11
1.2.2 Spark 的用途	12
1.2.3 Spark 编程接口	12
1.2.4 Spark 程序的提交类型	12
1.2.5 Spark 应用程序的输入输出类型	14
1.2.6 Spark 中的 RDD	14
1.2.7 Spark 与 Hadoop	14
1.3 Python 函数式编程	15
1.3.1 Python 函数式编程用到的数据结构	15
1.3.2 Python 对象序列化	18
1.3.3 Python 函数式编程基础	21
1.4 本章小结	23

第 2 章 部署 Spark

2.1 Spark 部署模式	25
2.1.1 本地模式	26
2.1.2 Spark 独立集群	26
2.1.3 基于 YARN 运行 Spark	27
2.1.4 基于 Mesos 运行 Spark	28
2.2 准备安装 Spark	28
2.3 获取 Spark	29
2.4 在 Linux 或 Mac OS X 上安装 Spark	30
2.5 在 Windows 上安装 Spark	32
2.6 探索 Spark 安装目录	34

2.7 部署多节点的 Spark 独立集群	35
2.8 在云上部署 Spark	37
2.8.1 AWS	37
2.8.2 GCP	39
2.8.3 Databricks	40
2.9 本章小结	41

第 3 章 理解 Spark 集群架构

3.1 Spark 应用中的术语	43
3.1.1 Spark 驱动器	44
3.1.2 Spark 工作节点与执行器	47
3.1.3 Spark 主进程与集群管理器	49
3.2 使用独立集群的 Spark 应用	51
3.3 在 YARN 上运行 Spark 应用的部署模式	51
3.3.1 客户端模式	52
3.3.2 集群模式	53
3.3.3 回顾本地模式	54
3.4 本章小结	55

第 4 章 Spark 编程基础

4.1 RDD 简介	57
4.2 加载数据到 RDD	59
4.2.1 从文件创建 RDD	59
4.2.2 从文本文件创建 RDD 的方法	61
4.2.3 从对象文件创建 RDD	64
4.2.4 从数据源创建 RDD	64
4.2.5 从 JSON 文件创建 RDD	67
4.2.6 通过编程创建 RDD	69
4.3 RDD 操作	70
4.3.1 RDD 核心概念	70
4.3.2 基本的 RDD 转化操作	75
4.3.3 基本的 RDD 行动操作	79
4.3.4 键值对 RDD 的转化操作	83

4.3.5 MapReduce 与单词计数练习	90
4.3.6 连接操作	93
4.3.7 在 Spark 中连接数据集	98
4.3.8 集合操作	101
4.3.9 数值型 RDD 的操作	103
4.4 本章小结	106

第二部分 基础拓展

第 5 章 Spark 核心 API

高级编程

5.1 Spark 中的共享变量	109
5.1.1 广播变量	110
5.1.2 累加器	114
5.1.3 练习：使用广播变量和累加器	117
5.2 Spark 中的数据分区	118
5.2.1 分区概述	118
5.2.2 掌控分区	119
5.2.3 重分区函数	121
5.2.4 针对分区的 API 方法	123
5.3 RDD 的存储选项	125
5.3.1 回顾 RDD 谱系	125
5.3.2 RDD 存储选项	126
5.3.3 RDD 缓存	129
5.3.4 持久化 RDD	129
5.3.5 选择何时持久化或缓存 RDD	132
5.3.6 保存 RDD 检查点	132
5.3.7 练习：保存 RDD 检查点	134
5.4 使用外部程序处理 RDD	136
5.5 使用 Spark 进行数据采样	137
5.6 理解 Spark 应用与集群配置	139
5.6.1 Spark 环境变量	139
5.6.2 Spark 配置属性	143
5.7 Spark 优化	146
5.7.1 早过滤，勤过滤	147
5.7.2 优化满足结合律的操作	147

5.7.3 理解函数和闭包的影响	149
5.7.4 收集数据的注意事项	150
5.7.5 使用配置参数调节和优化应用	150
5.7.6 避免低效的分区	151
5.7.7 应用性能问题诊断	153
5.8 本章小结	157

第 6 章 使用 Spark 进行 SQL 与 NoSQL 编程

6.1 Spark SQL 简介	159
6.1.1 Hive 简介	160
6.1.2 Spark SQL 架构	164
6.1.3 DataFrame 入门	166
6.1.4 使用 DataFrame	177
6.1.5 DataFrame 缓存、持久化与 重新分区	185
6.1.6 保存 DataFrame 输出	186
6.1.7 访问 Spark SQL	189
6.1.8 练习：使用 Spark SQL	192
6.2 在 Spark 中使用 NoSQL 系统	193
6.2.1 NoSQL 简介	194
6.2.2 在 Spark 中使用 HBase	195
6.2.3 练习：在 Spark 中使用 HBase	198
6.2.4 在 Spark 中使用 Cassandra	200
6.2.5 在 Spark 中使用 DynamoDB	202
6.2.6 其他 NoSQL 平台	204
6.3 本章小结	204

第 7 章 使用 Spark 处理流数据与消息

7.1 Spark Streaming 简介	207
7.1.1 Spark Streaming 架构	208
7.1.2 DStream 简介	209
7.1.3 练习：Spark Streaming 入门	216
7.1.4 状态操作	217
7.1.5 滑动窗口操作	219
7.2 结构化流处理	221

7.2.1 结构化流处理数据源	222
7.2.2 结构化流处理的数据输出池	223
7.2.3 输出模式	224
7.2.4 结构化流处理操作	225
7.3 在 Spark 中使用消息系统	226
7.3.1 Apache Kafka	227
7.3.2 练习：在 Spark 中使用 Kafka	232
7.3.3 亚马逊 Kinesis	235
7.4 本章小结	238

第 8 章 Spark 数据科学与机器学习简介

8.1 Spark 与 R 语言	241
8.1.1 R 语言简介	242
8.1.2 通过 R 语言使用 Spark	248

8.1.3 练习：在 RStudio 中使用 SparkR	255
8.2 Spark 机器学习	257
8.2.1 机器学习基础	257
8.2.2 使用 Spark MLlib 进行机器学习	260
8.2.3 练习：使用 Spark MLlib 实现推荐器	265
8.2.4 使用 Spark ML 进行机器学习	269
8.3 利用笔记本使用 Spark	273
8.3.1 利用 Jupyter (IPython) 笔记本使用 Spark	273
8.3.2 利用 Apache Zeppelin 笔记本使用 Spark	276
8.4 本章小结	277

Contents

I: Spark Foundations

1 Introducing Big Data, Hadoop, and Spark 3

Introduction to Big Data, Distributed Computing, and Hadoop 3

A Brief History of Big Data and Hadoop 4

Hadoop Explained 5

Introduction to Apache Spark 11

Apache Spark Background 11

Uses for Spark 12

Programming Interfaces to Spark 12

Submission Types for Spark Programs 12

Input/Output Types for Spark Applications 14

The Spark RDD 14

Spark and Hadoop 14

Functional Programming Using Python 15

Data Structures Used in Functional Python Programming 15

Python Object Serialization 18

Python Functional Programming Basics 21

Summary 23

2 Deploying Spark 25

Spark Deployment Modes 25

Local Mode 26

Spark Standalone 26

Spark on YARN 27

Spark on Mesos 28

Preparing to Install Spark 28

Getting Spark 29

Installing Spark on Linux or Mac OS X 30

Installing Spark on Windows 32

Exploring the Spark Installation 34

Deploying a Multi-Node Spark Standalone Cluster 35

Deploying Spark in the Cloud	37
Amazon Web Services (AWS)	37
Google Cloud Platform (GCP)	39
Databricks	40
Summary	41
3 Understanding the Spark Cluster Architecture	43
Anatomy of a Spark Application	43
Spark Driver	44
Spark Workers and Executors	47
The Spark Master and Cluster Manager	49
Spark Applications Using the Standalone Scheduler	51
Deployment Modes for Spark Applications Running on YARN	51
Client Mode	52
Cluster Mode	53
Local Mode Revisited	54
Summary	55
4 Learning Spark Programming Basics	57
Introduction to RDDs	57
Loading Data into RDDs	59
Creating an RDD from a File or Files	59
Methods for Creating RDDs from a Text File or Files	61
Creating an RDD from an Object File	64
Creating an RDD from a Data Source	64
Creating RDDs from JSON Files	67
Creating an RDD Programmatically	69
Operations on RDDs	70
Key RDD Concepts	70
Basic RDD Transformations	75
Basic RDD Actions	79
Transformations on PairRDDs	83
MapReduce and Word Count Exercise	90
Join Transformations	93
Joining Datasets in Spark	98
Transformations on Sets	101
Transformations on Numeric RDDs	103
Summary	106

II: Beyond the Basics

5 Advanced Programming Using the Spark Core API 109

- Shared Variables in Spark 109
 - Broadcast Variables 110
 - Accumulators 114
 - Exercise: Using Broadcast Variables and Accumulators 117
- Partitioning Data in Spark 118
 - Partitioning Overview 118
 - Controlling Partitions 119
 - Repartitioning Functions 121
 - Partition-Specific or Partition-Aware API Methods 123
- RDD Storage Options 125
 - RDD Lineage Revisited 125
 - RDD Storage Options 126
 - RDD Caching 129
 - Persisting RDDs 129
 - Choosing When to Persist or Cache RDDs 132
 - Checkpointing RDDs 132
 - Exercise: Checkpointing RDDs 134
- Processing RDDs with External Programs 136
- Data Sampling with Spark 137
- Understanding Spark Application and Cluster Configuration 139
 - Spark Environment Variables 139
 - Spark Configuration Properties 143
- Optimizing Spark 146
 - Filter Early, Filter Often 147
 - Optimizing Associative Operations 147
 - Understanding the Impact of Functions and Closures 149
 - Considerations for Collecting Data 150
 - Configuration Parameters for Tuning and Optimizing Applications 150
 - Avoiding Inefficient Partitioning 151
 - Diagnosing Application Performance Issues 153
- Summary 157

6 SQL and NoSQL Programming with Spark 159

- Introduction to Spark SQL 159
 - Introduction to Hive 160
 - Spark SQL Architecture 164

Getting Started with DataFrames	166
Using DataFrames	177
Caching, Persisting, and Repartitioning DataFrames	185
Saving DataFrame Output	186
Accessing Spark SQL	189
Exercise: Using Spark SQL	192
Using Spark with NoSQL Systems	193
Introduction to NoSQL	194
Using Spark with HBase	195
Exercise: Using Spark with HBase	198
Using Spark with Cassandra	200
Using Spark with DynamoDB	202
Other NoSQL Platforms	204
Summary	204
7 Stream Processing and Messaging Using Spark	207
Introducing Spark Streaming	207
Spark Streaming Architecture	208
Introduction to DStreams	209
Exercise: Getting Started with Spark Streaming	216
State Operations	217
Sliding Window Operations	219
Structured Streaming	221
Structured Streaming Data Sources	222
Structured Streaming Data Sinks	223
Output Modes	224
Structured Streaming Operations	225
Using Spark with Messaging Platforms	226
Apache Kafka	227
Exercise: Using Spark with Kafka	232
Amazon Kinesis	235
Summary	238
8 Introduction to Data Science and Machine Learning Using Spark	241
Spark and R	241
Introduction to R	242
Using Spark with R	248
Exercise: Using RStudio with SparkR	255

Machine Learning with Spark	257
Machine Learning Primer	257
Machine Learning Using Spark MLlib	260
Exercise: Implementing a Recommender Using Spark MLlib	265
Machine Learning Using Spark ML	269
Using Notebooks with Spark	273
Using Jupyter (IPython) Notebooks with Spark	273
Using Apache Zeppelin Notebooks with Spark	276
Summary	277

Spark Foundations

- 1 Introducing Big Data, Hadoop, and Spark 3
- 2 Deploying Spark 25
- 3 Understanding the Spark Cluster Architecture 43
- 4 Learning Spark Programming Basics 57