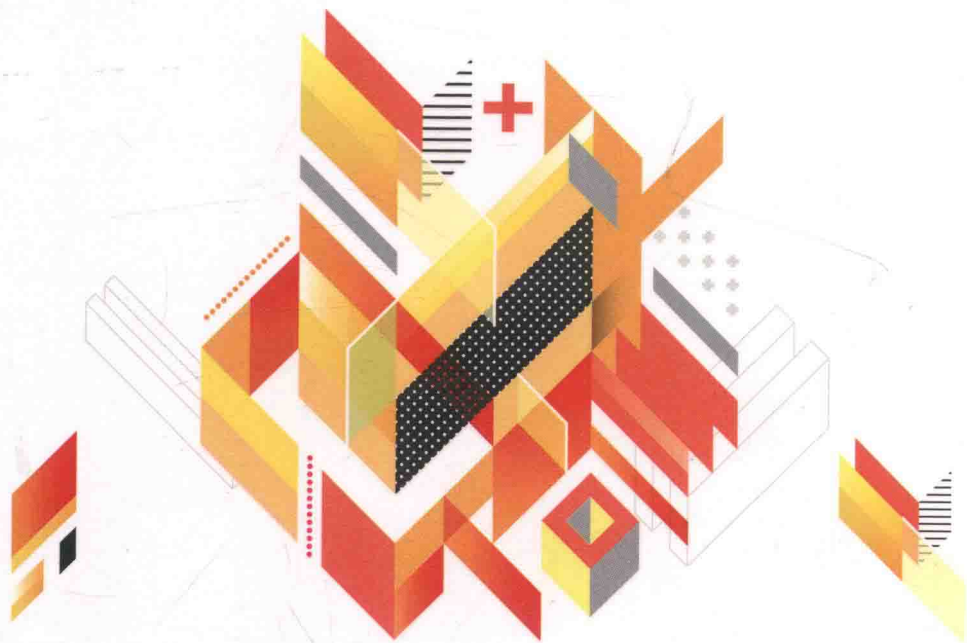




图解Spark

核心技术与案例实战

郭景瞻 编著

本书根据Spark最新版本，以源码为基础，结合实际案例，深入分析Spark作业调度、容错执行、存储管理以及运行架构等核心原理，同时对Spark生态圈BDAS组件做了详细的介绍。



中国工信出版集团



电子工业出版社
PUBLISHING HOUSE OF ELECTRONICS INDUSTRY
http://www.phei.com.cn

图解Spark

核心技术与案例实战

郭景瞻 编著

電子工業出版社

Publishing House of Electronics Industry

北京•BEIJING

内 容 简 介

本书以 Spark 2.0 版本为基础进行编写,全面介绍了 Spark 核心及其生态圈组件技术。主要内容包括 Spark 生态圈、实战环境搭建、编程模型和内部重要模块的分析,重点介绍了消息通信框架、作业调度、容错执行、监控管理、存储管理以及运行框架,同时还介绍了 Spark 生态圈相关组件,包括 Spark SQL 的即席查询、Spark Streaming 的实时流处理应用、MLbase/MLlib 的机器学习、GraphX 的图处理、SparkR 的数学计算和 Alluxio 的分布式内存文件系统等。

本书从 Spark 核心技术进行深入分析,重要章节会结合源代码解读其实现原理,围绕着技术原理介绍了相关典型实例,读者通过这些实例可以更加深入地理解 Spark 的运行机制。另外本书还应用了大量的图表进行说明,让读者能够更加直观地理解 Spark 相关原理。

本书不仅适合大数据、Spark 从业人员阅读,同时也适合大数据爱好者、架构师和软件开发人员阅读。通过本书,读者将能够很快地熟悉和掌握 Spark 大数据分析计算的利器,在生产中解决实际问题。

未经许可,不得以任何方式复制或抄袭本书之部分或全部内容。
版权所有,侵权必究。

图书在版编目(CIP)数据

图解 Spark: 核心技术与案例实战 / 郭景瞻编著. —北京: 电子工业出版社, 2017.1
ISBN 978-7-121-30236-7

I. ①图… II. ①郭… III. ①数据处理软件—图解 IV. ①TP274-64

中国版本图书馆 CIP 数据核字(2016)第 262015 号

责任编辑: 安 娜

印 刷: 北京京科印刷有限公司

装 订: 北京京科印刷有限公司

出版发行: 电子工业出版社

北京市海淀区万寿路 173 信箱 邮编: 100036

开 本: 787×980 1/16 印张: 30.25 字数: 570 千字

版 次: 2017 年 1 月第 1 版

印 次: 2017 年 1 月第 1 次印刷

印 数: 3000 册 定价: 99.00 元

凡所购买电子工业出版社图书有缺损问题,请向购买书店调换。若书店售缺,请与本社发行部联系,联系及邮购电话:(010) 88254888, 88258888。

质量投诉请发邮件至 zltz@phei.com.cn, 盗版侵权举报请发邮件至 dbqq@phei.com.cn。

本书咨询联系方式: 010-51260888-819, faq@phei.com.cn。

推荐序

移动互联网的兴起把我们带入了真正的大数据时代，各大互联网公司由于服务于海量用户，因此一般都储存了 EB 级的数据，而在一个典型的业务场景中，每次处理 TB 级的数据也是很常见的。同时由于竞争的加剧，互联网公司对业务的要求需要不断提高质量和降低响应速度，这就给大数据处理工具带来了非常大的挑战。目前的大数据生态圈以 Hadoop 为主，在数据存储、SQL 查询引擎、分布式计算、实时处理引擎和机器学习等方向先后诞生了一系列的开源项目，由于这些开源项目面向各自的领域，因而在设计的开发部署中，开发和运维的工程师就不得不同时面对不同的工具和环境，无形中大大增加了公司的成本和工程师的学习门槛。

Spark 就是在这样的背景下发展起来的，是目前为止唯一能够把交互式查询、实时处理、离线处理和机器学习无缝结合在一起的大数据产品。Spark 是基于内存的计算框架，其设计非常精巧，对于多次迭代的数据处理，例如机器学习和 SQL 查询，可以比 Hadoop 快很多倍，所以诞生不久即成为 Apache 的顶级项目，之后更是得到众多企业和开发者的拥护。一些大的互联网公司在过去几年虽然开发了自己的大数据处理框架，但最近几年也逐渐转向了 Spark。而对于那些对外提供云服务的厂商，Spark 更是成为一个标准配置环境，由此可见 Spark 的火热程度。

早期版本的 Spark 实现相当精简，然而随着版本的快速迭代和功能的不断增加，其实现已经变得相当复杂。由于分布式计算的复杂性，开发者和运维人员在实际使用过程中，经常会遇到集群甚至 Spark 本身相关的问题，在不了解工作原理的情况下很难快速定位和解决问题。本书作者在大数据领域深研究数年，自 Spark 诞生之日起就一直密切关注其发展，对其设计框架、运行机制和对外 API 都有较为深入的了解。作为京东大数据平台部门，我们需要解决业务部门在使用 Spark 过程中所遇到的各种复杂问题，在此过程中本书作者给予了我们莫大的帮助。

《图解 Spark：核心技术与案例实战》从生态系统讲起，先让读者对 Spark 生态圈有一个大概的了解。之后通过配置 Spark 开发环境，以及一个实际的例子告诉读者如何在 Spark 上快速开发。接下来作者详细介绍了 Spark 的编程模型和核心架构，这是本书的精华所在，也是真正了解 Spark 的必读内容，作者以平实的语言透彻地讲解了 RDD 含义、内部处理逻辑、任务执行的调度过程以及集群中 Driver 节点和 Worker 节点之间的交互等，至此读者可以清楚了解到 Spark

应用执行的完整过程。第 5 章详细介绍了存储原理及 Shuffle 过程，对于这些内容开发者往往不太在意，但对应用的整体性能有非常大的影响。第 6 章则以实际案例介绍了 Spark 的多种运行方式，读者从中可以了解到 Spark 是如何与自身资源管理框架、Yarn 集群或者 Mesos 集群进行交互的。

在对 Spark 应用有了整体的认识之后，作者又分别对 SQL 查询、流处理、机器学习、图计算和 SparkR 等核心子系统进行了深入解读，既介绍了各自的开发接口，又清楚地介绍了各个模块之间的关系。最后特别介绍了 Alluxio，其源于 Spark，可以为各种分布式系统提供抽象的文件存储服务。

大规模机器学习专家 京东大数据架构师 何云龙

前言

为什么要写这本书

在过去的十几年里，随着计算机的普遍应用和互联网的普及，使得数据呈现爆发式增长，在这个背景下，Doug Cutting 在谷歌的两篇论文（GFS 和 MapReduce）的启发下开发了 Nutch 项目。2006 年 Hadoop 脱离了 Nutch，成为 Apache 的顶级项目，带动了大数据发展的新十年。在此期间，大数据开源产品如雨后春笋般层出不穷，特别是 2009 年由加州大学伯克利分校 AMP 实验室开发的 Spark，它以内存迭代计算的高效和各组件所形成一站式解决平台成为这些产品的翘楚。

Spark 在 2013 年 6 月成为 Apache 孵化项目，8 个月后成为其顶级项目，并于 2014 年 5 月发布了 1.0 版本，在 2016 年 7 月正式发布了 2.0 版本。在这个过程中，Spark 社区不断壮大，成为了最为活跃的大数据社区之一。作为大数据处理的“利器”，Spark 在发展过程中不断地演进，因此各个版本存在较大的差异。市面上关于 Spark 的书已经不少，但是这些书所基于的 Spark 版本稍显陈旧，另外在介绍 Spark 的时候，未能把原理、代码和实例相结合，于是便有了本书，本书能够在剖析 Spark 原理的同时结合实际案例，从而让读者能够更加深入理解和掌握 Spark。

在本书中，首先对 Spark 的生态圈进行了介绍，讲述了 Spark 的发展历程，同时也介绍 Spark 实战环境的搭建；接下来从 Spark 的编程模型、作业执行、存储原理和运行架构等方面讲解了 Spark 内部核心原理；最后对 Spark 的各组件进行详细介绍，这些组件包括 Spark SQL 的即席查询、Spark Streaming 的实时流处理应用、MLbase/MLlib 的机器学习、GraphX 的图处理、SparkR 的数学计算和 Alluxio 的分布式内存文件系统等。

读者对象

（1）大数据爱好者

随着大数据时代的来临，无论是传统行业、IT 行业还是互联网等行业，都将涉及大数据技术，本书能够帮助这些行业的大数据爱好者了解 Spark 生态圈和发展演进趋势。通过本书，读

者不仅可以了解到 Spark 的特点和使用场景，而且如果希望继续深入学习 Spark 知识，那么本书也是很好的入门选择。

（2）Spark 开发人员

如果要进行 Spark 应用的开发，仅仅掌握 Spark 基本使用方法是不够的，还需深入了解 Spark 的设计原理、架构和运行机制。本书深入浅出地讲解了 Spark 的编程模型、作业运行机制、存储原理和运行架构等内容，通过对这些内容的学习，相信读者可以编写出更加高效的应用程序。

（3）Spark 运维人员

作为一名 Spark 运维人员，适当了解 Spark 的设计原理、架构和运行机制对于运维工作十分有帮助。通过对本书的学习，不仅能够更快地定位并排除故障，而且还能对 Spark 运行进行调优，让 Spark 运行得更加稳定和快速。

（4）数据科学家和算法研究

随着大数据技术的发展，实时流计算、机器学习、图计算等领域成为较热门的研究方向，而 Spark 有着较为成熟的生态圈，能够一站式解决类似场景的问题。这些研究人员可以通过本书加深对 Spark 的原理和应用场景的理解，从而能够更好地利用 Spark 各个组件进行数据计算和算法实现。

内容速览

本书分为三个部分，共计 12 章。

第一部分为基础篇（第 1~2 章），介绍了 Spark 诞生的背景、演进历程，以及 Spark 生态圈的组成，并详细介绍了如何搭建 Spark 实战环境。通过该环境不仅可以阅读 Spark 源代码，而且可以开发 Spark 应用程序。

第二部分为核心篇（第 3~6 章），讲解了 Spark 的编程模型、核心原理、存储原理和运行架构，在核心原理中对 Spark 通信机制、作业执行原理、调度算法、容错和监控管理等进行了深入分析，在分析原理和代码的同时结合实例进行演示。

第三部分为组件篇（第 7~12 章），介绍了 Spark 的各个组件，包括 Spark SQL 的即席查询、Spark Streaming 的实时流处理应用、MLbase/MLlib 的机器学习、GraphX 的图处理、SparkR 的数学计算和 Alluxio 的分布式内存文件系统等。

另外本书后面还包括 5 个附录：附录 A 为编译安装 Hadoop，附录 B 为安装 MySQL 数据库，附录 C 为编译安装 Hive，附录 D 为安装 ZooKeeper，附录 E 为安装 Kafka。由于本书篇幅有限，因此这些内容可到我的博客（<http://www.cnblogs.com/shishanyuan>）或博文视点网站（www.broadview.com.cn/30236）下载。

勘误和支持

由于笔者水平有限，加之编写时间跨度较长，同时 Spark 演进较快，因此在编写本书的过程中，难免会出现错误或者不准确的地方，恳请读者批评指正。如果本书存有错误，或者您有 Spark 的内容需要探讨，可以发送邮件到 jan98341@qq.com 与我联系，期待能够得到大家的反馈。

致谢

感谢中油瑞飞公司，让我接触到大数据的世界，并在工作的过程中深入了解 Spark。感谢吴建平、于鹏、李新宅、祝军、张文逵、马君博士、卢文君等领导同事，在本书编写过程中提供无私的帮助和宝贵的建议。

感谢京东商城的付彩宝、沈晓凯对我的工作和该书的支持，感谢付彩宝在繁忙的工作之余为本书写推荐，感谢京东数据挖掘架构师何云龙为本书作序，感谢大数据平台部的周龙波对该书提出了宝贵意见。

感谢 EMC 常雷博士为本书审稿并写推荐。

感谢 Alluxio 的 CEO 李浩源博士对本书的支持，感谢范斌在非常忙的工作中，抽出时间给 Alluxio 章节进行了审稿并提供了很好的建议。

非常感谢我的家人对我的理解和支持，特别是在写书过程中老婆又为我们家添了一位猴宝宝，让我拥有一对健康可爱的儿女，这些都给了我莫大的动力，让我的努力更加有意义。

谨以此书先给我亲爱的家人，你们是我努力的源泉。

郭景瞻

2016 年 11 月

目 录

第一篇 基础篇

第 1 章 Spark 及其生态圈概述	1
1.1 Spark 简介	1
1.1.1 什么是 Spark.....	1
1.1.2 Spark 与 MapReduce 比较.....	3
1.1.3 Spark 的演进路线图.....	4
1.2 Spark 生态系统	5
1.2.1 Spark Core.....	6
1.2.2 Spark Streaming	7
1.2.3 Spark SQL	9
1.2.4 BlinkDB	11
1.2.5 MLBase/MLlib.....	12
1.2.6 GraphX.....	12
1.2.7 SparkR.....	13
1.2.8 Alluxio.....	14
1.3 小结	15
第 2 章 搭建 Spark 实战环境	16
2.1 基础环境搭建	16
2.1.1 搭建集群样板机	17
2.1.2 配置集群环境	22
2.2 编译 Spark 源代码	25
2.2.1 配置 Spark 编译环境.....	26
2.2.2 使用 Maven 编译 Spark.....	27
2.2.3 使用 SBT 编译 Spark	29
2.2.4 生成 Spark 部署包.....	30
2.3 搭建 Spark 运行集群	31

2.3.1	修改配置文件	31
2.3.2	启动 Spark	33
2.3.3	验证启动	33
2.3.4	第一个实例	33
2.4	搭建 Spark 实战开发环境	35
2.4.1	CentOS 中部署 IDEA	36
2.4.2	使用 IDEA 开发程序	37
2.4.3	使用 IDEA 阅读源代码	42
2.5	小结	47

第二篇 核心篇

第3章	Spark 编程模型	48
3.1	RDD 概述	48
3.1.1	背景	48
3.1.2	RDD 简介	49
3.1.3	RDD 的类型	50
3.2	RDD 的实现	51
3.2.1	作业调度	51
3.2.2	解析器集成	52
3.2.3	内存管理	53
3.2.4	检查点支持	54
3.2.5	多用户管理	54
3.3	编程接口	55
3.3.1	RDD 分区 (Partitions)	55
3.3.2	RDD 首选位置 (PreferredLocations)	56
3.3.3	RDD 依赖关系 (Dependencies)	56
3.3.4	RDD 分区计算 (Iterator)	58
3.3.5	RDD 分区函数 (Partitioner)	58
3.4	创建操作	59
3.4.1	并行化集合创建操作	59
3.4.2	外部存储创建操作	61
3.5	转换操作	63
3.5.1	基础转换操作	63

3.5.2 键值转换操作	70
3.6 控制操作	77
3.7 行动操作	80
3.7.1 集合标量行动操作	80
3.7.2 存储行动操作	84
3.8 小结	87
第4章 Spark 核心原理	89
4.1 消息通信原理	90
4.1.1 Spark 消息通信架构	90
4.1.2 Spark 启动消息通信	91
4.1.3 Spark 运行时消息通信	94
4.2 作业执行原理	102
4.2.1 概述	102
4.2.2 提交作业	104
4.2.3 划分调度阶段	106
4.2.4 提交调度阶段	109
4.2.5 提交任务	112
4.2.6 执行任务	117
4.2.7 获取执行结果	119
4.3 调度算法	122
4.3.1 应用程序之间	122
4.3.2 作业及调度阶段之间	126
4.3.3 任务之间	130
4.4 容错及 HA	136
4.4.1 Executor 异常	136
4.4.2 Worker 异常	137
4.4.3 Master 异常	138
4.5 监控管理	139
4.5.1 UI 监控	139
4.5.2 Metrics	150
4.5.3 REST	152
4.6 实例演示	154
4.6.1 计算年降水实例	154

4.6.2 HA 配置实例	157
4.7 小结	160
第 5 章 Spark 存储原理	161
5.1 存储分析	161
5.1.1 整体架构	161
5.1.2 存储级别	167
5.1.3 RDD 存储调用	168
5.1.4 读数据过程	170
5.1.5 写数据过程	177
5.2 Shuffle 分析	186
5.2.1 Shuffle 简介	186
5.2.2 Shuffle 的写操作	186
5.2.3 Shuffle 的读操作	193
5.3 序列化和压缩	200
5.3.1 序列化	200
5.3.2 压缩	201
5.4 共享变量	202
5.4.1 广播变量	202
5.4.2 累加器	203
5.5 实例演示	204
5.6 小结	208
第 6 章 Spark 运行架构	209
6.1 运行架构总体介绍	209
6.1.1 总体介绍	209
6.1.2 重要类介绍	210
6.2 本地 (Local) 运行模式	211
6.2.1 运行模式介绍	211
6.2.2 实现原理	213
6.3 伪分布 (Local-Cluster) 运行模式	215
6.3.1 运行模式介绍	215
6.3.2 实现原理	216
6.4 独立 (Standalone) 运行模式	218

6.4.1	运行模式介绍	218
6.4.2	实现原理	219
6.5	YARN 运行模式	220
6.5.1	YARN 运行框架	220
6.5.2	YARN-Client 运行模式介绍	221
6.5.3	YARN-Client 运行模式实现原理	223
6.5.4	YARN-Cluster 运行模式介绍	227
6.5.5	YARN-Cluster 运行模式实现原理	229
6.5.6	YARN-Client 与 YARN-Cluster 对比	232
6.6	Mesos 运行模式	233
6.6.1	Mesos 介绍	233
6.6.2	粗粒度运行模式介绍	234
6.6.3	粗粒度实现原理	236
6.6.4	细粒度运行模式介绍	239
6.6.5	细粒度实现原理	240
6.6.6	Mesos 粗粒度和 Mesos 细粒度对比	243
6.7	实例演示	243
6.7.1	独立运行模式实例	243
6.7.2	YARN-Client 实例	247
6.7.3	YARN-Cluster 实例	250
6.8	小结	253

第三篇 组件篇

第 7 章	Spark SQL	255
7.1	Spark SQL 简介	255
7.1.1	Spark SQL 发展历史	255
7.1.2	DataFrame/Dataset 介绍	258
7.2	Spark SQL 运行原理	261
7.2.1	通用 SQL 执行原理	261
7.2.2	SparkSQL 运行架构	262
7.2.3	SQLContext 运行原理分析	265
7.2.4	HiveContext 介绍	276
7.3	使用 Hive-Console	278

7.3.1	编译 Hive-Console	278
7.3.2	查看执行计划	280
7.3.3	应用 Hive-Console	281
7.4	使用 SQLConsole	284
7.4.1	启动 HDFS 和 Spark Shell	284
7.4.2	与 RDD 交互操作	284
7.4.3	读取 JSON 格式数据	287
7.4.4	读取 Parquet 格式数据	288
7.4.5	缓存演示	289
7.4.6	DSL 演示	290
7.5	使用 Spark SQL CLI	290
7.5.1	配置并启动 Spark SQL CLI	291
7.5.2	实战 Spark SQL CLI	292
7.6	使用 Thrift Server	293
7.6.1	配置并启动 Thrift Server	293
7.6.2	基本操作	295
7.6.3	交易数据实例	296
7.6.4	使用 IDEA 开发实例	298
7.7	实例演示	299
7.7.1	销售数据分类实例	299
7.7.2	网店销售数据统计	303
7.8	小结	306
第 8 章	Spark Streaming	308
8.1	Spark Streaming 简介	308
8.1.1	术语定义	309
8.1.2	Spark Streaming 特点	312
8.2	Spark Streaming 编程模型	314
8.2.1	DStream 的输入源	314
8.2.2	DStream 的操作	315
8.3	Spark Streaming 运行架构	319
8.3.1	运行架构	319
8.3.2	消息通信	320
8.3.3	Receiver 分发	323

8.3.4	容错性	329
8.4	Spark Streaming 运行原理	331
8.4.1	启动流处理引擎	331
8.4.2	接收及存储流数据	334
8.4.3	数据处理	341
8.5	实例演示	346
8.5.1	流数据模拟器	346
8.5.2	销售数据统计实例	348
8.5.3	Spark Streaming+Kafka 实例	351
8.6	小结	356
第 9 章	Spark MLlib	358
9.1	Spark MLlib 简介	358
9.1.1	Spark MLlib 介绍	358
9.1.2	Spark MLlib 数据类型	360
9.1.3	Spark MLlib 基本统计方法	365
9.1.4	预言模型标记语言	369
9.2	线性模型	370
9.2.1	数学公式	370
9.2.2	线性回归	371
9.2.3	线性支持向量机	372
9.2.4	逻辑回归	373
9.2.5	线性最小二乘法、Lasso 和岭回归	373
9.2.6	流式线性回归	373
9.3	决策树	374
9.4	决策模型组合	375
9.4.1	随机森林	376
9.4.2	梯度提升决策树	377
9.5	朴素贝叶斯	377
9.6	协同过滤	378
9.7	聚类	380
9.7.1	K-means	380
9.7.2	高斯混合	382
9.7.3	快速迭代聚类	384

9.7.4	LDA	384
9.7.5	二分 K-means	385
9.7.6	流式 K-means	386
9.8	降维	386
9.8.1	奇异值分解降维	386
9.8.2	主成分分析降维	387
9.9	特征提取和变换	388
9.9.1	词频—逆文档频率	388
9.9.2	词向量化工具	389
9.9.3	标准化	390
9.9.4	范数化	390
9.10	频繁模式挖掘	391
9.10.1	频繁模式增长	391
9.10.2	关联规则挖掘	391
9.10.3	PrefixSpan	391
9.11	实例演示	392
9.11.1	K-means 聚类算法实例	392
9.11.2	手机短信分类实例	396
9.12	小结	401
第 10 章	Spark GraphX	402
10.1	GraphX 介绍	402
10.1.1	图计算	402
10.1.2	GraphX 介绍	403
10.1.3	发展历程	404
10.2	GraphX 实现分析	405
10.2.1	GraphX 图数据模型	406
10.2.2	GraphX 图数据存储	408
10.2.3	GraphX 图切分策略	410
10.2.4	GraphX 图操作	412
10.3	实例演示	418
10.3.1	图例演示	418
10.3.2	社区发现演示	425
10.4	小结	429

第 11 章 SparkR	430
11.1 概述	430
11.1.1 R 语言介绍	430
11.1.2 SparkR 介绍	431
11.2 SparkR 与 DataFrame	432
11.2.1 DataFrames 介绍	432
11.2.2 与 DataFrame 的相关操作	434
11.3 编译安装 SparkR	435
11.3.1 编译安装 R 语言	435
11.3.2 安装 SparkR 运行环境	437
11.3.3 安装 SparkR	438
11.3.4 启动并验证安装	439
11.4 实例演示	440
11.5 小结	444
第 12 章 Alluxio	445
12.1 Alluxio 简介	445
12.1.1 Alluxio 介绍	445
12.1.2 Alluxio 系统架构	446
12.1.3 HDFS 与 Alluxio	450
12.2 Alluxio 编译部署	451
12.2.1 编译 Alluxio	451
12.2.2 单机部署 Alluxio	453
12.2.3 集群模式部署 Alluxio	455
12.3 Alluxio 命令行使用	457
12.3.1 接口说明	457
12.3.2 接口操作示例	459
12.4 实例演示	462
12.4.1 启动环境	462
12.4.2 Alluxio 上运行 Spark	462
12.4.3 Alluxio 上运行 MapReduce	465
12.5 小结	466

本书附录部分请到博文视点网站下载 www.broadview.com.cn/30236。