

Hadoop专家

管理、调优与
Spark|YARN|HDFS安全

Expert Hadoop Administration:
Managing, Tuning, and Securing Spark,
YARN, and HDFS

[美] Sam R. Alapati 著

赵国贤 邓钊元 张京一 刘峰 等译

Hadoop专家

管理、调优与 Spark|YARN|HDFS安全

Expert Hadoop Administration:
Managing, Tuning, and Securing Spark, YARN, and HDFS

[美] Sam R. Alapati 著

赵国贤 邓钊元 张京一 刘峰 李小龙 译

电子工业出版社

Publishing House of Electronics Industry

北京•BEIJING

内 容 简 介

本书翻译自 Sam R. Alapati 的 *Expert Hadoop Administration*。Sam R. Alapati 是 Sabre 公司的首席 Hadoop 管理员，具有多年的 Hadoop 运维管理经验。他希望通过本书，为 Hadoop 集群开发与管理人员提供一些有益指导。

从事 Hadoop 的管理工作，首先要了解 Hadoop 的架构，只进行单纯的操作并不能被称为合格的管理员。基于此，本书在介绍 Hadoop 及其生态组件时，都会首先介绍其架构，以期读者能够从更高的层次认识管理工作。

本书首先介绍了 Hadoop 的整体架构及其部署与使用；然后着重介绍了两个重要的计算引擎 MapReduce 与 Spark；接着介绍了 Hadoop 的数据存储与安全、数据均衡等特性；最后则介绍了如何进行参数调优与故障排除。整个流程下来，读者能够建立起完整的关于 Hadoop 管理的体系架构。

本书为 Hadoop 管理员而编写，同时也适合 Hadoop 开发人员使用。

Authorized Translation from the English language edition, entitled *Expert Hadoop Administration: Managing, Tuning, and Securing Spark, YARN, and HDFS*, 0134597192 by Sam R. Alapati, published by Pearson Education, Inc., Copyright © 2017 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.

CHINESE SIMPLIFIED Language edition published by PUBLISHING HOUSE OF ELECTRONICS INDUSTRY, Copyright © 2019.

本书简体中文版专有出版权由 Pearson Education 授予电子工业出版社，未经许可，不得以任何方式复制或抄袭本书的任何部分。专有出版权受法律保护。

版权贸易合同登记号 图字：01-2017-3326

图书在版编目 (CIP) 数据

Hadoop专家：管理、调优与Spark|YARN|HDFS安全 / (美) 山姆·阿拉帕蒂 (Sam R. Alapati) 著；赵国贤等译. —北京：电子工业出版社，2019.3

书名原文：Expert Hadoop Administration: Managing, Tuning, and Securing Spark, YARN, and HDFS
ISBN 978-7-121-35669-8

I. ①H… II. ①山… ②赵… III. ①数据处理软件 IV. ①TP274

中国版本图书馆CIP数据核字 (2018) 第275589号

策划编辑：张春雨

责任编辑：牛 勇

印 刷：三河市良远印务有限公司

装 订：三河市良远印务有限公司

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：47.5 字数：1010千字

版 次：2019年3月第1版

印 次：2019年3月第1次印刷

定 价：168.00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 zltts@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。

译者序

承担本书翻译工作的主要人员是贝壳大数据架构相关团队，这个团队有着多年大数据的相关从业经验。本书很好地讲述了如何构建、优化、管理大数据智能计算平台。

在写下这篇译者序的时候，我更想把这个功劳归属于我们整个团队，我们团队负责公司大数据存储平台、计算平台、实时数据流平台的架构、性能优化、研发等，提供高效的大数据 olap 引擎，以及大数据工具链组件的研发，可为公司提供稳定、高效、开放的大数据基础组件与基础平台；专注于分布式计算、分布式存储以及大数据处理引擎的优化、架构等相关技术。整个翻译过程持续了一年多，在翻译中我们也感受到本书作者的专注与严谨。我们尽可能还原原作者的语义。作为一个做大数据相关工作的从业者，在翻译过程中自己也受益良多，也特别希望这本书能给大数据从业者赋能，为他们提供更好的助力。

在这个大数据与人工智能时代，Hadoop 作为一个基础平台，为多个公司提供了基础智能计算平台与大数据存储平台，本书正像一本手册一样，让我们能更好地利用好这个基础平台。

由于译者水平有限，本书难免有一些翻译错误，诚恳欢迎大家向我或者出版社反馈本书中的错误。

最后，我想要感谢参与本书翻译的刘峰、邓钊元、张京一、李小龙等同事，以及在翻译过程中帮助过我们的陈尔冬、杨菁伟、王涛、刘金国等领导与同事，还有很多其他帮助过我们的朋友，没有你们就不会有本书的出版。

序言

Apache Hadoop 2 和即将到来的 Hadoop 3 是在跨越 MapReduce 范式方面迈出的重要一步。其核心是新提出的 YARN 处理框架，该框架在 Hadoop 和 HDFS 之上提供了 API 和执行引擎，涵盖了之前的 MapReduce 模型。Hadoop 2 是对 Hadoop 1 的重大升级，因此在集群设置、管理和维护方面有较大改进。本书面向从事 Hadoop 2 生产集群的开发、操作和管理的人员。

Hadoop 2 和 3 的核心组件是 HDFS 和 YARN，除此之外，许多其他项目也被纳入 Hadoop 生产集群生态中。比如 Hive、Pig、Spark、Flume 及 Kafka 等经常与 Hadoop 核心组件配合使用，以提供更为完善的功能特性。本书涵盖了许多关于此类项目的介绍。

Sam Alapati 是 Sabre Holdings 公司的首席 Hadoop 管理员，过去 6 年一直从事 Hadoop 生产集群的维护管理工作。他是最有资格管理生产集群的人，并且他能把所有东西都整合到集群中。本书不仅仅是对 Hadoop 或 Spark 的简单介绍，而是提供了比较深入的体验内容，因此本书可以作为 Hadoop 管理员对 Hadoop 生产集群进行规范化、规模化、扩容以及提供安全性时的首选参考。

——Paul Dix，编辑

前言

Apache Hadoop 是一种流行的开源软件框架，其主要是在由普通商用硬件组成的集群中存储和处理海量数据。Hadoop 背后的主要思想是计算到数据，而非传统的数据到计算。良好的伸缩性是 Hadoop 的核心，Hadoop 之所以在当前的大数据领域备受欢迎，是因为普通商用服务器及开源性所带来的成本效益。

我从 2014 年秋季开始编写本书。Hadoop 2 在早前的几个月问世，新版本的 Hadoop 架构发生了许多有趣的变化。在此之前，有一本非常好的关于管理通用（不使用第三方供应商的工具）Hadoop 集群的书籍（Eric Sammer 的 *Hadoop Operations*）。但是，随着时间的推移，其在多个领域已经过时（该书发布于 2012 年）。Tom White 著的 *Hadoop: The Definitive Guide* 当然也是一部非常好的书籍，该书包含了一些 Hadoop 管理方面的有益探讨，但是相比于管理人员，Hadoop 的开发人员和架构师更适合阅读该书。于是我决定写一本书，该书应该成为关于集群管理、安全和优化方面的全面指南。

在本书的写作过程中，Spark 逐渐成为 Hadoop 最重要的处理框架之一。因此，我增加了 4 个章节来讨论 Spark 的架构、Spark 应用的本质及运行于 Hadoop 集群的 Spark 作业的管理和优化。

本书会直接通过 Hadoop 的配置文件来阐述 Hadoop 生态的管理、优化及安全。你可能想知道是否需要从底层开始学习 Hadoop 的管理。像许多管理 Hadoop 生态的人一样，我也使用第三方发行的 Hadoop，如 Cloudera 和 Hortonworks。当然，使用像 Cloudera Manager 或者 Apache Ambari 之类的工具来管理 Hadoop 集群是非常轻松的。但是，为了更好地管理 Hadoop 集群并最大限度地利用 Hadoop 集群，则需要了解这些管理工具管理集群背后的技术。只有从头开始构建一个集群并学习各种配置（如高可用性、高性能、安全性、加密等），才能够实现此目标。

Hadoop 具有大量的可配置属性。为了更好地利用 Hadoop 的强大性能，需要理解关键性能、安全性、高可用性以及其他相关的配置参数，并知道如何对其进行调优。为此，本书解释了所有与 Hadoop 管理相关的核心配置，并提供了大量的示例，以便你能够从

容地对集群进行配置，执行安全管理和优化工作。

Hadoop 是一个令人振奋的领域，其与“Hadoop 生态圈”下的软件进行交互。本书主要关注 Hadoop 核心本身，特别是 HDFS（Hadoop 分布式文件系统）及 YARN（Hadoop 处理框架）。本书也讨论了几个 Hadoop 生态圈的组件，如 Apache Sqoop、Apache Flume 和 Apache Spark 等，但重点是如何管理 Hadoop 架构本身。为此，我花费了大量时间讨论 HDFS 和 YARN 的架构体系。

谁适合阅读本书

本书主要是为 Hadoop 管理员而写。但是，并非全职的 Hadoop 管理员才能从本书中受益。如果你是一个大数据架构师、开发人员或者分析师，本书中的许多内容也适合你阅读。

本书的结构及内容

本书分为 5 个部分，共 21 章。

第 1 部分：Hadoop 架构与 Hadoop 集群介绍

- 第 1 章“Hadoop 与 Hadoop 环境介绍”从总体上介绍了 Hadoop 和大数据。由本章你可以了解到 Hadoop 与传统数据库的不同之处以及数据湖的概念。还可以了解到 Hadoop 与大数据和数据科学的契合之处。本章还介绍了 Hadoop 集群的概念，概述了 Hadoop 关键组件及 Hadoop 生态圈中的成员角色，如 ZooKeeper、Apache Sqoop、Apache Flume 和 Apache Kafka 等。

尽管 Hadoop1 现在已经成为历史，但它提供了一种方便的方法来追溯 Hadoop 到当前版本的演变历程，尤其是如何分离计算和调度，以及支持 MapReduce 之外的多个处理引擎。本书还讨论了 Hadoop1 和 Hadoop2 之间的主要区别，使你看清事情的本质，了解 Hadoop 的发展方向。

同时简要介绍了 MapReduce 和 Apache Spark 这两个 Hadoop 主要的计算框架，以及 Hive 和 Pig。本章还介绍了比较流行的 Hadoop 数据集成工具，如 Apache Flume 和 Apache Kafka。最后，总结了 Hadoop 管理员需要关注的领域，如资源分配、作业调度、性能调优以及安全性。

- 第 2 章“Hadoop 架构介绍”主要介绍了 Hadoop 的体系架构，并阐述了 HDFS 如何支持数据存储，以及提供数据处理功能的重要组件——YARN。
- 第 3 章“创建和配置一个简单的 Hadoop 集群”主要说明如何逐步地配置一个单节点的伪分布式集群。虽然无法使用单节点集群进行大规模的数据处理，但这

里主要是希望读者能够了解安装过程，而不是开始阶段就配置多个节点。在本章所学的所有内容，都涉及“真实”的多节点 Hadoop 集群的安装和配置。

- 第 4 章“规划和创建一个完全分布式集群”主要介绍了如何规划一个 Hadoop 集群以及如何对其进行调整。本章将一步一步地展示如何构建一个多节点 Hadoop 集群。

在学习了如何创建一个 Hadoop 集群后，需要了解如何修改 Hadoop 的默认配置。Hadoop 拥有大量的配置属性，可使用这些属性对存储、计算、资源分配和安全性等进行配置。

Hadoop 管理的一个关键点是了解如何使用大量配置参数来实现集群的配置、调整以及优化。本章将展示如何配置 Hadoop，以及如何对 Hadoop 服务、Web 接口及各种 Hadoop 端口进行配置。

第 II 部分：Hadoop 应用架构

- 第 5 章“在集群上运行一个应用——MapReduce 框架和 Hive、Pig”主要介绍了 MapReduce 的概念，其在许多年间都是 Hadoop 唯一可用的主要处理框架。在 Hadoop2 中，MapReduce 不再是唯一的计算框架，尽管其仍然在许多 Hadoop 环境中被重度使用。本章还介绍了著名的 WordCount 程序以及如何使用 MapReduce 执行它。

同时还介绍了两个在 Hadoop 中广泛使用的数据处理框架 Apache Hive 和 Apache Pig。

- 第 6 章“集群上的应用——Spark 框架介绍”主要介绍了 Apache Spark，其目标是接替 MapReduce 成为 Hadoop 主要的计算框架。本章重点介绍了 Spark 的安装与架构，以及如何将数据从各种数据源加载到 Spark。
- 第 7 章“运行 Spark 应用程序”介绍了 Spark 弹性分布式数据集（RDD）是什么，并展示了如何使用它。本章还介绍了如何通过 spark-submit 命令以交互式方式运行 Spark 作业。还将学习到配置 Spark 应用的各种方法以及如何监控 Spark 应用。

本章还介绍了用于流式数据处理的 Spark Streaming 以及用于结构化数据处理的 Spark SQL。

第 III 部分：管理和保护 Hadoop 数据和高可用性

- 第 8 章“NameNode 的作用和 HDFS 的工作原理”深入探讨了 NameNode 与 DataNode 之间的交互，以及如何在集群中配置机架感知。

数据复制是 HDFS 的名片。在本章中你将会了解到 HDFS 如何组织其数据以及

致谢

一本书的写作通常是一项团队工作，作者只是团队中的一员。我要感谢所有在本书撰写期间提供巨大帮助的人，首先是 Addison-Wesley 的执行编辑 Debra Williams Cauley，他负责监督这本书的编写和制作工作。Debra 可能是我遇到过的工作最勤奋、最认真的编辑，她对这个项目的奉献精神和管理一切的紧迫感对我处理这个项目的的方式产生了巨大影响，尤其是对于项目后期部分。

我非常感谢 Chris Zahn 对本书的精心编辑，她在整个艰难的过程中表现出极大的善良与优雅。Chris 的鼓励和支持给了我极大的帮助，他娴熟的编辑能力和敏锐而又不失大局的眼光给本书带来了不可估量的价值。从 Chris 的编辑中，我学到了很多关于正确格式和管理的事情。Chris 把事情理清楚后，就不大可能再有任何主要的文体错误！

我很幸运，有 4 位优秀的审稿人审阅过本书，他们都是来自 Hortonworks 的有丰富经验的专业人士。大数据顾问 Anubhav Awasthi，审阅了所有章节，发现了一些错误，并提出了几个重要建议，帮助我完成了对本书的改进。系统架构师 Karthik Varakantham 修改了一些文体和技术问题，并提出了一些非常有用的建议。高级顾问 Kannappan Natarasan 审阅了前几章，概述了前几章的内容，并提出了改进本书的建议。平台架构师 Ron Lee 提出了一些很好的建议，特别是对第 15~21 章部分。基于丰富的 Hadoop 环境的事件经验，Ron 针对本书给出了详细的评论，帮助我对本书进行了改进。

早些时候，Marina Stephens 和 John Guthrie 对本书的几个章节进行了详细的评审。依据他们的意见，我修改了行文风格和技术内容。感谢 Marina 和 John，感谢你们的辛苦工作，感谢你们发现许多错误，感谢你们提出的宝贵建议。

我是 Sabre 的一名大数据管理员，在这里我有幸与几位优秀的团队成员和管理人员一起工作。首先，我要感谢我的经理 Zeelani Shaik，感谢他在工作中一贯的礼貌、善良、理解和鼓励。Amjad Saeed 把我引荐到 Sabre，他的愉快、善良和优雅一直是我最大的快乐源泉。Zul Sidi，EDA 集团的副总裁，是整个团队的杰出领导者，他以自己的工作方式激励着我们，为我们树立了榜样。Zul 对我们做事方式方面的建议和开放态度，以及

他不断的鼓励和支持，帮助我和 EDA 团队的其他成员在他短暂的任职期间为我们的客户实现了许多重要的目标。

我也要感谢 Sujoe Bose 的善意和帮助，以及同事 Senthil Selvaraj 的帮助。我从 Sadu Hegde 那里学到了很多关于 Hadoop 的知识，我感谢他帮助我在 Sabre 起步。Mallik Dontula 在 Hadoop 和大数据的所有方面都是能手，我从他那里获得了宝贵的帮助和建议。Chris Morris 和 Larry Pritchett 不仅是好朋友，也是真正伟大的专业人士，我从与他们的共事中收益很多。我还要感谢 Aaron Patenaude 的慷慨，感谢他对我的一切帮助。我要感谢我的朋友 Winfield Geng 和 Bob Newman，他们在我需要的时候给予了我无私的帮助，他们是认真负责的，时刻提醒着我，感谢他们的建议和帮助。我要感谢 Mohammed Hossain 和 Andrew Ahmad 的友谊和帮助。我要感谢我的朋友和同事 Vinay Shetty 的大力支持。他不仅工作出色，而且他也很乐于助人。在过去的两年中，我非常感谢 Linda Phipps，感谢她为我所做的一切，她总是那么高兴！Lance Tripp 是鼓励我去寻找我现在在 Sabre 的职位的人。既然我热爱我的工作，我一定要感谢 Lance 聪明的招聘！

写一本书通常意味着你会从家里消失，尽管你的身体在家里。我感谢我的妻子 Valerie 的爱和善良，感谢我的孩子 Nina、Nicholas 和 Shannon 的爱，他们一直支持我的写作和研究工作。也谢谢 Dale、Shawn 和 Keith，他们一直和我很亲近。我也很感谢 Dixon 家的亲情和善良。

Stephanie Dixon、Clarence 和 Elaine，他们一直支持我所做的一切。我的兄弟 Hari 和 Bujji，我的嫂子 Aruna 和 Vanaja，我的侄女 Aparna 和 Soumya，我的侄子 Teja 对我来说意味着一切，所以感谢你们所有人！特别感谢我的侄子 Ashwin 和他的妻子 SreeVidya，在我在 San Jose 参加 Hadoop 会议期间，他们的盛情款待使我油然而生了几个关键想法。

关于作者

Sam R. Alapati 是 Sabre 的首席 Hadoop 管理员，公司总部位于得克萨斯州的南湖，他每天都要管理多个 Hadoop 集群。作为 Sabre 企业数据分析（EDA）部门所有 Hadoop 管理相关工作的负责人，Sam 管理并优化了与 Hadoop 相关的多个关键数据科学和数据分析工作的流程。Sam 还是一名 Oracle 数据库管理专家，具有丰富的关系型数据库和 SQL 的相关知识，因而他能成功地完成 Hadoop 相关的项目。Sam 在数据库和中间件领域取得了多项成就，包括在过去 14 年出版了 18 本受欢迎的书籍，主要是关于 Oracle 数据库管理和 Oracle Weblogic Server 方面的。Sam 也是《现代 Linux 管理》（O'Reilly, 2017）一书的作者。Sam 多年来在配置、体系结构和管理 Hadoop 性能方面的从业经历使他认识到，许多 Hadoop 管理员和开发人员都希望有一个方便的指南，比如本书，以便在创建、管理、保护和优化 Hadoop 基础架构时参考。

目录

第 I 部分 Hadoop架构与Hadoop集群介绍

第1章 Hadoop与Hadoop环境介绍.....	3
Hadoop简介.....	4
Hadoop 的特性.....	5
Hadoop 与大数据.....	5
Hadoop 的典型应用场景.....	6
传统数据库系统.....	7
数据湖.....	9
大数据、数据科学和 Hadoop.....	10
Hadoop集群与集群计算.....	11
集群计算.....	11
Hadoop 集群.....	12
Hadoop组件和Hadoop生态.....	14
Hadoop管理员需要做些什么.....	16
Hadoop 管理——新的范式.....	17
关于 Hadoop 管理你需要知道的.....	18
Hadoop 管理员的工具集.....	19
Hadoop 1和Hadoop 2的关键区别.....	19
架构区别.....	20
高可用性.....	20
多计算引擎.....	21

分离处理和调度.....	21
Hadoop 1 和 Hadoop 2 中的资源分配	22
分布式数据处理: MapReduce和Spark、Hive、Pig.....	22
MapReduce	22
Apache Spark	23
Apache Hive.....	24
Apache Pig	24
数据整合: Apache Sqoop、Apache Flume和Apache Kafka.....	25
Hadoop管理中的关键领域.....	26
集群存储管理.....	26
集群资源分配.....	26
作业调度.....	27
Hadoop 数据安全.....	27
总结.....	28
 第2章 Hadoop架构介绍.....	 31
Hadoop与分布式计算.....	31
Hadoop 架构.....	32
Hadoop 集群.....	33
主节点和工作节点.....	33
Hadoop 服务.....	34
数据存储——Hadoop分布式文件系统.....	35
HDFS 特性	35
HDFS 架构	36
HDFS 文件系统	38
NameNode 操作	41
利用YARN (Hadoop操作系统) 进行数据处理.....	45
YARN 的架构.....	46
ApplicationMaster 如何与 ResourceManager 协作进行资源分配.....	51
总结.....	54
 第3章 创建和配置一个简单的Hadoop集群.....	 55
Hadoop发行版本和安装类型	56

Hadoop 发行版本	56
Hadoop 安装类型	57
设置一个伪分布式Hadoop集群	58
满足操作系统的要求	58
修改内核参数	59
设置 SSH	64
Java 需求	65
安装 Hadoop	66
创建必要的 Hadoop 用户	66
创建必要的目录	67
Hadoop初始配置	67
环境变量配置文件	69
只读默认配置文件	70
site 专用配置文件	70
其他 Hadoop 相关的配置文件	71
配置文件的优先级	72
可变扩展和配置参数	74
配置 Hadoop 守护进程环境变量	74
配置 Hadoop 的核心属性 (使用 core-site.xml 文件)	76
配置 MapReduce (使用 mapred-site.xml 文件)	78
配置 YARN (使用 yarn-site.xml 文件)	79
配置 HDFS (使用 hdfs-site.xml 文件)	80
操作新的Hadoop集群	82
格式化分布式文件系统	82
设置环境变量	82
启动 HDFS 和 YARN 服务	83
验证服务启动	85
关闭服务	85
总结	86
 第4章 规划和创建一个完全分布式集群	 87
规划Hadoop集群	88
集群规划注意事项	88

安排服务器.....	90
节点选择的标准.....	90
从单机架到多机架.....	91
调整 Hadoop 集群.....	91
CPU、内存和存储选择的一般性原则.....	92
主节点的特殊要求.....	95
关于服务器大小的几点建议.....	96
集群增长.....	97
大型集群指南.....	97
创建一个多节点集群.....	98
如何设置测试集群.....	98
修改Hadoop的配置.....	102
更改 HDFS 的配置 (hdfs-site.xml 文件).....	102
更改 YARN 的配置.....	105
修改 MapReduce 的配置.....	109
启动集群.....	110
使用脚本启动和关闭集群.....	112
快速检查新集群的文件系统.....	113
配置Hadoop服务、Web界面和端口.....	114
服务配置和 Web 界面.....	115
设置 Hadoop 服务的端口.....	117
Hadoop 客户端.....	120
总结.....	122

第 II 部分 Hadoop应用架构

第5章 在集群上运行一个应用——MapReduce框架和Hive、Pig.....	125
MapReduce框架.....	125
MapReduce 模型.....	126
MapReduce 怎样工作.....	127
MapReduce 作业处理.....	129
一个简单的 MapReduce 程序.....	130
通过运行 WordCount 程序理解 Hadoop 作业的处理过程.....	132

MapReduce 输入 / 输出目录	133
Hadoop 如何展示作业细节	133
Hadoop Streaming.....	135
Apache Hive.....	137
Hive 数据组织	138
使用 Hive 表	138
将数据导入 Hive	138
使用 Hive 查询	139
Apache Pig.....	139
Pig 执行模型	140
一个简单的 Pig 示例	140
总结.....	141
 第6章 集群上的应用——Spark框架介绍	 143
Spark是什么	144
为什么使用 Spark	145
速度.....	145
易用性.....	147
通用框架.....	148
Spark 和 Hadoop.....	148
Spark技术栈	149
安装Spark	151
Spark 示例	152
Spark 的主要文件和目录	153
编译 Spark 二进制文件	153
减少 Spark 日志	153
Spark运行模式	154
本地模式.....	154
集群模式.....	154
集群管理器.....	154
独立集群管理器.....	155
基于 Apache Mesos 的 Spark.....	157
基于 YARN 的 Spark	158

YARN 和 Spark 如何协同合作	159
设置基于 Hadoop 集群的 Spark	159
Spark和数据获取	159
从 Linux 文件系统加载数据	160
从 HDFS 加载数据	160
从关系型数据库获取数据	161
总结	162
第7章 运行Spark应用程序	163
Spark编程模型	163
Spark 编程和 RDD	164
Spark 编程	166
Spark应用程序	167
RDD 基础	168
创建 RDD	168
RDD 操作	171
RDD 持久化	173
Spark应用的结构	174
Spark 术语	174
Spark 应用程序的组件	174
交互式运行Spark应用程序	175
Spark shell 和 Spark 应用程序	176
Spark shell	176
使用 Spark shell	176
Spark 集群执行概述	179
创建和提交Spark应用	180
构建 Spark 应用	180
在独立的 Spark 集群上运行应用	180
使用 spark-submit 执行应用	181
在 Mesos 上运行 Spark 应用	183
在 Hadoop YARN 集群上运行 Spark 应用	183
使用 JDBC/ODBC 服务	186
配置Spark应用	187