



Core Hadoop

Hadoop核心技术

翟周伟◎著



机械工业出版社
China Machine Press

Core Hadoop

Hadoop核心技术

翟周伟◎著



机械工业出版社
China Machine Press

图书在版编目 (CIP) 数据

Hadoop 核心技术 / 翟周伟著. —北京: 机械工业出版社, 2015.3
(大数据技术丛书)

ISBN 978-7-111-49468-3

I. H… II. 翟… III. 数据处理软件 IV. TP274

中国版本图书馆 CIP 数据核字 (2015) 第 044068 号

Hadoop 核心技术

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: 姜 影

责任校对: 董纪丽

印 刷: 三河市宏图印务有限公司

版 次: 2015 年 3 月第 1 版第 1 次印刷

开 本: 186mm×240mm 1/16

印 张: 25.25

书 号: ISBN 978-7-111-49468-3

定 价: 69.00 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: hzjg@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

Preface 前言

为什么要写这本书

如今是一个数据爆炸的时代，个人的图片、视频、文档等数据很容易就可达到数百 GB 的规模，而企业的数据规模增长得更快，以互联网搜索引擎公司为例，往往需要存储爬虫获取的所有站点的原始网页数据，同时还需要记录数以亿计网民的搜索点击行为及日志等重要数据，这些数据的规模可以达到 PB，甚至是 EB 级别。为了解决这些数据的存储和相关计算问题就必须构建一个强大且稳定的分布式集群系统来作为搜索引擎的基础架构支撑平台，但是对大多数的互联网公司而言，研发一个这样的高效能系统往往要支付高昂的费用。然而值得庆幸的是 Google 在 2004 年公布了关于其基础架构的两篇核心论文 GFS (The Google File System) 和 MapReduce (MapReduce: Simplified Data Processing on Large Clusters)，正是这两篇论文奠定了 Hadoop 的理论和实践基础，而后顶级工程师 Doug Cutting 将 Google 的 GFS 分布式文件系统和 MapReduce 并行计算模型实现且命名为 Hadoop，并将 Hadoop 开放源代码贡献给开源世界。经过多年的发展，如今已经形成了以 Hadoop 为核心的大数据生态系统，开创了通用海量数据处理基础架构平台的先河。

Hadoop 是一个非常优秀的分布式计算系统，利用通用的硬件就可以构建一个强大、稳定、简单，并且高效的分布式集群计算系统，完全可以满足互联网公司基础架构平台的需求，付出相对低廉的代价就可以轻松处理超大规模的数据。如今 Hadoop 已经被国内外各大公司广泛使用，在国内以百度、腾讯、阿里等为代表的著名互联网公司也都利用 Hadoop 构建底层的大数据基础架构平台，而移动、联通、电信等电信企业，以及电力、银行等国内传统企业也在大规模使用 Hadoop。因此学习并掌握 Hadoop 技术已经是进入这些企业的基本技能之一。相应地，学习和使用 Hadoop 技术的爱好者和开发者也越来越多。本书正是在这样的背景下创作出来的，希望可以帮助更多的人学习并掌握 Hadoop 技术，从而推动 Hadoop 技术在中国的

推广，进而推动中国信息产业的发展。

Hadoop 如今已经发行了多个版本，并且产生了各种以 Hadoop 为基础的应用系统，但是就核心技术而言 Hadoop 主要有两个重要版本：第一个就是以 MapReduce 模型为核心技术的版本，目前开源稳定版本为 Hadoop-1.X（包括 Hadoop-0.20.X，Hadoop-0.21.X，Hadoop-0.22.X 等）；第二个就是以 YARN 计算框架为核心技术的新版本，目前最新版本为 Hadoop-2.X（包括 Hadoop-0.23.X 等）。目前以 MapReduce 模型为核心技术的 Hadoop 版本已经在工业界的商业系统运行多年，不论是在稳定性还是高效性方面都已经得到了认可，而 YARN 计算框架虽然有更好的计算模式，但是在稳定性方面还没有真正被大规模的商业系统所证实，更多的是作为实验系统或者新特性被调研使用。因此将本书以 MapReduce 模型为核心技术的版本 Hadoop-1.X 作为研究和讲解对象，分为三大部分由浅入深进行讲解，包括基础篇、高级篇和实战篇。为了使读者更好地理解并掌握 Hadoop 核心技术，本书对核心实现机制都结合了源代码进行深入分析，并在实战篇中结合实际应用讲解了如何使用 Hadoop。

读者对象

（1）Hadoop 技术学习者和爱好者

Hadoop 作为一个分布式计算框架已经成为了互联网公司基础架构系统的标配，也吸引了越来越多的开发者和爱好者的关注。本书在写作上由纲及目、由浅入深，一步一步带领读者从 Hadoop 技术的入门开始，逐步到核心原理的实现机制以及 Hadoop 的实战应用，因此本书可以帮助 Hadoop 技术学习者和爱好者快速入门，对 Hadoop 的学习起到一个比较全面、深入而又不乏实战性的导向作用。

（2）Hadoop MapReduce 并行应用开发者

Hadoop 可以说是一个复杂而又简单的系统，说其复杂是因为 Hadoop 同时实现了分布式文件系统和并行计算框架，在实现机制上考虑了各种容错情况，因此不可谓不复杂；说其简单是因为 Hadoop 向并行应用开发者提供了非常简单的编程接口——Map 和 Reduce 函数操作，因此，应用开发者只需要会使用这两个编程接口就可以很方便地编写处理大规模数据的并行应用程序。但是如果开发者想要编写高质量的并行程序，而只懂得基本的 MapReduce 编程方法肯定是不够的，这相当于知其然而不知其所以然，因此开发人员还需要知其所以然——必须理解 Hadoop 的核心实现机制，需要对 Hadoop 框架有一个全面而又深入的认识和理解，只有这样 Hadoop 应用开发者才能编写出高效而又简洁的 MapReduce 应用程序。

（3）Hadoop 集群运维人员

Hadoop 运维工程师可以说是复合型人才，因为 Hadoop 的运维不仅需要熟悉 Linux 系统的运维方法，还需要非常熟悉 Hadoop 的各种管理命令，不仅如此，还要了解 Hadoop 集群搭

建、权限管理、作业调度管理与维护等。本书在 Hadoop 命令系统和集群搭建方面专门编排了章节进行讲解,因此可以作为 Hadoop 运维工程师的工具书。Hadoop 运维工程师的主要职责虽然是通过集群运维保证 Hadoop 的高可靠性和高可用性,但是仍然需要对 Hadoop 的核心实现机制,甚至是某些实现细节进行深入理解,只有这样才可以在集群出现异常时快速找出问题,起到锦上添花的作用,因此本书对 Hadoop 运维工程来说也是一本值得精读的参考书。

(4) 分布式系统的相关研发人员

Hadoop 可以说是分布式系统领域中的经典之作,分布式文件系统和分布式计算系统中的核心理论方法都在 Hadoop 中有很好的实现。因此,通过对 Hadoop 的学习不但可以理解分布式系统中的理论方法,而且还可以深入理解这些理论方法具体的实现机制。

(5) 大数据技术学习者和大数据工程师

如果从大数据的角度来讲,Hadoop 无疑是目前使用最为广泛的大数据平台。特别是在互联网领域,Hadoop 数据平台作为底层基础架构系统支撑着大多数的数据统计分析应用,因此对于大数据技术学习者和大数据工程师而言,学习 Hadoop 技术是进入大数据领域的一个很好的切入点。

如何阅读本书

本书分为三篇。

第一篇为基础篇(第1~6章),从认识 Hadoop 开始,讲解 Hadoop 的前世今生以及使用领域,然后正式介绍 Hadoop 的基本使用,帮助读者了解 Hadoop 的背景知识和简单使用方法,接着通过 HDFS 分布式文件系统和 MapReduce 并行计算模型从理论和实现机制的角度对 Hadoop 核心技术进行讲解,最后对 Hadoop 的命令系统进行了系统的介绍。对于初级和中级读者而言,第一篇的内容需要重点阅读和学习,这篇是 Hadoop 核心技术的基础,只有基础知识扎实后才能更好地掌握 Hadoop 的高级功能和精髓。

第二篇为高级篇(第7~9章),从原理与实现的角度对 Hadoop 的核心功能进行了深入的研究,涵盖 MapReduce 深度分析、Hadoop Streaming 和 Pipes 原理解析,以及 Hadoop 作业调度器系统的深入研究和讲解。本篇内容适合在阅读了基础篇的基础上或者已经对 Hadoop 的核心原理有了一定理解的基础上进行阅读。

第三篇为实战篇(第10~12章),从实战的角度进行讲解,首先讲述 Hadoop 集群搭建技术,然后对 Streaming 和 Pipes 编程进行了实战级的应用讲解,Streaming 编程接口是一个非常简单且高效的 MapReduce 编程方式,由于不限制编程语言,因此 Streaming 的使用比 Java 原生接口应用得还要广泛,由此可见,学习并掌握 Streaming 编程技术非常有助于软件工程师

的 Hadoop 应用技术的提高。第 12 章讲解了 Hadoop MapReduce 应用开发实战,从整体的并行应用开发角度进行讲解,对实际开发过程中的常用功能使用和常见问题解决都进行了介绍。这部分内容适合在实际工作中使用 Hadoop 开发应用的工程师阅读和学习。

勘误和支持

由于作者的水平有限且编写时间仓促,书中难免会出现一些错误或者不准确的地方,恳请读者批评指正。为此,特意创建一个在线支持与应急方案的二级网站:<http://book.zhaizhouwei.cn>。读者可以将书中的错误发布在 Bug 勘误表页面中,此外如果你遇到任何问题,也可以访问 Q&A 页面,我将尽量在线上为读者提供最满意的解答。如果你有更多的宝贵意见,也欢迎发送邮件至邮箱 zhaizhouwei@163.com,期待能够得到你们的真挚反馈。

致谢

感谢北京邮电大学吕玉琴教授、刘刚博士、侯斌博士、李威海博士等老师,是他们在读研的时候悉心指导我进入大数据以及统计挖掘领域,并使我能在 2009 年就开始学习并使用 Hadoop 技术,才使得本书今日的出版成为可能。特别感谢百度网页搜索部技术总监沈抖博士在百忙之中抽出时间对本书进行通读并指正,在多个章节提出了非常专业的指导建议。在此真心的感谢他们。

感谢机械工业出版社华章公司的杨福川老师和姜影老师,从本书的初期策划到最终的出版过程中始终支持我的写作,是他们的鼓励和帮助引导我能顺利完成全部书稿。

最后感谢我的爸爸、妈妈等家人,感谢他们给我的鼓励并时刻给予我信心和力量!感谢我的夫人欧阳茹给我生活上的悉心照顾以及在琐碎事情上的宽容大度。

谨以此书献给我最亲爱的家人、同事,以及众多热爱 Hadoop 的朋友们!

翟周伟

推荐阅读



推荐阅读



推荐阅读



VMware Virtual SAN权威指南

作者：(美) Cormac Hogan 等 ISBN: 978-7-111-48023-5 定价：59.00元

不论您是虚拟化新手，还是存储专家，这本书是有关VMware Virtual SAN最权威的解读，是实现软件定义存储最有效的指南。

——任道远，VMware中国研发中心总经理

VMware资深虚拟存储专家亲笔撰写，全球第一本全面、系统讲解Virtual SAN技术的权威著作，Amazon全5星评价。

从Virtual SAN的部署、安装、配置到虚拟机存储管理、架构细节和日常管理、维护等方面，深入探讨Virtual SAN的各项技术细节，并用多个实例详细讲解Virtual SAN群集的设计和实现。

本书专为管理员、咨询师和架构师所著，在书中Cormac Hogan和Duncan Epping既介绍了Virtual SAN如何实现基于对象的存储和策略平台，这些功能简化了虚拟机存储的放置，还介绍了Virtual SAN如何与vSphere协同工作，大幅提高系统弹性、存储横向扩展和QoS控制的能力。

VMware网络技术：原理与实践

作者：(美) Christopher Wahl 等 ISBN: 978-7-111-47987-1 定价：59.00元

资深虚拟化技术专家亲笔撰写，CCIE认证专家Ivan Pepelnjak作序鼎力推荐，Amazon广泛好评
既详细讲解物理网络的基础知识，又通过丰富实例深入探究虚拟交换机的功能和设计，
全面阐释虚拟网络环境构建的各种技术细节、方法及最佳实践

本书针对VMware专业人员，阐述了现代网络的核心概念，并介绍了如何在虚拟网络环境设计、配置和故障检修中应用这些概念。作者凭借其在虚拟化项目实施方面的丰富经验，从网络模型、常见网络层次的介绍开始，由浅入深地介绍了现代网络的基本概念，并自然地过渡到虚拟交换等虚拟化环境中与物理网络最为关联的部分，最后扩展到实际的设计用例，详细介绍了不同实用场景、不同的硬件配置下，虚拟化环境构建的考虑因素和具体实施方案。

Contents 目 录

前 言

基 础 篇

第 1 章 认识 Hadoop 2

1.1 缘于搜索的小象 2

1.1.1 Hadoop 的身世 2

1.1.2 Hadoop 简介 3

1.1.3 Hadoop 发展简史 6

1.2 大数据、Hadoop 和云计算 7

1.2.1 大数据 7

1.2.2 大数据、Hadoop 和云计算的 关系 8

1.3 设计思想与架构 9

1.3.1 数据存储与切分 9

1.3.2 MapReduce 模型 11

1.3.3 MPI 和 MapReduce 13

1.4 国外 Hadoop 的应用现状 13

1.5 国内 Hadoop 的应用现状 17

1.6 Hadoop 发行版 20

1.6.1 Apache Hadoop 20

1.6.2 Cloudera Hadoop 20

1.6.3 Hortonworks Hadoop 发行版 21

1.6.4 MapR Hadoop 发行版 22

1.6.5 IBM Hadoop 发行版 24

1.6.6 Intel Hadoop 发行版 24

1.6.7 华为 Hadoop 发行版 25

1.7 小结 26

第 2 章 Hadoop 使用之初体验 27

2.1 搭建测试环境 27

2.1.1 软件与准备 27

2.1.2 安装与配置 28

2.1.3 启动与停止 29

2.2 算法分析与设计 31

2.2.1 Map 设计 31

2.2.2 Reduce 设计 32

2.3 实现接口 32

2.3.1 Java API 实现 33

2.3.2 Streaming 接口实现 36

2.3.3 Pipes 接口实现 38

2.4 编译 40

2.4.1 基于 Java API 实现的编译 40

2.4.2 基于 Streaming 实现的编译	40	3.4.9 健壮性	59
2.4.3 基于 Pipes 实现的编译	41	3.4.10 负载均衡	60
2.5 提交作业	41	3.4.11 升级和回滚机制	62
2.5.1 基于 Java API 实现作业提交	41	3.5 HDFS 权限管理	64
2.5.2 基于 Streaming 实现作业提交	42	3.5.1 用户身份	64
2.5.3 基于 Pipes 实现作业提交	43	3.5.2 系统实现	65
2.6 小结	44	3.5.3 超级用户	65
		3.5.4 配置参数	65
第3章 Hadoop 存储系统	45	3.6 HDFS 配额管理	66
3.1 基本概念	46	3.7 HDFS 的缺点	67
3.1.1 NameNode	46	3.8 小结	68
3.1.2 DataNode	46		
3.1.3 客户端	47	第4章 HDFS 的使用	69
3.1.4 块	47	4.1 HDFS 环境准备	69
3.2 HDFS 的特性和目标	48	4.1.1 HDFS 安装配置	69
3.2.1 HDFS 的特性	48	4.1.2 HDFS 格式化与启动	70
3.2.2 HDFS 的目标	48	4.1.3 HDFS 运行检查	70
3.3 HDFS 架构	49	4.2 HDFS 命令的使用	71
3.3.1 Master/Slave 架构	49	4.2.1 fs shell	71
3.3.2 NameNode 和 Secondary NameNode 通信模型	51	4.2.2 archive	77
3.3.3 文件存取机制	52	4.2.3 distcp	78
3.4 HDFS 核心设计	54	4.2.4 fsck	81
3.4.1 Block 大小	54	4.3 HDFS Java API 的使用方法	82
3.4.2 数据复制	55	4.3.1 Java API 简介	82
3.4.3 数据副本存放策略	56	4.3.2 读文件	82
3.4.4 数据组织	57	4.3.3 写文件	86
3.4.5 空间回收	57	4.3.4 删除文件或目录	90
3.4.6 通信协议	58	4.4 C 接口 libhdfs	91
3.4.7 安全模式	58	4.4.1 libhdfs 介绍	91
3.4.8 机架感知	59	4.4.2 编译与部署	91
		4.4.3 libhdfs 接口介绍	92

4.4.4	libhdfs 使用举例	95
4.5	WebHDFS 接口	97
4.5.1	WebHDFS REST API 简介	97
4.5.2	WebHDFS 配置	98
4.5.3	WebHDFS 使用	98
4.5.4	WebHDFS 错误响应和查询参数	101
4.6	小结	103
第 5 章	MapReduce 计算框架	104
5.1	Hadoop MapReduce 简介	104
5.2	MapReduce 模型	105
5.2.1	MapReduce 编程模型	105
5.2.2	MapReduce 实现原理	106
5.3	计算流程与机制	108
5.3.1	作业提交和初始化	108
5.3.2	Mapper	110
5.3.3	Reducer	111
5.3.4	Reporter 和 OutputCollector	112
5.4	MapReduce 的输入 / 输出格式	113
5.4.1	输入格式	113
5.4.2	输出格式	118
5.5	核心问题	124
5.5.1	Map 和 Reduce 数量	124
5.5.2	作业配置	126
5.5.3	作业执行和环境	127
5.5.4	作业容错机制	129
5.5.5	作业调度	131
5.6	有用的 MapReduce 特性	132
5.6.1	计数器	132
5.6.2	DistributedCache	134

5.6.3	Tool	135
5.6.4	IsolationRunner	136
5.6.5	Profiling	136
5.6.6	MapReduce 调试	136
5.6.7	数据压缩	137
5.6.8	优化	138
5.7	小结	138

第 6 章 Hadoop 命令系统 139

6.1	Hadoop 命令系统的组成	139
6.2	用户命令	141
6.3	管理员命令	144
6.4	测试命令	148
6.5	应用命令	156
6.6	Hadoop 的 streaming 命令	163
6.6.1	streaming 命令	163
6.6.2	参数使用分析	164
6.7	Hadoop 的 pipes 命令	168
6.7.1	pipes 命令	168
6.7.2	参数使用分析	169
6.8	小结	170

高级篇

第 7 章 MapReduce 深度分析 172

7.1	MapReduce 总结构分析	172
7.1.1	数据流向分析	172
7.1.2	处理流程分析	174
7.2	MapTask 实现分析	176
7.2.1	总逻辑分析	176

7.2.2	Read 阶段	178
7.2.3	Map 阶段	178
7.2.4	Collector 和 Partitioner 阶段	180
7.2.5	Spill 阶段	181
7.2.6	Merge 阶段	185
7.3	ReduceTask 实现分析	186
7.3.1	总逻辑分析	186
7.3.2	Shuffle 阶段	187
7.3.3	Merge 阶段	189
7.3.4	Sort 阶段	190
7.3.5	Reduce 阶段	191
7.4	JobTracker 分析	192
7.4.1	JobTracker 服务分析	192
7.4.2	JobTracker 启动分析	193
7.4.3	JobTracker 核心子线程分析	195
7.5	TaskTracker 分析	201
7.5.1	TaskTracker 启动分析	201
7.5.2	TaskTracker 核心子线程分析	205
7.6	心跳机制实现分析	207
7.6.1	心跳检测分析	207
7.6.2	TaskTracker.transmitHeart- Beat()	207
7.6.3	JobTracker.heartbeat()	209
7.6.4	JobTracker.processHeartbeat()	212
7.7	作业创建分析	213
7.7.1	初始化分析	214
7.7.2	作业提交分析	215
7.8	作业执行分析	217
7.8.1	JobTracker 初始化	218
7.8.2	TaskTracker.startNewTask()	220
7.8.3	TaskTracker.localizeJob()	220

7.8.4	TaskRunner.run()	221
-------	------------------	-----

7.8.5	MapTask.run()	222
-------	---------------	-----

7.9	小结	223
-----	----	-----

第 8 章 Hadoop Streaming 和 Pipes

原理与实现

8.1	Streaming 原理浅析	224
8.2	Streaming 实现架构	226
8.3	Streaming 核心实现机制	227
8.3.1	主控框架实现	227
8.3.2	用户进程管理	228
8.3.3	框架和用户程序的交互	229
8.3.4	PipeMapper 和 PiperReducer	230
8.4	Pipes 原理浅析	231
8.5	Pipes 实现架构	233
8.6	Pipes 核心实现机制	234
8.6.1	主控类实现	234
8.6.2	用户进程管理	235
8.6.3	PipesMapRunner	235
8.6.4	PipesReducer	238
8.6.5	C++ 端 HadoopPipes	238
8.7	小结	239

第 9 章 Hadoop 作业调度系统

9.1	作业调度概述	241
9.1.1	相关概念	241
9.1.2	作业调度流程	242
9.1.3	集群资源组织与管理	243
9.1.4	队列控制和权限管理	244
9.1.5	插件式调度框架	245
9.2	FIFO 调度器	246

9.2.1	基本调度策略	246
9.2.2	FIFO 实现分析	247
9.2.3	FIFO 初始化与停止	248
9.2.4	作业监听控制	249
9.2.5	任务分配算法	250
9.2.6	配置与使用	251
9.3	公平调度器	254
9.3.1	产生背景	254
9.3.2	主要功能	255
9.3.3	基本调度策略	255
9.3.4	FairScheduler 实现分析	257
9.3.5	FairScheduler 启停分析	258
9.3.6	作业监听控制	260
9.3.7	资源池管理	260
9.3.8	作业更新策略	262
9.3.9	作业权重和资源量的计算	266
9.3.10	任务分配算法	267
9.3.11	FairScheduler 配置参数	268
9.3.12	使用与管理	270
9.4	容量调度器	272
9.4.1	产生背景	272
9.4.2	主要功能	272
9.4.3	基本调度策略	274
9.4.4	CapacityScheduler 实现分析	274
9.4.5	CapacityScheduler 启停分析	275
9.4.6	作业监听控制	277
9.4.7	作业初始化分析	277
9.4.8	任务分配算法	278
9.4.9	内存匹配机制	279
9.4.10	配置与使用	280
9.5	调度器对比分析	283
9.5.1	调度策略对比	283

9.5.2	队列和优先级	283
9.5.3	资源分配保证	283
9.5.4	作业限制	284
9.5.5	配置管理	284
9.5.6	扩展性支持	284
9.5.7	资源抢占和延迟调度	284
9.5.8	优缺点分析	285
9.6	其他调度器	285
9.6.1	HOD 调度器	285
9.6.2	LATE 调度器	286
9.7	小结	288

实战篇

第 10 章	Hadoop 集群搭建	290
10.1	Hadoop 版本的选择	290
10.2	集群基础硬件需求	291
10.2.1	内存	291
10.2.2	CPU	292
10.2.3	磁盘	292
10.2.4	网卡	293
10.2.5	网络拓扑	293
10.3	集群基础软件需求	294
10.3.1	操作系统	294
10.3.2	JVM 和 SSH	295
10.4	虚拟化需求	295
10.5	事前准备	296
10.5.1	创建安装用户	296
10.5.2	安装 Java	297
10.5.3	安装 SSH 并设置	297
10.5.4	防火墙端口设置	298

10.6	安装 Hadoop	298
10.6.1	安装 HDFS	299
10.6.2	安装 MapReduce	299
10.7	集群配置	300
10.7.1	配置管理	300
10.7.2	环境变量配置	301
10.7.3	核心参数配置	302
10.7.4	HDFS 参数配置	303
10.7.5	MapReduce 参数配置	306
10.7.6	masters 和 slaves 配置	313
10.7.7	客户端配置	313
10.8	启动和停止	314
10.8.1	启动 / 停止 HDFS	314
10.8.2	启动 / 停止 MapReduce	315
10.8.3	启动验证	315
10.9	集群基准测试	316
10.9.1	HDFS 基准测试	316
10.9.2	MapReduce 基准测试	317
10.9.3	综合性能测试	318
10.10	集群搭建实例	319
10.10.1	部署策略	319
10.10.2	软件和硬件环境	320
10.10.3	Hadoop 安装	321
10.10.4	配置 core-site.xml	321
10.10.5	配置 hdfs-site.xml	322
10.10.6	配置 mapred-site.xml	322
10.10.7	SecondaryNameNode 和 Slave	324
10.10.8	配置作业队列	324
10.10.9	配置第三方调度器	325
10.10.10	启动与验证	327

10.11	小结	327
-------	----------	-----

第 11 章 Hadoop Streaming 和 Pipes 编程实战

11.1	Streaming 基础编程	328
11.1.1	Streaming 编程入门	328
11.1.2	Map 和 Reduce 数目	331
11.1.3	队列、优先级及权限	332
11.1.4	分发文件和压缩包	333
11.1.5	压缩参数的使用	336
11.1.6	本地作业的调试	338
11.2	Streaming 高级应用	338
11.2.1	参数与环境变量传递	339
11.2.2	自定义分隔符	340
11.2.3	自定义 Partitioner	343
11.2.4	自定义计数器	347
11.2.5	处理二进制数据	347
11.2.6	使用聚合函数	351
11.3	Pipes 编程接口	352
11.3.1	TaskContext	352
11.3.2	Mapper	353
11.3.3	Reducer	354
11.3.4	Partitioner	354
11.3.5	RecordReader	355
11.3.6	RecordWriter	356
11.4	Pipes 编程应用	357
11.5	小结	359

第 12 章 Hadoop MapReduce 应用开发

12.1	开发环境准备	360
------	--------------	-----

12.2	Eclipse 集成环境开发	361	12.4.3	本地压缩库	375
12.2.1	构建 MapReduce Eclipse IDE	361	12.4.4	使用压缩	376
12.2.2	开发示例	363	12.5	排序应用	378
12.3	MapReduce Java API 编程	368	12.5.1	Hadoop 排序问题	378
12.3.1	Mapper 编程接口	369	12.5.2	二次排序	378
12.3.2	Reducer 编程接口	370	12.5.3	比较器和组合排序	380
12.3.3	驱动类编写	372	12.5.4	全局排序	381
12.3.4	编译运行	373	12.6	多路输出	382
12.4	压缩功能使用	374	12.7	常见问题与处理方法	384
12.4.1	Hadoop 数据压缩	374	12.7.1	常见的开发问题	384
12.4.2	压缩特征与性能	374	12.7.2	运行时错误问题	386
			12.8	小结	387