

PACKT  
PUBLISHING



简单、实用的代码示例，快速解决诸多Hadoop相关技术问题

# Hadoop 实战手册

Hadoop Real-World Solutions Cookbook

[美] Jonathan R. Owens   Jon Lentz   Brian Femiano   著  
傅杰   赵磊   卢学裕   译



人民邮电出版社  
POSTS & TELECOM PRESS



# Hadoop 实战手册

[美] Jonathan R. Owens   Jon Lentz   Brian Femiano   著  
傅杰   赵磊   卢学裕   译

人民邮电出版社  
北 京

## 图书在版编目 (C I P) 数据

Hadoop 实战手册 / (美) 欧文斯 (Owens, J. R.),  
(美) 伦茨 (Lentz, J.), (美) 费米亚诺 (Femiano, B.)  
著; 傅杰, 赵磊, 卢学裕译. — 北京: 人民邮电出版  
社, 2014. 3

书名原文: Hadoop real-world solutions cookbook  
ISBN 978-7-115-33795-5

I. ①H… II. ①欧… ②伦… ③费… ④傅… ⑤赵…  
⑥卢… III. ①数据处理软件—技术手册 IV.  
①TP274-62

中国版本图书馆CIP数据核字(2013)第277092号

## 版权声明

Copyright ©2013 Packt Publishing. First published in the English language under the title *Hadoop Real-World Solutions Book*.

All Rights Reserved.

本书由英国 Packt Publishing 公司授权人民邮电出版社出版。未经出版者书面许可, 对本书的任何部分不得以任何方式或任何手段复制和传播。

版权所有, 侵权必究。

◆ 著 [美] Jonathan R. Owens Jon Lentz Brian Femiano

译 傅杰 赵磊 卢学裕

责任编辑 杨海玲

责任印制 程彦红 杨林杰

◆ 人民邮电出版社出版发行 北京市丰台区成寿寺路11号

邮编 100164 电子邮件 315@ptpress.com.cn

网址 <http://www.ptpress.com.cn>

北京天宇星印刷厂印刷

◆ 开本: 800×1000 1/16

印张: 16.25

字数: 318 千字

2014 年 3 月第 1 版

印数: 1-3 500 册

2014 年 3 月北京第 1 次印刷

著作权合同登记号 图字: 01-2013-4468 号



定价: 59.00 元

读者服务热线: (010)81055410 印装质量热线: (010)81055316

反盗版热线: (010)81055315

# 内容提要

这是一本 Hadoop 实用手册，主要针对实际问题给出相应的解决方案。本书特色是以实践结合理论分析，手把手教读者如何操作，并且对每个操作都做详细的解释，对一些重要的知识点也做了必要的拓展。全书共包括 3 个部分，第一部分为基础篇，主要介绍 Hadoop 数据导入导出、HDFS 的概述、Pig 与 Hive 的使用、ETL 和简单的数据处理，还介绍了 MapReduce 的调试方式；第二部分为数据分析高级篇，主要介绍高级聚合、大数据分析等技巧；第三部分为系统管理篇，主要介绍 Hadoop 的部署的各种模式、添加新节点、退役节点、快速恢复、MapReduce 调优等。

本书适合各个层次的 Hadoop 技术人员阅读。通过阅读本书，Hadoop 初学者可以使用 Hadoop 来进行数据处理，Hadoop 工程师或者数据挖掘工程师可以解决复杂的业务分析，Hadoop 系统管理员可以更好地进行日常运维。本书也可作为一本 Hadoop 技术手册，针对要解决的相关问题，在工作中随时查阅。

# 译者序

随着 Hadoop 技术在互联网公司的广泛应用，普及程度越来越高，自学 Hadoop 的 Java 程序员也越来越多。大多数人（包括译者本人）自学 Hadoop 的都是从“部署 Hadoop 环境+运行 WordCount 例子”开始，而且大多数自学者也都终止在 WordCount。因没有具体的应用场景而感到学习没有方向，没有成就感。

本书特色是以实践结合理论分析，手把手教读者如何操作，并且对每个操作都做详细的解释，对一些重要的知识点也做了必要的拓展。此外，书中的教学源代码都可以在官网上下载到。本书整体可分成 3 部分，第一部分为基础篇包含第 1 章、第 2 章、第 3 章、第 4 章、第 8 章内容，主要介绍 Hadoop 数据导入导出、HDFS 的概述、Pig 与 Hive 的使用、ETL 和简单的数据处理，还介绍了 MapReduce 的调试方式；第二部分为数据分析高级篇包含第 5 章、第 6 章、第 7 章、第 10 章内容，主要介绍高级聚合、大数据分析等技巧；第三部分为系统管理篇，包含第 9 章，主要介绍 Hadoop 的部署的各种模式、添加新节点、退役节点、快速恢复、MapReduce 调优等。

如果你是 Hadoop 初学者，建议你先阅读第一部分内容，完成这部分内容的学习以后，你基本上可以使用 Hadoop 来进行数据处理。

如果你已经是 Hadoop 工程师或者数据挖掘工程师，可以系统地学习第二部分内容，当然也可以根据需要进行查阅学习。完成这部分内容的学习，有助于解决一些复杂的业务分析。

如果你是 Hadoop 系统管理员，建议你阅读第三部分内容，当然你也可以阅读第一部分的内容，这样更有助于进行日常运维。

本书也可作为一本手册，在教学、工作中随时查阅，解决相关问题。

# 译者简介

**傅杰** 硕士，毕业于清华大学高性能所，现就职于优酷土豆集团，任数据平台架构师，负责集团大数据基础平台建设，支撑其他团队的存储与计算需求，包含 Hadoop 基础平台、日志采集系统、实时计算平台、消息系统、天机镜系统等。个人专注于大数据基础平台架构及安全研究，积累了丰富的平台运营经验，擅长 Hadoop 平台性能调优、JVM 调优及诊断各种 MapReduce 作业，还担任 China Hadoop Submit 2013 大会专家委员、优酷土豆大数据系列课程策划&讲师、EasyHadoop 社区讲师。

**赵磊** 硕士，毕业于中国科学技术大学，现就职于优酷土豆集团，任数据挖掘算法工程师，负责集团个性化推荐和无线消息推送系统的搭建和相关算法的研究。个人专注于基于大数据的推荐算法的研究与应用，积累了丰富的大数据分析与数据挖掘的实践经验，对分布式计算和海量数据处理有深刻的认识。

**卢学裕** 硕士，毕业于武汉大学，曾供职腾讯公司即通部门，现就职于优酷土豆集团，担任大数据技术负责人，负责优酷土豆集团大数据系统平台、大数据分析、数据挖掘和推荐系统。有丰富的 Hadoop 平台使用及优化经验，尤其擅长 MapReduce 的性能优化。基于 Hadoop 生态系统构建了优酷土豆的推荐系统，BI 分析平台。

# 前言

本书能帮助开发者更方便地使用 Hadoop，从而熟练地解决问题。读者会更加熟悉 Hadoop 相关的各种工具从而进行最佳的实践。

本书指导读者使用各种工具解决各种问题。这些工具包括：Apache Hive、Pig、MapReduce、Mahout、Giraph、HDFS、Accumulo、Redis 以及 Ganglia。

本书提供了深入的解释以及代码实例。每章的内容包含一组问题集的描述，并对面临的技术挑战提出了解决方案，最后完整地解决了这些问题。每节将单一问题分解成不同的步骤，这样更容易按照步骤执行相关操作。本书覆盖的内容包括：关于 HDFS 的导入、导出数据，使用 Giraph 进行图分析，使用 Hive、Pig 以及 MapReduce 进行批量数据分析，使用 Mahout 进行机器学习方法，调试并修改 MapReduce 作业的错误，使用 Apache Accumulo 对结构数据进行列存储与检索。

本书的示例中涉及的 Hadoop 技术同样也可以应用于读者自己所面对的问题。

## 本书涵盖哪些内容

第 1 章“Hadoop 分布式文件系统——导入和导出数据”，包含了从一些流行的数据库导入导出数据的方法，包括 MySQL、MongoDB、Greenplum 以及 MSSQL Server。此外，还包括一些辅助工具，例如 Pig、Flume 以及 Sqoop。

第 2 章“HDFS”，介绍从 HDFS 读入或写出数据，介绍了如何使用不同的序列化库，包含 Avro、Thrift 以及 Protocol Buffers。同样包含如何设置数据块大小、备份数以及是否



需要进行 LZO 压缩。

第 3 章“抽取和转换数据”，包含对不同数据源类型进行基本的 Hadoop ETL 操作。不同的工具包括 Hive、Pig 以及 MapRedcue JAVA API，用于批量处理数据，输出一份或多份转换数据。

第 4 章“使用 Hive、Pig 以及 MapReduce 处理常见的任务”，关注于如何利用这些工具提供的函数快速解决不同种类的问题。其中包括字符串连接，外部表的映射，简单表的连接，自定义函数以及基于集群进行分发操作。

第 5 章“高级连接操作”，介绍了 MapReduce、Hive 以及 Pig 中复杂而有效的连接技术。章节的内容包括 Pig 中的归并连接、复制连接以及倾斜连接。同样也包含 Hive 的 map 端连接以及全外连接的内容。同时本章也包括对于外部数据存储，如何使用 Redis 进行数据连接。

第 6 章“大数据分析”，介绍了如何使用 Hadoop 解答关于你的不同数据的各种查询。一些关于 Hive 的例子将展示在不同的分析中如何正确地实现并重复使用用户自定义函数 (UDF)。此外其中有关于 Pig 的两小节，介绍了对于 Audioscrobbler 数据集的各类分析。关于 MapReduce Java API 的一小节介绍了 Combiner 的使用。

第 7 章“高级大数据分析”，展示了使用 Apache Giraph 以及 Mahout 处理不同类型的图分析和机器学习问题。

第 8 章“调试”，帮助你定位并测试 MapReduce 作业。这些例子使用 MRUnit 以及更利于测试的本地模式。此外还强调了使用 counter 以及更新任务状态的重要性，这样做有助于监控 MapReduce 作业。

第 9 章“系统管理”，主要关注如何性能调优 Hadoop 中不同的配置项。下面的内容包含在内：基本的初始化，XML 配置项调整，定位坏数据节点，处理 NameNode 故障，使用 Ganglia 进行性能监控。

第 10 章“使用 Apache Accumulo 进行持久化”，展示了使用 NoSQL 数据存储 Apache Accumulo 带来的很多特性和功能。这些章节利用了这些独特的特性，包括 iterator、combiner、扫描授权以及约束，同时也给出了对于有效建立地理空间行健值以及使用 MapReduce 执行批量分析的例子。

## 阅读需要的准备

读者需要访问一个伪分布式（单台机器）或完全分布式（多台机器）的集群用于执行



本书中的例子。章节中使用的各种工具需要在集群上安装并正确地配置。此外，本书提供的代码使用不同的语言编写，因此读者能访问的机器最好安装了合适的开发工具。

## 读者范围

本书使用简要的代码作为例子，展示了可以使用 Hadoop 解决各类现实中的问题。本书的目的是使不同水平的开发者都能方便地使用 Hadoop 及其工具。Hadoop 的初学者能通过本书加速学习进度并了解现实中 Hadoop 应用的例子。对于有更多经验的 Hadoop 开发者，通过本书，会对许多工具以及技术有新的认识或有一个更清晰的框架，这些东西之前可能你只听说过但并没有真正理解其中的价值。

## 书中的约定

在你阅读本书时，你会发现书中有各种样式的文本，这些不同样式的文本是用来区分不同类型的信息的。下面是一些不同样式文本的实例，以及相应的说明。

文本中的代码如下所示：“所有的 Hadoop 文件系统 shell 命令行都使用统一的形式 `hadoop fs -COMMAND`。”

代码块如下所示：

```
weblogs = load '/data/weblogs/weblog_entries.txt' as
    (md5:chararray,
     url:chararray,
     date:chararray,
     time:chararray,
     ip:chararray);

md5_grp = group weblogs by md5 parallel 4;

store md5_grp into '/data/weblogs/weblogs_md5_groups.bcp';
```

如果我们希望让代码块中的一特殊部分引起你的注意，相关行或条目会以黑体印刷：

```
weblogs = load '/data/weblogs/weblog_entries.txt' as
    (md5:chararray,
     url:chararray,
     date:chararray,
     time:chararray,
```

```
ip:chararray);
```

```
md5_grp = group weblogs by md5 parallel 4;
```

```
store md5_grp into '/data/weblogs/weblogs_md5_groups.bcp';
```

命令行输入或输出都是以如下样式书写的：

```
hadoop distcp -m 10 hdfs://namenodeA/data/weblogs hdfs://namenodeB/data/weblogs
```

新的术语以及重要文字会以加粗的字体出现。你在屏幕上（如菜单或者对话框中）见到的文字都会是像后面这样出现在正文中：“为了构建 JAR 文件，下载 Jython Java 安装程序，并执行该程序，从安装菜单中选择 **Standalone** 选项。”



方框中出现的为警告或重要的注解。



提示或小技巧出现在这里。

## 读者反馈

我们总是欢迎来自读者的反馈。请告诉我们你觉得这本书怎么样，你喜欢哪些内容，不喜欢哪些内容。读者的反馈对于帮助我们写出那些对读者真正有用的内容至关重要。

如果是给我们反馈一些普通的信息，你可以给 [feedback@packtpub.com](mailto:feedback@packtpub.com) 这个邮箱发一封邮件即可，记得在你邮件的标题中提及相应的书名。

如果你是某一方面的专家并且对于写作或者撰稿有兴趣的话，你可以访问 [www.packtpub.com/authors](http://www.packtpub.com/authors)，读读我们的作者指南。

## 客户支持

现在你成为了 Packt 出版社的一名尊敬的用户，为使你的购买物超所值，我们为你准备了一系列的东西。

## 下载本书中的示例代码

你可以登录账号下载到所有从 <http://www.packtpub.com> 购买的 Packt 的书的示例代码文件。如果你从别的地方购买了此书，可以访问 <http://www.packtpub.com/support> 登记信息，文件将通过邮件直接发给你。

## 勘误

尽管我们已经非常小心谨慎，以确保内容的准确性，错误还是不可避免。如果你在书中发现了错误（这种错误可能是文字或者代码方面的），若你能向我们报告这些错误，我们将感激不尽。这样做，可以使其他读者免受这些错误带来的困扰，帮助改善该书的下一个版本。如果你发现了什么错误，请访问 <http://www.packtpub.com/support>，选择书名，点击“提交错误”链接，然后输入你发现错误的详细内容，通知我们。一旦你指出的错误得到确认，你提交的内容就会被采纳，并加入一个已经存在的勘误列表中。你也可以访问这个链接 <http://www.packtpub.com/support>，选择书名，查看相应书本已有的勘误表。

## 关于盗版

所有媒体的版权材料在互联网上被盗版是一个日趋严重的问题。在 Packt，我们对于版权与许可证的保护工作是十分看重的。如果你发现任何非法复制我们作品的现象，无论以任何形式，只要在互联网上，请提供给我们对应的网络地址或网站名称，我们会立即对其行为进行纠正。

请通过 [copyright@packtpub.com](mailto:copyright@packtpub.com) 联系我们，并附上可疑盗版材料的链接。

我们感谢你对保护作者权益的帮助，你的协助同时也保障了我们带给你更多有价值内容的能力。

## 如果你有疑问

如果你对书的某些方面有疑问，你可以通过 [questions@packtpub.com](mailto:questions@packtpub.com) 联系我们，我们会尽最大的努力为你解答。

# 作者简介

**Jonathan R. Owens:** 软件工程师，拥有 Java 和 C++ 技术背景，最近主要从事 Hadoop 及相关分布式处理技术工作。

目前就职于 comScore 公司，为核心数据处理团队成员。comScore 是一家知名的从事数字测评与分析的公司，公司使用 Hadoop 及其他定制的分布式系统对数据进行聚合、分析和管理工作，每天处理超过 400 亿单的交易。

---

感谢我的父母 James 和 Patricia Owens，感谢他们对我的支持以及从小给我介绍科技知识。

---

**Jon Lentz:** comScore 核心数据处理团队软件工程师。comScore 是一家知名的在线受众测评与分析的公司。他更倾向于使用 Pig 脚本来解决问题。在加入 comScore 之前，他主要开发优化供应链和分配固定收益证券的软件。

---

我的女儿 Emma 出生在本书写作的过程中。感谢她深夜的陪伴！

---

**Brian Femiano:** 本科毕业于计算机专业，并且从事相关专业软件开发工作 6 年，最近两年主要利用 Hadoop 构建高级分析与大数据存储。他拥有商业领域的相关经验，以及丰富的政府合作经验。他目前就职于弗吉尼亚的 Potomac Fusion 公司，这家公司主要从事可扩展算法的开发，并致力于学习并改进政府领域中最先进和最复杂的数据集。他通过教授课程和会议培训在公司内部普及 Hadoop 和云计算相关的技术。

---

我要感谢合著者的耐心，以及他们为你们可以在书上看到的代码做出的努力。同时也要感谢 Potomac Fusion 的许多同事，他们驾驭最前沿技术的才能和激情，以及促进知识传播的精神一直鼓舞着我。

---

# 审阅者简介

**Edward J. Cody:** 作家、演讲师，大数据仓库、Oracle 商业智能和 Hyperion EPM 实现的专家，是 *The Business Analyst's Guide to Oracle Hyperion Interactive Reporting 11* 的作者，以及 *The Oracle Hyperion Interactive Reporting 11 Expert Guide* 的合著者。在过去的职业生涯中，他给商业公司和联邦政府都做过咨询，目前主要从事大型 EPM、BI 和大数据仓库的实现研究。

---

本书的作者做了很好的工作，同时我要感谢 Packt 出版社让我有机会参与本书的编辑出版工作。

---

**Daniel Jue:** Sotera Defense Solutions 公司的高级软件工程师，Apache 软件基金会成员之一，他曾在和平与战乱地区为 ACSIM、DARPA 和许多联邦机构工作，致力于揭示蕴藏在“大数据”背后的动力学与异常现象。Daniel 拥有马里兰大学帕克分校计算机专业的学士学位，并同时专注于物理学与天文学的研究，他目前的兴趣是将分布式人工智能技术应用到自适应异构的云计算中。

---

感谢我漂亮的妻子 Wendy，以及我的双胞胎儿子 Christopher 和 Jonathan，在我审阅本书的过程中，他们给予了我关爱与耐心。非常感谢 Brian Femiano、Bruce Miller 和 Jonathan Larson，他们让我接触到了许多伟大的想法、观点，使我深受启发。

---

**Bruce Miller:** Sotera Defense Solutions 公司高级软件工程师，目前受雇于美国国防部高级研究计划署（DARPA），并专注于大数据的软件开发 10 余年。工作之余喜欢用 Haskell 和 Lisp 语言进行编程，并将其应用于解决一些实际的问题。

# 目录

|                                           |    |
|-------------------------------------------|----|
| 第 1 章 Hadoop 分布式文件系统——导入和导出数据 .....       | 1  |
| 1.1 介绍 .....                              | 1  |
| 1.2 使用 Hadoop shell 命令导入和导出数据到 HDFS ..... | 2  |
| 1.3 使用 distcp 实现集群间数据复制 .....             | 7  |
| 1.4 使用 Sqoop 从 MySQL 数据库导入数据到 HDFS .....  | 9  |
| 1.5 使用 Sqoop 从 HDFS 导出数据到 MySQL .....     | 12 |
| 1.6 配置 Sqoop 以支持 SQL Server .....         | 15 |
| 1.7 从 HDFS 导出数据到 MongoDB .....            | 17 |
| 1.8 从 MongoDB 导入数据到 HDFS .....            | 20 |
| 1.9 使用 Pig 从 HDFS 导出数据到 MongoDB .....     | 23 |
| 1.10 在 Greenplum 外部表中使用 HDFS .....        | 24 |
| 1.11 利用 Flume 加载数据到 HDFS 中 .....          | 26 |
| 第 2 章 HDFS .....                          | 28 |
| 2.1 介绍 .....                              | 28 |
| 2.2 读写 HDFS 数据 .....                      | 29 |
| 2.3 使用 LZO 压缩数据 .....                     | 31 |
| 2.4 读写序列化文件数据 .....                       | 34 |
| 2.5 使用 Avro 序列化数据 .....                   | 37 |
| 2.6 使用 Thrift 序列化数据 .....                 | 41 |
| 2.7 使用 Protocol Buffers 序列化数据 .....       | 44 |
| 2.8 设置 HDFS 备份因子 .....                    | 48 |



|                                                            |            |
|------------------------------------------------------------|------------|
| 2.9 设置 HDFS 块大小 .....                                      | 49         |
| <b>第 3 章 抽取和转换数据 .....</b>                                 | <b>51</b>  |
| 3.1 介绍 .....                                               | 51         |
| 3.2 使用 MapReduce 将 Apache 日志转换为 TSV 格式 .....               | 52         |
| 3.3 使用 Apache Pig 过滤网络服务器日志中的爬虫访问量 .....                   | 54         |
| 3.4 使用 Apache Pig 根据时间戳对网络服务器日志数据排序 .....                  | 57         |
| 3.5 使用 Apache Pig 对网络服务器日志进行会话分析 .....                     | 59         |
| 3.6 通过 Python 扩展 Apache Pig 的功能 .....                      | 61         |
| 3.7 使用 MapReduce 及二次排序计算页面访问量 .....                        | 62         |
| 3.8 使用 Hive 和 Python 清洗、转换地理事件数据 .....                     | 67         |
| 3.9 使用 Python 和 Hadoop Streaming 执行时间序列分析 .....            | 71         |
| 3.10 在 MapReduce 中利用 MultipleOutputs 输出多个文件 .....          | 75         |
| 3.11 创建用户自定义的 Hadoop Writable 及 InputFormat 读取地理事件数据 ..... | 78         |
| <b>第 4 章 使用 Hive、Pig 和 MapReduce 处理常见的任务 .....</b>         | <b>85</b>  |
| 4.1 介绍 .....                                               | 85         |
| 4.2 使用 Hive 将 HDFS 中的网络日志数据映射为外部表 .....                    | 86         |
| 4.3 使用 Hive 动态地为网络日志查询结果创建 Hive 表 .....                    | 87         |
| 4.4 利用 Hive 字符串 UDF 拼接网络日志数据的各个字段 .....                    | 89         |
| 4.5 使用 Hive 截取网络日志的 IP 字段并确定其对应的国家 .....                   | 92         |
| 4.6 使用 MapReduce 对新闻档案数据生成 n-gram .....                    | 94         |
| 4.7 通过 MapReduce 使用分布式缓存查找新闻档案数据中包含关键词的行 .....             | 98         |
| 4.8 使用 Pig 加载一个表并执行包含 GROUP BY 的 SELECT 操作 .....           | 102        |
| <b>第 5 章 高级连接操作 .....</b>                                  | <b>104</b> |
| 5.1 介绍 .....                                               | 104        |
| 5.2 使用 MapReduce 对数据进行连接 .....                             | 104        |
| 5.3 使用 Apache Pig 对数据进行复制连接 .....                          | 108        |
| 5.4 使用 Apache Pig 对有序数据进行归并连接 .....                        | 110        |
| 5.5 使用 Apache Pig 对倾斜数据进行倾斜连接 .....                        | 111        |
| 5.6 在 Apache Hive 中通过 map 端连接对地理事件进行分析 .....               | 113        |
| 5.7 在 Apache Hive 通过优化的全外连接分析地理事件数据 .....                  | 115        |

|                                                            |     |
|------------------------------------------------------------|-----|
| 5.8 使用外部键值存储 (Redis) 连接数据 .....                            | 118 |
| 第 6 章 大数据分析 .....                                          | 123 |
| 6.1 介绍 .....                                               | 123 |
| 6.2 使用 MapReduce 和 Combiner 统计网络日志数据集中的独立 IP 数 .....       | 124 |
| 6.3 运用 Hive 日期 UDF 对地理事件数据集中的时间日期进行转换与排序 .....             | 129 |
| 6.4 使用 Hive 创建基于地理事件数据的每月死亡报告 .....                        | 131 |
| 6.5 实现 Hive 用户自定义 UDF 用于确认地理事件数据的来源可靠性 .....               | 133 |
| 6.6 使用 Hive 的 map/reduce 操作以及 Python 标记最长的无暴力发生的时间区间 ..... | 136 |
| 6.7 使用 Pig 计算 Audioscrobbler 数据集中艺术家之间的余弦相似度 .....         | 141 |
| 6.8 使用 Pig 以及 datafu 剔除 Audioscrobbler 数据集中的离群值 .....      | 145 |
| 第 7 章 高级大数据分析 .....                                        | 147 |
| 7.1 介绍 .....                                               | 147 |
| 7.2 使用 Apache Giraph 计算 PageRank .....                     | 147 |
| 7.3 使用 Apache Giraph 计算单源最短路径 .....                        | 150 |
| 7.4 使用 Apache Giraph 执行分布式宽度优先搜索 .....                     | 158 |
| 7.5 使用 Apache Mahout 计算协同过滤 .....                          | 165 |
| 7.6 使用 Apache Mahout 进行聚类 .....                            | 168 |
| 7.7 使用 Apache Mahout 进行情感分类 .....                          | 171 |
| 第 8 章 调试 .....                                             | 174 |
| 8.1 介绍 .....                                               | 174 |
| 8.2 在 MapReduce 中使用 Counters 监测异常记录 .....                  | 174 |
| 8.3 使用 MRUnit 开发和测试 MapReduce .....                        | 177 |
| 8.4 本地模式下开发和测试 MapReduce .....                             | 179 |
| 8.5 运行 MapReduce 作业跳过异常记录 .....                            | 182 |
| 8.6 在流计算作业中使用 Counters .....                               | 184 |
| 8.7 更改任务状态显示调试信息 .....                                     | 185 |
| 8.8 使用 illustrate 调试 Pig 作业 .....                          | 187 |
| 第 9 章 系统管理 .....                                           | 189 |
| 9.1 介绍 .....                                               | 189 |
| 9.2 在伪分布模式下启动 Hadoop .....                                 | 189 |