

Hadoop

硬实战

[美] Alex Holmes 著
梁李印 宁青 杨卓萃 译

Hadoop in Practice

含85个技术点



电子工业出版社

PUBLISHING HOUSE OF ELECTRONICS INDUSTRY
<http://www.phei.com.cn>

Hadoop in Practice

Hadoop硬实战

[美] Alex Holmes 著

梁李印 宁青 杨卓萃 译

电子工业出版社

Publishing House of Electronics Industry

北京•BEIJING

内 容 简 介

Hadoop 是一个开源的 MapReduce 平台,设计运行在大型分布式集群环境中,为开发者进行数据存储、管理以及分析提供便利的方法。本书详细讲解了 Hadoop 和 MapReduce 的基本概念,并收集了 85 个问题及其解决方案。在关键问题领域对基础概念和实战方法做了权衡。

本书适合使用 Hadoop 进行数据存储、管理和分析的技术人员使用。

Original English Language edition published by Manning Publications, 178 South Hill Drive, Westampton, NJ 08060 USA. Copyright ©2012 by Manning Publications. Simplified Chinese-language edition copyright ©2015 by Publishing House of Electronics Industry. All rights reserved.

本书简体中文版专有出版权由 Manning Publications 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:01-2013-4143

图书在版编目(CIP)数据

Hadoop 硬实战 / (美)霍姆斯(Holmes,A.)著;梁李印,宁青,杨卓萃译. —北京:电子工业出版社,2015.1
书名原文: Hadoop in practice
ISBN 978-7-121-25072-9

I. ①H… II. ①霍… ②梁… ③宁… ④杨… III. ①数据处理软件 IV. ①TP274

中国版本图书馆 CIP 数据核字(2014)第 288388 号

策划编辑:张春雨

责任编辑:刘 舫

印 刷:北京中新伟业印刷有限公司

装 订:三河市皇庄路通装订厂

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:33.5

字数:750 千字

版 次:2015 年 1 月第 1 版

印 次:2015 年 1 月第 1 次印刷

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。

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

服务热线:(010) 88258888。

To Michal, Marie, Oliver, Ollie, Mish, and Anch

译者序

在淘宝数据平台架构组工作了两年，深知一套 Hadoop 平台从运维到开发实属不易。本书是一本实战性的书，讲的大多都是术，解决生产问题的术。通过阅读本书，读者可以较为轻松地完成环境搭建，故障诊断与解决，数据的开发等。阅读本书同样需要读者有一定的 Hadoop 理论知识，否则只知其然不知其所以然。同时本书作者也强烈推荐阅读 *Hadoop: The Definitive Guid* 一书来深入了解 Hadoop 原理。

《Hadoop 硬实战》中文译本的问世离不开朋友们的帮忙——梁李印（无影），数据平台 Hadoop 专家，专攻 MapReduce 框架开发，其维护的 Hadoop 代码在 7000 多台节点的集群上日日夜夜地运行着；杨卓萃（卓萃），Hive 专家，每晚有上万个 Hive 作业通过其维护的 Hive 代码生成。

同时还要感谢平安科技的汪洋、王鹏冲、张建龙对我工作上的支持，让我抽出更多时间来做其他的事情。感谢淘宝数据平台的王磊（泽远）、王勇（图海）、李庐阳（刘顺）、吴威（无畏）和其他同学，据说“双十一”的销售额已多达 571 亿元，你们的股票涨了多少？感谢思科合肥 Grant Pan、Chris Gao、Liyang Tang，我想你们了。

宁青

2014 年 11 月 11 日于观澜湖畔

前言

我是在 2008 年秋天参加 VeriSign 的互联网抓取和分析项目的时候开始接触 Hadoop 的，我的团队与 Doug Cutting 以及 Nutch 项目的人就如何有效地存储和管理 TB 级的抓取数据和分析数据得出相似的结论。当时我们已经建立了自己的分布式系统，但是这个系统无法系统地将抓取的数据与新添加的数据合并。

在研究了 Hadoop 项目后，我们发现它很适合我们的需求——它支持大数据存储，并提供数据合并机制。在几个月的时间内，我们创建并开发了包含多个 MapReduce 作业的 MapReduce 应用，并将这个应用与我们自己的 MapReduce workflow 管理系统部署在一个拥有 18 个节点的小集群上。通过这个小集群可以观察我们的 MapReduce 作业是如何在几分钟内处理完数据的。当然，我们没想到的是，在调试和优化 MapReduce 作业上会花费这么多时间；更没想到的是，我们担负起生产管理员这一新职责时，在支持生产的头几个月内我们遇到大量的磁盘故障问题。

随着对 Hadoop 熟悉程度的提高，我们运用 Hadoop 继续建立了更多的功能以帮助处理大规模数据集。我们还开始在公司内部宣传 Hadoop 的优点，并发动其他面临处理大数据的项目使用 Hadoop。

在运用 Hadoop（尤其是在处理 MapReduce 时）的过程中我们遇到的最大挑战是重新学习如何使用它解决问题。MapReduce 有自己并行处理进程的方法，这种处理方式与我们通常使用的 JVM 程序完全不同。最大的障碍就是训练我们的大脑去熟悉 MapReduce 的处理方法，Chuck Lam 在 2010 年编写出版的 *Hadoop in Action* 一书详细讲解了 MapReduce 的相关信息。

当你习惯了使用 MapReduce 后，接下来就需要学习如何使用 Hadoop，如何将数据导入和导出 HDFS，如何在 Hadoop 中高效处理数据。Hadoop 的这些方面还没有引起很广泛的关注，这也是我要撰写本书的潜在原因——主要介绍 Hadoop 的一些高级应用，并涉及 Hadoop 的一些难点问题。

我相信目前已经有很多读者有 Hadoop 的相关应用经验，我写本书的目的只是为了将我的个人经验转化为书本知识。我对本书中的实例进行了验证，这并不是一个愉快的过程，但是编写本书的过程中所运用到的新方法和工具使我对 Hadoop 有了进一步的了解，希望读者通过本书可以获取 Hadoop 的更多知识。

致谢

首先，我最想感谢的是推举我写这本书的 Michael Noll，他也是最先收到本书的样本的，并帮忙修改了本书的组织结构。在编写本书的过程中，他对我的支持和鼓舞程度无法用言语形容。

我还要感谢本书的策划编辑 Cynthia Kane，在编写本书的过程中，她对这本书提出了宝贵的建议，最大的建议就是教我使用可视化工具解释复杂的概念。

我还想感谢本书的所有评审：Aleksi Sergeevich, Alexander Luya, Asif Jan, Ayon Sinha, Bill Graham, Chris Nauroth, Eli Collins, Ferdy Galema, Harsh Chouraria, Jeff Goldschrafe, Maha Alabduljalil, Mark Kemna, Oleksey Gayduk, Peter Krey, Philipp K. Janert, Sam Ritchie, Soren Macbeth, Ted Dunning, Yunkai Zhang 和 Zhenhua Guo。

本书的技术主编 Jonathan Seidman 在本书出版之前对本书进行了整体审查。非常感谢 Josh Wills, Crunch 的创始人，他审查了有关 Crunch 的章节。还要感谢 Josh Patterson，他帮忙审查了 Mahout 章节。

与 Manning 的工作人员一起工作很愉快，特别是 Troy Mott, Katie Tennant, Nick Chase, Tara Walsh, Bob Herbstman, Michael Stephens, Marjan Bace 和 Maureen Spencer。

最后，特别感谢我的妻子，Michal，她不得不忍受疯狂工作的我，在编写本书的过程中，她是我最大的鼓舞源。

关于本书

Hadoop 的创始人 Doug Cutting，喜欢称 Hadoop 为大数据的内核，我也比较赞成。凭借其分布式存储和计算能力，Hadoop 是从根本上处理庞大数据集的技术。Hadoop 在结构化数据（RDBMS）和非结构化数据（日志文件、XML、文本文件）之间建立了连接桥，使这些数据之间可以自由连接。它也已经从合并 OLTP 和日志文件的传统应用实例发展为更复杂的应用，如使用 Hadoop 创建数据仓库（如 Facebook）和挖掘数据中信息的数据科学领域。

本书收集了一些中高级的 Hadoop 实例，并以问题 / 解决方案的形式展示出来。85 个技术点中的每个技术点都提出了你将面临的具体任务，如使用 Flume 将日志文件导入 Hadoop，或者使用 Mahout 进行预测分析。每个问题都按步骤进行讲解，一步步参看这些讲解，你会慢慢了解 Hadoop 的大数据世界。

本书的目标读者是了解 Hadoop 的基础概念并运用过一些 Hadoop 实例的用户。Manning 出版的由 Chuck Lam 编写的 *Hadoop in Action* 一书讲解了 Hadoop 的基础知识。

本书中的大部分技术都是基于 Java 的，这意味着读者需要熟悉 Java。由 Joshua Bloch 编写的 *Effective Java* 一书的第二版适合所有级别的 Java 用户参看。

概览

本书的 13 个章节划分为 5 个部分。

第 1 部分只包含一章，主要包含本书的简介。讲解了 Hadoop 的基础概念和如

何在单个主机上启动和运行 Hadoop, 并讲解了如何编写和执行一个 MapReduce 作业。

第 2 部分“数据逻辑”由两章组成, 其中涵盖了用于处理数据的相关技术和工具、如何将数据导入和导出 Hadoop, 以及如何处理不同的数据格式。将数据导入 Hadoop 通常是在使用 Hadoop 过程中遇到的第一个难题, 第 2 章讲解了处理通用数据源的工具。第 3 章包含了如何在 MapReduce 中处理常用的数据格式, 如 XML 和 JSON, 以及讲解了哪些数据格式更适合大数据集。

第 3 部分“大数据模式”, 讲解了提高处理大数据集的技术。第 4 章讲解了如何优化 MapReduce 连接和排序操作, 第 5 章讲解了如何处理大量小文件和压缩。第 6 章描述了如何调试 MapReduce 性能问题和提高作业运行速度的技术。

第 4 部分“数据科学”, 深入研究了从数据中获取有用信息的工具和方法。第 7 章包含如何使用 MapReduce 表示图形这样的数据结构, 并分析了用于处理图像数据的算法。第 8 章描述了流行的数据统计和挖掘平台 R 如何与 Hadoop 集成。第 9 章描述了如何使用 Mahout 结合 MapReduce 进行大规模扩展预测分析。

第 5 部分“驯服大象”, 简化了 MapReduce 的技术。第 10 章和第 11 章分别讲解了 Hive 和 Pig, 它们都是 MapReduce 的域特定语言 (DSL), 为 MapReduce 提供上层抽象。第 12 章研究了 Crunch 和 Cascading, 是提供它们自己的 MapReduce 抽象的 Java 库。第 13 章讲解了编写单元测试的技术和调试 MapReduce 问题的相关技术。

附录 A 中包含了 Hadoop 和本书中涵盖的其他技术的安装指南。附录 B 包含了第 2 章中介绍的低层次的 Hadoop 导入 / 导出机制中使用到的工具。附录 C 讲解了 HDFS 如何支持读写操作, 而附录 D 介绍了作者在第 4 章中编写的一些 MapReduce 合并框架。

编码规范和下载

本书中的所有文本和实例都适用于 Hadoop 0.20.x (和 1.x), 且大部分代码都是使用新的 MapReduce API `org.apache.hadoop.mapreduce` 编写的。少量使用旧的 `org.apache.hadoop.mapreduce` 包编写的实例通常是使用旧 API 的第三方库或工具生成的结果。

本书中使用的所有代码可以在 GitHub 上的 <https://github.com/alexholmes/hadoop-book> 中获取, 也可以在出版商的网站 www.manning.com/HadoopinPractice 上获取。

需要通过 Java 1.6 以上和 Maven 3.0 以上的版本编译这些代码。git 是一个源码控制管理系统, GitHub 提供托管 git 库的服务。Maven 用于构建系统。

你可以使用下面的命令复制 (下载) 我的 GitHub 库:

```
$ git clone git://github.com/alexholmes/hadoop-book.git
```

当数据源建立以后，就可以编译代码了：

```
$ cd hadoop-book
$ mvn package
```

这将创建一个 Java JAR 文件：target/hadoop-book-1.0.0-SNAPSHOT-jar-with-dependencies.jar。通过 bin/run.sh 运行代码。

如果在 CDH 分布式系统中运行这些代码，不需要配置就可运行脚本。如果在其他系统上运行这些代码，需要设置 HADOOP_HOME 环境变量指向你的 Hadoop 安装目录。

bin/run.sh 脚本读取的第一个参数是实例的完整的 Java 类名，其次读取的是预期的示例类的任一参数。请按照下面的步骤来运行第 1 章中倒序索引的 MapReduce 代码：

```
$ hadoop fs -mkdir /tmp
$ hadoop fs -put test-data/ch1/* /tmp/

# replace the path below with the location of your Hadoop installation
# this isn't required if you are running CDH3
export HADOOP_HOME=/usr/local/hadoop

$ bin/run.sh com.manning.hip.ch1.InvertedIndexMapReduce \
/tmp/file1.txt /tmp/file2.txt output
```

如果没有安装 Hadoop，上面的代码将不能运行。请参看第 1 章中有关 CDH 的安装信息，或者附录 A 中的 Apache 指南。

第三方库

为了方便，我们使用了一些第三方库，它们被包含在 Maven 创建的 jar 中，不需要对它们做其他处理。下面的表格中包含了代码实例中使用到的库。

通用第三方库

库	链接	描述
Apache Commons IO	http://commons.apache.org/io/	用 Java 写的工具包，提供输入输出流支持。你将经常使用 IOUtils 来关闭连接并将内容从文件读取到内存中
Apache Commons Lang	http://commons.apache.org/lang/	字符串、日期、集合类工具包。你将经常使用 StringUtils 类来做切分

数据集

本书中的实例使用到 3 个数据集，所有的数据集都很小，这样方便处理，可以在 <https://github.com/alexholmes/hadoop-book/tree/master/test-data> 目录中的 GitHub 库里获取数据集的副本。某个章节具体所需要的数据存放在 GitHub 目录中对应的章节子目录中。

纳斯达克金融股票

我从 Infochimps 中下载了纳斯达克日交易数据 (<http://mng.bz/xjwc>)，过滤了一些数据，只保留了 5 支股票从 2000 年到 2009 年的交易数据。本书中使用的数据可从 GitHub 上的 <https://github.com/alexholmes/hadoop-book/blob/master/test-data/stocks.txt> 中获取。

这些数据是 CSV 格式的，数据中的字段是按照以下顺序排列的：

```
Symbol,Date,Open,High,Low,Close,Volume,Adj Close
```

Apache 日志数据

我在 Apache Common Log Format 创建了一个日志文件样例（参看 <http://mng.bz/L4S3>），以及一些伪 E 类 IP 地址和一些虚拟的资源 and 响应代码。可以在 GitHub 上的 <https://github.com/alexholmes/hadoop-book/blob/master/test-data/apachelog.txt> 中获取这个文件。

名称

通过政府的人口普查文件 <http://mng.bz/LuFB> 检索名称，这个文件可以通过以下网址 <https://github.com/alexholmes/hadoop-book/blob/master/test-data/names.txt> 获取。

获取帮助

毫无疑问，在使用 Hadoop 的过程中肯定会遇到一些问题，幸运的是，可以在维基和用户社区中得到解答。

Hadoop 的维基地址是 <http://wiki.apache.org/hadoop/>，其中包含了有用的介绍和安装说明以及故障排除说明。

Hadoop Common、HDFS 和 MapReduce 邮件列表可以在 http://hadoop.apache.org/mailling_lists.html 获取。

Search Hadoop 是一个有用的网站，列出了 Hadoop 和 Hadoop 生态系统项目，它还提供全文本搜索功能：<http://search-Hadoop.com/>。

为了获取 Hadoop 的最新消息，你必须订阅一些有用的博客。下面列出的博客

是我的最爱：

- Cloudera 主要写 Hadoop 实际应用相关的知识：<http://www.cloudera.com/blog/>。
- Hortonworks 博客值得一读，主要讨论应用和 Hadoop 的未来项目路线图：<http://hortonworks.com/blog/>。
- Michael Noll 是第一个提供 Hadoop 详细安装说明的博主，他主要讲解 Hadoop 在实际应用中的挑战：<http://www.michael-noll.com/blog/>。

还有一些活跃的 Hadoop Twitter 用户你需要关注，其中包括 Arun Murthy (@acmurthy)、Tom White (@tom_e_white)、Eric Sammer (@esamme)、Doug Cutting (@cutting) 和 Todd Lipcon (@tlipcon)，Hadoop 项目本身的 Twitter 账号是 @Hadoop。

关于作者

Alex Holmes 是一个拥有 15 年大数据分布式 Java 系统开发经验的高级软件工程师。在最近 4 年中，他通过解决跨多个项目的大数据问题获取了 Hadoop 专业知识。他曾经就职于 JavaOne 和 Jazoon，目前是 VeriSign 的技术总监。

Alex 维护了一个与 Hadoop 相关的博客：<http://grepalex.com>，他的 Twitter 地址是 https://twitter.com/grep_alex。

关于封面插图

本书封面上的人物是“来自 Kistanja 的年轻人 Dalmatia”。插图来自 Nikola Arsenovic 的“19 世纪中叶的克罗地亚传统服装”画册，该画册在 2003 年由克罗地亚的 Split 民族博物馆出版，这幅插图是从 Split 民族博物馆的一个图书管理员那里获取的，这个博物馆坐落在中世纪的罗马镇中心：是公元 304 年皇帝戴克退位后的先宫废墟。这本书包括了来自克罗地亚的不同区域的人物彩色插图，这些人物插图穿着不同的服饰并拿着不同的生活用品。

Kistanja 是克罗地亚 Bukovica 地区的一个小镇，位于达尔马提亚北部，该地区拥有丰富的罗马和威尼斯历史。mamok 这个词在克罗地亚表示学士、男友或追求者——一个到了适婚年龄的青年男子——封面上的男子衣冠楚楚，着白色衬衣和五颜六色的绣花背心，显然是他最好的衣服，这些衣服一般是在去教堂或喜庆场合穿的，或者是去拜访一位年轻的女士的时候穿的。

在过去的 200 年中，人们的着装和生活方式发生了变化，多元化的区域文化随着时间消失了。现在很难从生活习惯分辨不同区域的人，现在村庄或城镇之间间隔距离很近。也许我们已经从文化多样性转变为更多样化的个人生活——肯定是一个更多样化和快节奏的科技生活。

Manning 出版社通过在封面上呈现两个世纪之前的丰富的区域文化生活来庆祝今天计算机科技的创造性和主动性，并将读者带回到来自旧书和收藏品中的插图所显示的生活中。

目录

前言	XV
致谢	XVII
关于本书	XIX
第 1 部分 背景和基本原理	1
1 跳跃中的 Hadoop	3
1.1 什么是 Hadoop	4
1.1.1 Hadoop 的核心组件	5
1.1.2 Hadoop 生态圈	9
1.1.3 物理架构	10
1.1.4 谁在使用 Hadoop	12
1.1.5 Hadoop 的局限性	13
1.2 运行 Hadoop	14
1.2.1 下载并安装 Hadoop	14
1.2.2 Hadoop 的配置	15
1.2.3 CLI 基本命令	17
1.2.4 运行 MapReduce 作业	18
1.3 本章小结	24

第 2 部分 数据逻辑..... 25

2 将数据导入导出 Hadoop.....27

2.1 导入导出的关键要素	29
2.2 将数据导入 Hadoop.....	30
2.2.1 将日志文件导入 Hadoop	31
技术点 1 使用 Flume 将系统日志文件导入 HDFS.....	33
2.2.2 导入导出半结构化和二进制文件	42
技术点 2 自动复制文件到 HDFS 的机制	43
技术点 3 使用 Oozie 定期执行数据导入活动	48
2.2.3 从数据库中拉数据	52
技术点 4 使用 MapReduce 将数据导入数据库	53
技术点 5 使用 Sqoop 从 MySQL 导入数据	58
2.2.4 HBase	68
技术点 6 HBase 导入 HDFS	68
技术点 7 将 HBase 作为 MapReduce 的数据源	70
2.3 将数据导出 Hadoop.....	73
2.3.1 将数据导入本地文件系统	73
技术点 8 自动复制 HDFS 中的文件	73
2.3.2 数据库	74
技术点 9 使用 Sqoop 将数据导入 MySQL.....	75
2.3.3 Hbase	78
技术点 10 将数据从 HDFS 导入 HBase	78
技术点 11 使用 HBase 作为 MapReduce 的数据接收器	79
2.4 本章小结	81

3 数据序列化——处理文本文件及其他格式的文件83

3.1 了解 MapReduce 中的输入和输出	84
3.1.1 数据输入	85
3.1.2 数据输出	89
3.2 处理常见的序列化格式	91
3.2.1 XML	91

技术点 12 MapReduce 和 XML	91
3.2.2 JSON.....	95
技术点 13 MapReduce 和 JSON	95
3.3 大数据的序列化格式	99
3.3.1 比较 SequenceFiles、Protocol Buffers、Thrift 和 Avro	99
3.3.2 Sequence File.....	101
技术点 14 处理 SequenceFile.....	103
3.3.3 Protocol Buffers.....	109
技术点 15 整合 Protocol Buffers 和 MapReduce	110
3.3.4 Thrift.....	117
技术点 16 使用 Thrift.....	117
3.3.5 Avro	119
技术点 17 MapReduce 的下一代数据序列化技术	120
3.4 自定义文件格式	127
3.4.1 输入输出格式	127
技术点 18 输入和输出格式为 CSV 的文件	128
3.4.2 output committing 的重要性	136
3.5 本章小结	136

第 3 部分 大数据模式.....137

4 处理大数据的 MapReduce 模式 139

4.1 Join.....	140
4.1.1 Repartition Join	141
技术点 19 优化 repartition join	142
4.1.2 Replicated Join	146
4.1.3 Semi-join	147
技术点 20 实现 semi-join	148
4.1.4 为你的数据挑选最优的合并策略	154
4.2 排序	155
4.2.1 二次排序	156
技术点 21 二次排序的实现	157