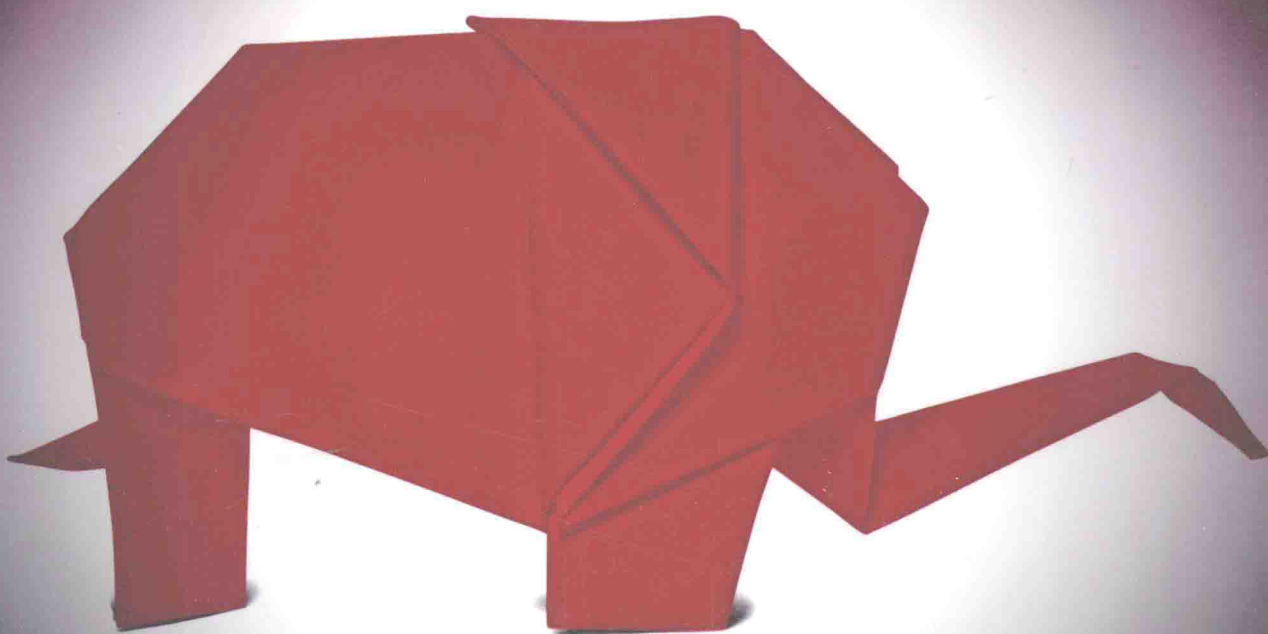




Professional Hadoop Solutions

Hadoop高级编程

——构建与实现大数据解决方案

Boris Lublinsky

[美] Kevin T. Smith 著

Alexey Yakubovich

穆玉伟 靳晓辉 译



清华大学出版社

Hadoop 高级编程

——构建与实现大数据解决方案

Boris Lublinsky
[美] Kevin T. Smith 著
Alexey Yakubovich
穆玉伟 靳晓辉 译

清华大学出版社

Boris Lublinsky, Kevin T. Smith, Alexey Yakubovich

Professional Hadoop Solutions

EISBN: 978-1-118-61193-7

Copyright © 2013 by John Wiley & Sons, Inc., Indianapolis, Indiana

All Rights Reserved. This translation published under license.

本书中文简体字版由 Wiley Publishing, Inc. 授权清华大学出版社出版。未经出版者书面许可, 不得以任何方式复制或抄袭本书内容。

北京市版权局著作权合同登记号 图字: 01-2013-8907

Copies of this book sold without a Wiley sticker on the cover are unauthorized and illegal.

本书封面贴有 Wiley 公司防伪标签, 无标签者不得销售。

版权所有, 侵权必究。侵权举报电话: 010-62782989 13701121933

图书在版编目(CIP)数据

Hadoop高级编程——构建与实现大数据解决方案 / (美)卢博林斯基(Lublinsky, B.), (美)史密斯(Smith, K. T.), (美)雅库伯维奇(Yakubovich, A.) 著; 穆玉伟, 靳晓辉 译; —北京: 清华大学出版社, 2014

书名原文: Professional Hadoop Solutions

ISBN 978-7-302-36906-6

I. ①H… II. ①卢… ②史… ③雅… ④穆… ⑤靳… III. ①数据处理软件—程序设计 IV. ①TP274

中国版本图书馆 CIP 数据核字(2014)第 131674 号

责任编辑: 王 军 刘伟琴

装帧设计: 孔祥峰

责任校对: 曹 阳

责任印制: 何 芊

出版发行: 清华大学出版社

网 址: <http://www.tup.com.cn>, <http://www.wqbook.com>

地 址: 北京清华大学学研大厦 A 座

邮 编: 100084

社 总 机: 010-62770175

邮 购: 010-62786544

投稿与读者服务: 010-62776969, c-service@tup.tsinghua.edu.cn

质 量 反 馈: 010-62772015, zhiliang@tup.tsinghua.edu.cn

印 装 者: 清华大学印刷厂

经 销: 全国新华书店

开 本: 185mm×260mm

印 张: 28

字 数: 681 千字

版 次: 2014 年 7 月第 1 版

印 次: 2014 年 7 月第 1 次印刷

印 数: 1~4000

定 价: 59.80 元



译者序

自 2004 年左右 Google 发布 GFS 和 MapReduce 论文以来，伴随着国内移动通信和互联网的爆炸式发展，人们开始被信息所淹没，大数据开始推动各个领域发生深刻的变化，也产生了全新的机遇。Hadoop 生态系统作为 GFS 和 MapReduce 最成功的开源实现和演化，为我们分析和挖掘大数据提供了强大而易用的工具。

大数据如此重要，译者在工作中也不可避免地使用 Hadoop 等系统来存储用户数据，并在这些海量数据的基础上进行大量的运算和挖掘，以分析用户行为、支持产品决策。在学习或使用 Hadoop、GFS、MapReduce 等众多分布式存储/计算平台时，自己的一个最大体会就是中文资料太少，好的资料基本上都是英文资料。因而，看到本书英文版之后，就非常迫切地希望能将它翻译为中文，以便和国内同仁们一起分享。

本书的作者都是大数据和 Hadoop 领域的资深专家，有着多年的相关领域经验。本书各章提供了很多例子，除了一些入门级的例子之外，很多例子都来自实际项目，让我们有机会领略 Hadoop 在实际环境中的企业级应用。此外，本书紧跟最新技术，视野开阔，书中介绍的很多最新的项目，例如 Oozie、DSL 等都代表了 Hadoop 生态系统演化的最新流行方向，目前还很少有其他的中文出版物进行过系统阐述。这也是译者翻译并力荐本书的主要原因。

翻译的过程很有收获，但又颇为辛苦，译者在此过程中力求做到专业术语的准确、一致，对一些难以用中文描述的概念，也进行了反复推敲，力争做到准确而易于理解。尽管如此，由于时间仓促以及水平有限，错误之处在所难免，敬请广大读者批评指正。

参与本书翻译工作的还有李露熙、张宇欣和龙伟。

当人们都在讨论大数据的时候，真正理解大数据的精髓，并掌握 Hadoop 中一系列的强大工具来驾驭大数据，是一件非常快乐而有意义的事情。希望读者通过阅读本书能早日步入 Hadoop 的殿堂，领略大数据之美！

译者

2014 年 5 月于北京

作者简介

Boris Lublinsky 是诺基亚的一名资深架构师，他在诺基亚积极参与了无数聚焦于技术架构、面向服务架构(Service-Oriented Architecture, SOA)和集成项目等企业级应用的所有阶段。他也是诺基亚架构委员会的一名活跃会员。Boris 是多家行业杂志中 80 多篇出版物的作者，同时也是 *Service-Oriented Architecture and Design Strategies* (Indianapolis: Wiley, 2008)一书的合著者。此外，他是 InfoQ 在 SOA 和大数据领域的一名编辑，并经常在行业会议上演讲。在过去的两年里，他参与了多个基于 Hadoop 和 Amazon Web Services (AWS) 应用的设计和实现。他现在是芝加哥地区 Hadoop 用户组的活跃成员、工作组织者和撰稿人。

Kevin T. Smith 是 Novetta Solutions 公司应用事业解决方案(ASM)部门的技术解决方案和推广总监，他提供策略性的技术领导并为客户开发创新性的、聚焦于数据的、高度安全的解决方案。他经常在技术会议上演讲，是无数 Web 服务、云计算、大数据和网络安全相关技术文章的作者。他撰写了很多技术书籍，包括 *Applied SOA: Service-Oriented Architecture and Design Strategies* (Indianapolis: Wiley, 2008)、*The Semantic Web: A Guide to the Future of XML, Web Services, and Knowledge Management* (Indianapolis: Wiley, 2003)、*Professional Portal Development with Open Source Tools* (Indianapolis: Wiley, 2004)、*More Java Pitfalls* (Indianapolis: Wiley, 2003)以及其他书籍。

Alexey Yakubovich是Hortonworks的一名系统架构师。他在Hadoop/Big Data环境中为不同的公司和项目工作了5年：PB量级存储、流程自动化、自然语言处理(Natural Language Processing, NLP)、来自移动设备的数据流的数据科学和社交媒体。更早些时候，他工作于SOA、Java 2企业版(J2EE)、分布式应用和代码生成等技术领域。他通过解决Hilbert第一问题的最后部分获得了数学博士学位。他是MDA OMG工作组的一名成员，并参与和出席美国芝加哥地区的Hadoop用户组。

技术编辑简介

Michael C. Daconta 是 Advanced Technology for InCadence Strategic Solutions (<http://www.incadencecorp.com>) 的副总裁，他目前为政府和商业客户指导多个高级技术项目。他是一位著名的作者、演讲者和专栏作家，撰写或与其他人共同撰写了 11 部技术书籍(涉及语义 Web、XML、XUL、Java、C++ 和 C 等主题)、无数的杂志文章和在线专栏文章。他还为 *Government Computer News* 撰写月度“Reality Check”专栏文章。他获得了纽约大学的计算机学士学位、Nova Southeastern 大学的计算机科学硕士学位。

Michael Segel 从事 IT 行业已经超过了 20 年。他从 2009 年起集中关注大数据领域，同时拥有 MapR 和 Cloudera 认证。Segel 创立了美国芝加哥地区 Hadoop 用户组，并活跃于美国芝加哥大数据社区。他获得了俄亥俄州立大学工程学院 CIS 专业的理学学士学位。在空闲时，他把时间花在遛狗上。

Ralph Perko 是美国西北太平洋国家实验室可视化分析组的一名软件架构师和工程师。他目前参与多个 Hadoop 项目，帮助发明、开发用于处理和可视化大规模数据集的新方法。他拥有 16 年的软件开发经验。

致 谢

感谢和我一起工作的许多同事。他们总是督促我发挥到极限、询问我的解决方案，同时也让他们自己进步。本书采纳了他们的很多想法，希望能因此而更好。

非常感谢 Dean Wampler，他为第 13 章提供了 Hadoop 特定领域语言的内容。

感谢本书的合著者 Kevin 和 Alexey，他们把自己的 Hadoop 视角带到本书中，使得本书的内容更好、更平衡。

感谢技术编辑们对本书内容的很多有价值的建议。

感谢 Wiley 的编辑人员 Bob Elliott 决定出版本书；感谢我们的编辑 Kevin Shafer 不懈地工作来修正任何不一致之处；还感谢很多其他人，他们创造了您这时正在阅读的最终产品。

——Boris Lublinsky

首先，我想感谢本书的合著者以及 Wrox 出版社伟大团队付出的努力。这是我的第 7 个出书项目，我想要特别感谢我的好妻子 Gwen，以及我亲爱的孩子们 Isabella、Emma 和 Henry，他们容忍我很多个日日夜夜和周末一直都呆在自己的电脑前。

真诚地感谢阅读这些章节早期草稿的很多人，他们提供了很多建议、编辑、深入的观察和想法——特别是 Mike Daconta、Ralph Perko、Praveena Raavichara、Frank Tyler 和 Brian Uri。我对自己的公司 Novetta Solutions 心存感激，尤其想感谢 Joe Pantella 和 Novetta 执行团队的其他成员，他们支持我撰写本书。

有几件事情可能已经影响了我今年的新书撰写。有一段时间，华盛顿 Redskins(译者注：一支橄榄球队)看起来势不可挡，如果他们确实能进入“超级碗”(译者注：美国职业橄榄球队冠军赛)，这将会危及本书的截止时间。然而，Redskins 的赛季和另一个我参与其中的赛季被突然结束了。因而，我们要为本书的按时完成而感谢 Redskins 球队和 POJ 球队。

最后，特别感谢 CiR、Manera，以及过去、现在和未来的那个人。

——Kevin T. Smith

感谢诺基亚的同事们，撰写本书时我还在那里工作，他们的建议和知识为我撰写的几章提供了非常专业的素材。

感谢本书的合著者 Boris 和 Kevin，他们使本书以及我的参与成为可能。

感谢 Wiley 的编辑人员出版本书，他们提供了所有必要的帮助和指导来使本书变得更 valuable。

——Alexey Yakubovich

前言

在这个技术不停地改变的快节奏世界中，我们已经被信息淹没。我们在不断生成和存储大量的数据。随着网络上的设备不断丰富，我们已经看到了信息格式和数据多样性的惊人增长——大数据。

但是我们要直面它——如果我们忠实于自己的话，多数组织还不能积极有效地管理大量的数据，我们也还不能使这些信息发挥优势，从而更好地做出决策，更聪明地做生意。我们已经被大批量的数据所淹没，但同时我们又渴求知识。这会导致公司损失生产力、损失机会，并损失收入。

在过去的10年中，很多技术承诺帮助处理和分析我们拥有的大批量信息，而这些技术多数都出现了不足。我们知道这一点，因为作为专注于数据的程序员，我们已经都尝试过了。很多方法都是受专利保护的，导致供应商被锁定。一些方法看起来很有希望，但无法扩展以处理大型数据集，还有很多是宣传很好但不能满足预期，或者在关键时刻还没有准备好。

然而，**Apache Hadoop** 登场之后，一切都不一样了。当然这里也有宣传的因素，但它是一个开源项目，已经在大规模可扩展商业应用中取得了不可思议的成功。尽管学习曲线有些陡峭，但这是我们第一次可以轻松地编写程序并对大规模数据进行分析——以一种以前我们做不到的方式。**MapReduce** 算法可以使开发人员处理分布在可扩展机器集群上的数据，基于该算法，我们以过去无法做到的方式，在进行复杂数据分析方面取得了很大的成功。

关于 **Hadoop** 的书并不缺少。人们写了很多，而且其中很多书都很好。那么，为什么还要编写本书呢？当我们开始使用 **Hadoop** 时，希望有一本书不只是介绍 **API**，还要解释 **Hadoop** 生态系统的诸多部分如何共同工作，并可用于构建企业级的解决方案。我们在寻找这样一本书，它可以带领读者学习数据设计及其对实现的影响，并解释 **MapReduce** 的工作原理，以及如何用 **MapReduce** 来重新表述特定的业务问题。我们在寻找如下问题的答案：

- **MapReduce** 的强项和弱点是什么，以及我们如何自定义它以便更好地满足自己的需求？
- 为什么我们需要在 **MapReduce** 之上有一个额外的协调层，以及 **Oozie** 是怎么满足这个需求的？
- 我们如何使用特定领域语言(Domain-Specific Language, DSL)来简化 **MapReduce** 开发？

- 每个人都在讲的实时 Hadoop 是什么，它可以做什么，以及它不能做什么？它的工作原理是什么？
- 我们如何确保 Hadoop 应用程序的安全，我们需要考虑什么，我们必须考虑什么安全隐患，以及处理这些问题有哪些方法？
- 我们如何将自己的 Hadoop 应用程序迁移到云中，以及这样做有哪些重要的考虑因素？

当开始 Hadoop 探险时，我们不得不夜以继日地浏览整个 Internet 和 Hadoop 源代码，与人们交谈并尝试使用这些代码来找到这些问题的答案。然后我们决定通过撰写本书来分享自己的发现和经历，并希望本书能够成为读者理解和使用 Hadoop 的一个良好的开端。

本书读者对象

本书是由程序员写给程序员的。本书的作者是开发企业级解决方案的技术人员，我们对于本书的目标是为使用 Hadoop 的其他开发人员提供可靠的、实用的建议。本书的目标人群是试图更好地理解 and 利用 Hadoop——不只是做简单的数据分析，同时也将 Hadoop 用作企业级应用的基础——的软件架构师和开发人员。

因为 Hadoop 是一个基于 Java 的框架，所以本书包含了大量的代码示例，需要读者熟悉 Java。此外，作者假设读者在一定程度上熟悉 Hadoop，并有一些入门级的 MapReduce 知识。

尽管本书在设计上希望读者从头到尾以逐个模块的方式阅读，但某些章节可能更适合特定的人群。想要理解 Hadoop 数据存储能力的数据设计人员更可能从第 2 章中受益。从 MapReduce 开始的程序员们最有可能关注第 3~第 5 章，以及第 13 章。意识到不使用像 Oozie 这样的 Workflow 系统会导致复杂性的开发人员，最可能关注第 6~第 8 章。那些对实时 Hadoop 感兴趣的人们会关注第 9 章。有兴趣在实现中使用亚马逊云的人们可能会关注第 11 章，而在乎安全性的人们可能关注第 10 章和第 12 章。

本书涵盖的内容

现在，每个人都在处理大数据。很多企业正在最大限度地实施大规模可扩展性数据分析工作，其中多数企业都在尝试使用 Hadoop 来完成该目标。本书集中讲述构建基于 Hadoop 高级企业级应用的架构和方法，并为此涵盖了如下主要 Hadoop 组件：

- 基于 Hadoop 的企业级应用的架构蓝图
- 基础的 Hadoop 数据存储和组织系统
- Hadoop 的主要执行框架(MapReduce)
- Hadoop 的 Workflow/Coordinator 服务器(Oozie)
- 实现基于 Hadoop 的实时系统的技术

- 在云环境中运行 Hadoop 的方式
- 确保 Hadoop 应用安全的技术和架构

本书的组织结构

本书被编排为 13 章。

第 1 章(“大数据和 Hadoop 生态系统”)介绍了大数据,以及 Hadoop 用作大数据实现的方法。在该章中,我们学习 Hadoop 如何解决大数据带来的挑战,以及哪些 Hadoop 核心组件可以共同工作来创建丰富的 Hadoop 生态系统,适合解决很多现实世界的问题。我们也学习了多个可用的 Hadoop 发行版,以及新出现的用于大数据应用的架构模式。

任何大数据实现的基础都是数据存储设计。第 2 章(“Hadoop 数据存储”)涵盖了 Hadoop 支持的分布式数据存储。该章讨论了两个主要的 Hadoop 数据存储机制——HDFS 和 HBase——的架构和 API,并提供了何时使用哪种机制的一些建议。这里我们学习了 HDFS(联盟)和 HBase 新文件格式以及协处理器的最新发展。该章也涵盖了 HCatalog(Hadoop 元数据管理解决方案)和 Avro(一个序列化/组装框架),以及它们在 Hadoop 数据存储中扮演的角色。

作为主要的 Hadoop 执行框架,MapReduce 是本书的主要议题之一,包含在第 3~第 5 章中。

第 3 章(“使用 MapReduce 处理数据”)介绍了 MapReduce 框架。涵盖了 MapReduce 架构、它的主要组件和 MapReduce 编程模型。该章也重点介绍了 MapReduce 应用的设计、设计模式以及 MapReduce 的一般注意事项。

第 4 章(“自定义 MapReduce 执行”)建立在第 3 章的基础之上,涵盖了自定义 MapReduce 执行的重要方法。我们学习 MapReduce 执行可以被自定义的一些方面,并使用工作代码示例来揭示如何做到这一点。

最后,在第 5 章(“构建可靠的 MapReduce 应用程序”)中,我们学习了构建可靠的 MapReduce 应用程序的途径,包括测试和调试,以及使用内置的 MapReduce 工具(例如,日志和计数器)来查看 MapReduce 的内部执行。

尽管 MapReduce 自身具有强大的功能,但实际的解决方案通常需要将多个 MapReduce 应用组合到一起,这涉及很多的复杂性。通过使用 Hadoop Workflow/Coordinator 引擎——Oozie——可以显著地简化这种复杂性,Oozie 将在第 6~第 8 章中描述。

第 6 章(“使用 Oozie 自动化数据处理”)介绍了 Oozie。这里我们学习 Oozie 的整体架构、它的主要组件,以及每个组件的编程语言。我们也学习 Oozie 的整体执行模型,以及可以和 Oozie 服务器交互的方式。

第 7 章(“使用 Oozie”)建立在第 6 章所学知识的基础之上并展示了一个实用的从无到有使用 Oozie 开发实际应用的示例。该示例演示了如何在解决方案中使用不同的 Oozie 组件,并演示了设计和实现途径。

最后,第 8 章(“高级 Oozie 特性”)讨论一些高级特性,并展示了扩展 Oozie 和将它与

其他企业级应用集成的方法。在该章中，我们学习了一些开发人员需要知道的建议和技巧——例如，Oozie 代码的动态生成如何帮助开发人员克服一些现存的、任何其他方法都不能解决的 Oozie 缺点。

当今与大数据相关的最热门的趋势之一是进行“实时分析”的能力。该主题在第 9 章（“实时 Hadoop”）中讨论。该章开始提供了一些目前使用的实时 Hadoop 应用示例，并展示了这些实现的整体架构需求。我们将学习构建这样一些实现的 3 种主要方法——基于 HBase 的应用程序、实时查询和基于流的处理。

该章提供了两个基于 HBase 的实时应用——一个假想的图片管理系统和一个基于 Lucene、使用 HBase 作为后端的搜索引擎。我们也会学习实现实时查询的整体架构，以及两个具体的产品——Apache Drill 和 Cloudera 的 Impala——实现实时查询的方式。该章同时涵盖了另一种类型的实时应用——复杂事件处理，包括它的整体架构，以及 HFlume 和 Storm 实现该架构的方式。最后，该章提供了实时查询、复杂事件处理和 MapReduce 之间的一个对比。

Hadoop 应用程序开发中一个经常被忽略但至关重要的主题是 Hadoop 安全。第 10 章（“Hadoop 安全”）深入讨论了与大数据分析和 Hadoop——确切地说是 Hadoop 安全模型和最佳实践——相关的安全性考虑。这里我们学习 Rhino 项目——一个支持开发人员扩展 Hadoop 安全能力（包括加密、认证、授权、单点登录（Single-Sign-On, SSO）和审计）的框架。

基于云的 Hadoop 使用需要有趣的架构决策。第 11 章（“在 AWS 上运行 Hadoop 应用”）描述了这些挑战，并涵盖了在亚马逊 Web 服务（Amazon Web Service, AWS）云上运行 Hadoop 的不同方法。该章也讨论了一些权衡选择并考察了一些最佳实践。我们将学习弹性 MapReduce（Elastic MapReduce, EMR）和可以用于补充 Hadoop 功能的额外 AWS 服务（例如 S3、CloudWatch、Simple Workflow 等）。

除了确保 Hadoop 自身的安全，Hadoop 实现通常和其他企业级组件集成——数据经常被导入到 Hadoop 并被导出。第 12 章（“为 Hadoop 实现构建企业级安全解决方案”）涵盖了如何确保使用 Hadoop 的企业级应用尽可能安全，并提供了示例和最佳实践。

作为本书的最后一章，第 13 章（“Hadoop 的未来”）浏览了一些当前和未来将发生的 Hadoop 行业趋势和创新。这里我们学习了可用来简化 MapReduce 开发的 Hadoop DSL 及其使用，以及新的 MapReduce 资源管理系统（YARN）和 MapReduce 运行时扩展（Tez）。我们也学习了最重要的 Hadoop 发展方向和趋势。

使用本书需要的条件

本书中展示的所有代码都用 Java 实现。因此，要使用这些代码，读者需要 Java 编译器和开发环境。所有的开发都在 Eclipse 中完成，但由于每个项目都有一个 Maven pom 文件，因而把代码迁移到任何您所选择的开发环境都足够简单。

所有的数据访问和 MapReduce 代码都在 Hadoop 1（Cloudera CDH 3 发行版和 Amazon

EMR)和 Hadoop 2(Cloudera CDH 4 发行版)上测试通过。因此, 这些代码应当可以在任何 Hadoop 发行版上运行。Oozie 代码在最新版的 Oozie(例如, 可以从 Cloudera CDH 4.1 发行版获得)上测试通过。

示例的源代码都以 Eclipse 项目进行组织(每章一个项目), 并可以从 Wrox 网站 www.wrox.com/go/prohadoopsolutions 下载获得。

源代码

在学习本书中的示例时, 读者可以选择手工输入所有的代码, 也可以使用随书的源代码文件。本书的源代码可以从 www.wrox.com 下载获得。具体对本书而言, 专门的代码下载位于 www.wrox.com/go/prohadoopsolutions 的 Download Code 链接。

读者也可以在 www.wrox.com 按照本书英文版的 ISBN(本书英文版的 ISBN 是 978-1-118-61193-7)搜索本书来找到这些代码。当前所有的 Wrox 图书代码下载的完整列表位于 www.wrox.com/dynamic/books/download.aspx。

在整个选定的各章中, 读者也可以在需要时, 从代码清单的标题和文本中找到代码文件名称的引用。

www.wrox.com 上的大多数代码都压缩在 .ZIP、.RAR 文档或适合特定平台的类似文档中。一旦下载了这些代码, 只需要使用合适的压缩工具将其解压即可。

注意:

由于许多图书的标题都很类似, 因此按 ISBN 搜索是最简单的, 本书英文版的 ISBN 是 978-1-118-61193-7。

另外, 读者也可以通过网址 www.wrox.com/dynamic/books/download.aspx 进入 Wrox 代码下载主页面, 找到本书以及所有其他 Wrox 图书的可用代码。

勘误表

尽管我们已经尽了各种努力来保证文章或代码中不出现错误, 但是错误总是难免的, 如果您在本书中发现了错误, 例如拼写错误或代码错误, 请告诉我们, 我们将非常感激。通过勘误表, 可以让其他读者避免受挫, 当然, 这还有助于提供更高质量的信息。

要找到本书英文版的勘误表, 可以登录 <http://www.wrox.com/go/prohadoopsolutions>, 单击 Errata 链接。在这个页面上可以查看到 Wrox 编辑已提交和粘贴的所有勘误项。

如果您发现的错误在我们的勘误表里还没有出现的话, 请登录 www.wrox.com/contact/techsupport.shtml 并完成那里的表格, 把您发现的错误发送给我们。我们会检查您的反馈信息, 如果正确, 我们将在本书的勘误表页面张贴该错误消息, 并在本书的后续版本加以修订。

p2p.wrox.com

要与作者和同行讨论，请加入 p2p.wrox.com 上的 P2P 论坛。这个论坛是一个基于 Web 的系统，便于您张贴与 Wrox 图书相关的消息和相关技术，与其他读者和技术用户交流心得。该论坛提供了订阅功能，当您感兴趣的主题在论坛上有新帖子发布时，系统会向您发送电子邮件。Wrox 的作者、编辑和其他业界专家和读者都会到这个论坛上来探讨问题。

在 <http://p2p.wrox.com> 上，有许多不同的论坛，它们不仅有助于阅读本书，还有助于开发自己的应用程序。要加入论坛，可以遵循下面的步骤：

- (1) 进入 p2p.wrox.com，单击 Register 链接。
- (2) 阅读使用协议，并单击 Agree 按钮。
- (3) 填写加入该论坛所需要的信息和自己希望提供的其他可选信息，单击 Submit 按钮。
- (4) 您会收到一封电子邮件，其中的信息描述了如何验证账户，完成加入过程。

注意：

不加入 P2P 也可以阅读论坛上的消息，但要张贴自己的消息，就必须先加入该论坛。

加入论坛后，就可以张贴新消息，响应其他用户张贴的消息。可以随时在 Web 上阅读消息。如果要让该网站给自己发送特定论坛中的消息，可以单击论坛列表中该论坛名旁边的 Subscribe to this Forum 图标。

要想了解更多的有关论坛软件的工作情况，以及 P2P 和 Wrox 图书的许多常见问题的解答，就一定要阅读 FAQ，只需要在任意 P2P 页面上单击 FAQ 链接即可。

目 录

第 1 章 大数据和 Hadoop 生态系统1	第 3 章 使用 MapReduce 处理数据 55
1.1 当大数据遇见 Hadoop.....2	3.1 了解 MapReduce55
1.1.1 Hadoop: 直面大 数据的挑战3	3.1.1 MapReduce 执行管道 56
1.1.2 商业世界中的数据科学4	3.1.2 MapReduce 中的运行时 协调和任务管理 59
1.2 Hadoop 生态系统.....6	3.2 第一个 MapReduce 应用程序.....61
1.3 Hadoop 核心组件.....7	3.3 设计 MapReduce 实现69
1.4 Hadoop 发行版.....9	3.3.1 将 MapReduce 用作 并行处理框架 70
1.5 使用 Hadoop 开发 企业级应用.....10	3.3.2 使用 MapReduce 进行 简单的数据处理 71
1.6 小结14	3.3.3 使用 MapReduce 构建 连接 72
第 2 章 Hadoop 数据存储 15	3.3.4 构建迭代式 MapReduce 应用程序 77
2.1 HDFS 15	3.3.5 是否使用 MapReduce 82
2.1.1 HDFS 架构 15	3.3.6 常见的 MapReduce 设计陷阱 83
2.1.2 使用 HDFS 文件 19	3.4 小结84
2.1.3 Hadoop 特定的文件类型 21	第 4 章 自定义 MapReduce 执行 85
2.1.4 HDFS 联盟和高可用性 26	4.1 使用 InputFormat 控制 MapReduce 执行 85
2.2 HBase.....28	4.1.1 为计算密集型应用 程序实现 InputFormat 87
2.2.1 HBase 架构.....28	4.1.2 实现 InputFormat 以 控制 Map 的数量 93
2.2.2 HBase 结构设计.....34	4.1.3 实现用于多个 HBase 表 的 InputFormat 99
2.2.3 HBase 编程.....35	4.2 使用自定义 RecordReader 以自己的方式读取数据 102
2.2.4 HBase 新特性.....42	
2.3 将 HDFS 和 HBase 的组合 用于高效数据存储.....45	
2.4 使用 Apache Avro.....45	
2.5 利用 HCatalog 管理元数据49	
2.6 为应用程序选择合适的 Hadoop 数据组织形式51	
2.7 小结53	

4.2.1 实现基于队列的 RecordReader.....	102	6.2.2 Oozie 的恢复能力.....	164
4.2.2 为 XML 数据实现 RecordReader.....	105	6.2.3 Oozie Workflow 作业的 生命周期.....	164
4.3 使用自定义输出格式 组织输出数据.....	109	6.3 Oozie Coordinator.....	165
4.4 使用自定义记录写入器以 自己的方式写入数据.....	119	6.4 Oozie Bundle	170
4.5 使用组合器优化 MapReduce 执行	121	6.5 用表达式语言对 Oozie 进行参数化.....	174
4.6 使用分区器控制 Reducer 执行.....	124	6.5.1 Workflow 函数.....	175
4.7 在 Hadoop 中使用非 Java 代码.....	128	6.5.2 Coordinator 函数.....	175
4.7.1 Pipes.....	128	6.5.3 Bundle 函数.....	175
4.7.2 Hadoop Streaming	128	6.5.4 其他 EL 函数	175
4.7.3 使用 JNI.....	129	6.6 Oozie 作业执行模型	176
4.8 小结	131	6.7 访问 Oozie	179
第 5 章 构建可靠的 MapReduce 应用程序	133	6.8 Oozie SLA	180
5.1 单元测试 MapReduce 应用程序.....	133	6.9 小结.....	185
5.1.1 测试 Mapper.....	136	第 7 章 使用 Oozie	187
5.1.2 测试 Reducer.....	137	7.1 使用探测包验证位置 相关信息正确性.....	187
5.1.3 集成测试	138	7.2 设计基于探测包的地点 正确性验证.....	188
5.2 使用 Eclipse 进行本地 应用程序测试.....	139	7.3 设计 Oozie Workflow	190
5.3 将日志用于 Hadoop 测试.....	141	7.4 实现 Oozie Workflow 应用程序.....	193
5.4 使用作业计数器报告指标.....	146	7.4.1 实现数据准备 Workflow.....	193
5.5 MapReduce 中的防御性 编程.....	149	7.4.2 实现考勤指数和聚类 探测包串 Workflow.....	201
5.6 小结	151	7.5 实现 Workflow 行为.....	203
第 6 章 使用 Oozie 自动化数据 处理.....	153	7.5.1 发布来自 java 动作的 执行上下文	204
6.1 认识 Oozie.....	154	7.5.2 在 Oozie Workflow 中 使用 MapReduce 作业	204
6.2 Oozie Workflow	155	7.6 实现 Oozie Coordinator 应用程序.....	207
6.2.1 在 Oozie Workflow 中 执行异步操作	159	7.7 实现 Oozie Bundle 应用程序.....	212
		7.8 部署、测试和执行 Oozie 应用程序.....	213
		7.8.1 部署 Oozie 应用程序.....	213

7.8.2 使用 Oozie CLI 执行 Oozie 应用程序.....	215	9.3 使用专门的实时 Hadoop 查询系统.....	295
7.8.3 向 Oozie 作业传递参数.....	218	9.3.1 Apache Drill.....	296
7.9 使用 Oozie 控制台获取 Oozie 应用程序信息.....	221	9.3.2 Impala.....	298
7.9.1 了解 Oozie 控制台界面.....	221	9.3.3 实时查询和 MapReduce 的对比.....	299
7.9.2 获取 Coordinator 作业信息.....	225	9.4 使用基于 Hadoop 的事件 处理系统.....	300
7.10 小结.....	227	9.4.1 HFlame.....	301
第 8 章 高级 Oozie 特性.....	229	9.4.2 Storm.....	302
8.1 构建自定义 Oozie Workflow 动作.....	230	9.4.3 事件处理和 MapReduce 的对比.....	305
8.1.1 实现自定义 Oozie Workflow 动作.....	230	9.5 小结.....	305
8.1.2 部署 Oozie 自定义 Workflow 动作.....	235	第 10 章 Hadoop 安全.....	307
8.2 向 Oozie Workflow 添加 动态执行.....	237	10.1 简要的历史: 理解 Hadoop 安全的挑战.....	308
8.2.1 总体实现方法.....	237	10.2 认证.....	309
8.2.2 一个机器学习模型、参数 和算法.....	240	10.2.1 Kerberos 认证.....	310
8.2.3 为迭代过程定义 Workflow.....	241	10.2.2 委派安全凭据.....	318
8.2.4 动态 Workflow 生成.....	244	10.3 授权.....	323
8.3 使用 Oozie Java API.....	247	10.3.1 HDFS 文件访问权限.....	323
8.4 在 Oozie 应用中使用 uber jar 包.....	251	10.3.2 服务级授权.....	327
8.5 数据吸收传送器.....	256	10.3.3 作业授权.....	329
8.6 小结.....	263	10.4 Oozie 认证和授权.....	329
第 9 章 实时 Hadoop.....	265	10.5 网络加密.....	331
9.1 现实世界中的实时应用.....	266	10.6 使用 Rhino 项目增强 安全性.....	332
9.2 使用 HBase 来实现 实时应用.....	266	10.6.1 HDFS 磁盘级加密.....	333
9.2.1 将 HBase 用作图片 管理系统.....	268	10.6.2 基于令牌的认证和 统一的授权框架.....	333
9.2.2 将 HBase 用作 Lucene 后端.....	275	10.6.3 HBase 单元格级安全.....	334
		10.7 将所有内容整合起来——保 证 Hadoop 安全的最佳实践.....	334
		10.7.1 认证.....	335
		10.7.2 授权.....	335
		10.7.3 网络加密.....	336
		10.7.4 敬请关注 Hadoop 的 增强功能.....	336

10.8 小结	336	12.1.1 认证	380
第 11 章 在 AWS 上运行 Hadoop		12.1.2 授权	380
应用	337	12.1.3 保密性	380
11.1 初识 AWS	338	12.1.4 完整性	381
11.2 在 AWS 上运行 Hadoop 的 可选项	339	12.1.5 审计	381
11.2.1 使用 EC2 实例的 自定义安装	339	12.2 Hadoop 安全没有为企业级 应用原生地提供哪些机制	381
11.2.2 弹性 MapReduce	339	12.2.1 面向数据的访问控制	382
11.2.3 做出选择前的额外 考虑	339	12.2.2 差分隐私	382
11.3 理解 EMR-Hadoop 的关系	340	12.2.3 加密静止的数据	383
11.3.1 EMR 架构	341	12.2.4 企业级安全集成	384
11.3.2 使用 S3 存储	343	12.3 保证使用 Hadoop 的企业 级应用安全的方法	384
11.3.3 最大化 EMR 的使用	343	12.3.1 使用 Accumulo 进行 访问控制保护	385
11.3.4 利用 CloudWatch 和 其他 AWS 组件	345	12.3.2 加密静止数据	394
11.3.5 访问和使用 EMR	346	12.3.3 网络隔离和分隔方案	395
11.4 使用 AWS S3	351	12.4 小结	397
11.4.1 理解桶的使用	352	第 3 章 Hadoop 的未来	399
11.4.2 使用控制台浏览内容	354	13.1 使用 DSL 简化 MapReduce 编程	400
11.4.3 在 S3 中编程访问文件	355	13.1.1 什么是 DSL	400
11.4.4 使用 MapReduce 上传 多个文件到 S3	365	13.1.2 Hadoop 的 DSL	401
11.5 自动化 EMR 作业流 创建和作业执行	367	13.2 更快、更可扩展的数据 处理	412
11.6 管理 EMR 中的作业执行	372	13.2.1 Apache YARN	412
11.6.1 在 EMR 集群上使用 Oozie	372	13.2.2 Tez	414
11.6.2 AWS 简单工作流	374	13.3 安全性的改进	415
11.6.3 AWS 数据管道	375	13.4 正在出现的趋势	415
11.7 小结	376	13.5 小结	416
第 12 章 为 Hadoop 实现构建 企业级安全解决方案	377	附录 有用的阅读	417
12.1 企业级应用的安全顾虑	378		