

云计算与大数据实验教材系列

Hadoop

核心技术与实验

主编 饶文碧 袁景凌 张露 熊盛武 刘荣英



WUHAN UNIVERSITY PRESS

武汉大学出版社



Hadoop

核心技术与实验

内容提要

本书是一本学习Hadoop数据处理技术的入门级实验教材，包括Hadoop知识点、关键技术和实验案例等内容。知识点部分对学生在Hadoop应用实践中碰到的概念、模型和机制进行简要的介绍；关键技术与经典实验案例相结合，帮助学生通过实践掌握研发的关键技术。

本书最大的特点是把知识点、关键技术和实验案例相结合，在对Hadoop技术体系进行描述基础上，还为主要知识点设计了经典的小案例，易于理解，可操作性强。

- 责任编辑 / 张 欣
- 责任校对 / 李孟潇
- 版式设计 / 马 佳
- 封面设计 / 罗 珏

ISBN 978-7-307-12840-8



9 787307 128408 >

定价: 36.00元

云计算与大数据实验教材系列

Hadoop

核心技术与实验

主编 饶文碧 袁景凌 张露 熊盛武 刘荣英



WUHAN UNIVERSITY PRESS

武汉大学出版社

图书在版编目(CIP)数据

Hadoop 核心技术与实验/饶文碧等主编. —武汉: 武汉大学出版社, 2017.4

云计算与大数据实验教材系列

ISBN 978-7-307-12840-8

I. H… II. 饶… III. 数据处理软件—高等学校—教材 IV. TP274

中国版本图书馆 CIP 数据核字(2017)第 025309 号

责任编辑:张欣

责任校对:李孟潇

版式设计:马佳

出版发行: 武汉大学出版社 (430072 武昌 珞珈山)

(电子邮件: cbs22@whu.edu.cn 网址: www.wdp.com.cn)

印刷:湖北恒泰印务有限公司

开本: 787×1092 1/16 印张:14.75 字数:348千字 插页:1

版次: 2017年4月第1版 2017年4月第1次印刷

ISBN 978-7-307-12840-8 定价:36.00元

版权所有, 不得翻印;凡购我社的图书, 如有质量问题, 请与当地图书销售部门联系调换。

前 言

Hadoop 是 Apache 软件基金会旗下的一个开源分布式计算平台，以 Hadoop 分布式文件系统 HDFS 和 MapReduce 为核心的 Hadoop 为用户提供了系统底层细节透明的分布式基础架构，用户可以在不了解分布式底层细节的情况下，开发分布式程序，充分利用集群的威力进行高速运算和存储。Hadoop 设计之初的目标定位于高可靠性、高可扩展性、高容错性和高效性，正是这些设计上与生俱来的优点，才使得 Hadoop 一出现就受到众多大公司的青睐，同时也引起了研究界的普遍关注。到目前为止，Hadoop 技术在互联网领域已经得到了广泛的运用，越来越多的企业将 Hadoop 技术作为进入大数据领域的必备技术。

本书目的：

由于 Hadoop 拥有可计量、成本低、高效、可信等突出特点，基于 Hadoop 的应用已经遍地开花，尤其是在互联网领域。本书目的在于帮助初学者和有一定 Hadoop 基础但是缺乏实践经验的读者理解 Hadoop，并用它解决大数据问题。

本书针对初学者既需要系统地了解 Hadoop 整个技术体系，又需要通过实践掌握研发的关键技术的特点，在撰写上采用了理论与案例相结合的方式，在对 Hadoop 技术体系进行较为全面的描述基础上，还为主要知识点设计了经典的小案例，易于理解，可操作性强。

本书内容：

第 1 章：Hadoop 基础，主要包括 Hadoop 简介、Hadoop 的安装与配置和 Hadoop 的常用插件。

第 2 章：MapReduce 开发，主要包括 MapReduce 计算模型、如何在 Hadoop 中开发 MapReduce 的应用程序、MapReduce 应用案例和 MapReduce 工作机制。

第 3 章：Hadoop 进阶，主要包括 Hadoop I/O 操作、Hadoop 集群管理和下一代 Hadoop MapReduce 框架的设计细节。

第 4 章：Hadoop 实战，通过两个经典案例的剖析让读者深入了解开发 Hadoop 程序的步骤和思路。

本书特点：

本书知识点的讲解都配有相应的案例，以加深读者的理解；书中案例值得借鉴，还保留了部分源码，初学者可以根据这些案例动手实践。

本书目标读者：

本书理论知识与项目实战相结合，适合 Hadoop 初学者和有一定 Hadoop 基础但是缺乏

实践经验的读者阅读，也适合作为高等院校相关课程的教学参考书。

本书由饶文碧、袁景凌、张露、熊盛武、刘荣英撰写，江泉、谭永强、李梦璇硕士参与了实验调试及相关文档的撰写。

编 者

2017 年 1 月

序 言 二

民主恳谈是浙江省温岭市具有原创性的基层民主形式。但是,创建民主恳谈的初衷并不是为了构建一种新型的民主形式,而是探索新形势下如何加强和改进农村基层的思想政治工作。后来,随着工作的不断推进,民主恳谈会开始深化,转到了基层民主政治建设的层面,成为中国基层协商民主的最初形态。温岭民主恳谈会不仅于2004年3月获得了“中国地方政府创新奖”,而且受到国内外媒体和学术界的广泛关注,可以说是改革开放以来受到最多关注和研究最多的案例之一。

我本人于2005年开始介入温岭民主恳谈会活动,不仅作为观察者进行观察,而且参与了该市泽国镇民主恳谈会相关制度设计、组织工作等。但是,同温岭市委党校朱圣明先生相比较,我只能算是半个亲历者。在十多年前,我结识了温岭市委党校的朱圣明先生。在我所接触到的地方党校的教员中,像朱圣明先生那样执着于学术研究,将自己的研究建立在自己家乡丰厚的地方实验基础上,并且十年来不弃不舍,着实不多见。如果不主动介绍,不听其发言,不看其论著,很容易将朱先生当作当地的一位普通农民。但是,每次参加温岭民主恳谈会活动,尤其泽国恳谈会活动后的小规模的学术研讨,朱先生都给我留下了深刻的印象:一是他的大嗓门,发言时音量大;二是他观察仔细,提出或分析问题非常到位。因此,我每次到温岭总是非常期待朱先生的出现。现在呈现在我面前的《民主恳谈:中国基层协商民主的温岭实践》,就是他十年来的观察、思考和分析的结果,可谓水到渠成,而由他对温岭实践作系统的梳理与分析是恰当的。他邀请我为他的书写个序。这个邀请为难我,因为我从没有给人写过序,也从来没有想起为某一书写序。但是,恭敬不如从命,勉强为之。

本书共六章,其结构可分为两个部分,第一部分即本书的第一章,属于概述性的文字,作者对温岭民主恳谈会的创新意义,为何能够创新,民主恳谈会的演变历程等作了概括,让读者有一个概貌性的了解。第二部分是本书的主体,即从第二章到第六章,作者将温岭民主恳谈会的内容分为五类,并且给予详细的讨论与分析,当然在这五大类型中,作者把重点放在了参与式预算上。较之其他学者,得生活和工作在温岭的便利条件,朱先生提供了更为具体、丰富的经验材料和数据。

同时,因其本地人的便利,朱先生现场经历了各种类型民主恳谈的完整过程,他的研究与仅听转述者回忆的研究不同,融入了许多现场感,这是非外人所能够了解和体会到的。

从比较的视角来看,无论较之亚洲还是北欧,中国协商民主的建构与进展可以用“神速”一词来形容,在中共十八大后短短三年多时间里,就构建出中国协商民主的体系。我们很难在西方国家中见到类似这样从立法机构到政府、到政党、再到社会的全方位、多维度的协商体系。但是,我们通过《民主恳谈:中国基层协商民主的温岭实践》发现,这种全方位、多维度的协商体系在温岭的实践已有年头了。以温岭“民主恳谈会”为代表的中国基层、地方公民协商等协商民主机制,已经构成了协商民主经验的重要来源,必将进一步丰富世界协商民主的实践,必将为世界民主政治的发展做出中国的贡献,同时也必将有助于国家治理体系和治理能力现代化建设的推进。当然,这些富有特色与成果的协商民主,需要我们从中国自己的情/语境(context)加以理解。在中国,是什么条件或因素可以进一步推进中国协商民主发展、深化,协商民主蕴涵着怎样的政治、社会意义,依然是研究中国协商民主政治的一些重要议题。我相信《民主恳谈:中国基层协商民主的温岭实践》的出版能够激发出相关的研讨。此外,我本人也期待着朱先生能够进一步提炼温岭协商民主的经验,放在比较的维度,在理论上作出同国际学术界的对话,丰富世界协商民主理论的研究。

浙江大学政治学系系主任 郎友兴

2016年10月2日

前言

协商民主是当代中国政治发展的一个重要内容,也是中国治理现代化的一个重要组成部分。基层协商民主是社会主义协商民主最活跃的实践形式,地方党委、政府在日常公共事务决策中与社会的沟通衔接构成了基层协商民主的重要内容。温岭民主恳谈是中国基层协商民主的典范,经过十多年孜孜不倦的艰苦探索,现已创新和发展出对话型民主恳谈、决策型民主恳谈、党内民主恳谈、参与式预算、工资集体协商五种类型的民主恳谈。

对话型民主恳谈一般发生在决策之前,是民主恳谈会的组织者与参与者之间就某项社会公众广泛关注的议题进行对话、沟通、讨论、协商以求得共识,从而落实人民群众的知情权、参与权、表达权和监督权。“对话型民主恳谈可在市级、镇(街道)、村(社区)的党组织、人大、政府、政协、社会团体、事业企业单位、群众自治组织以及其他各社会组织中开展,起收集民意、沟通信息、协商问题、协调矛盾的作用。”^①

决策型民主恳谈更多地与公共政策制定和公共事务相联系,譬如城镇规划、村庄整治、校网调整、公共自行车站点的布局、公交车线路的安排等等。在相互联系、相互依存的社会中,人们的活动具有“公共性”,事关“公共性”的决策最容易形成聚焦,又具有争议性,因此决策型民主恳谈是重点。

在中国现有的政治生态中,中国共产党作为执政党,其组织的触角遍及社会的每一个角落。所以,通过党内民主恳谈可以把社会舆情反映上来,把党代表的智慧和积极性充分发挥出来。党内民主恳谈的实施,推进了党务公开,畅通了党内诉求渠道,保障了党员的主体地位,提高了党委决策的民主化、科学化和制度化水平。

参与式预算是民主恳谈深化和发展的产物,是指公民以民主恳谈为主要形式广泛参与政府年度预算方案协商讨论,人大代表审查政府财政预算并决定预算的修正和调整,进而实现实质性参与和监督政府预算执行的协商民主实践。参与式预算把“钱袋子”化作具体的分配、建设和民生等问题,让老百姓参与进来,使预算

① 引自中共温岭市委《关于全面深化民主恳谈推进协商民主制度化发展的意见》(2013年4月28日)。

结构的安排更加合理。

温岭的工资集体协商同样以民主恳谈这一独特的载体为平台,2003 年从新河羊毛衫行业开始,经过多年坚持不懈的推广与完善,通过企业、村(居、社区)、行业、区域性工资集体协商制度建设,形成了四轮联动,横向到边、纵向到底,“行业谈标准、区域谈底线、企业谈增幅”的工资集体协商新模式。为劳动者和企业双方提供了一种制度化、组织化的共决途径,使劳资冲突中的核心——劳动定额(工价)问题得到了较好的解决。

1999 年是民主恳谈的创生之年,这一年中共温岭市委以论坛的形式为老百姓搭建了一个诉求表达的“平台”;2001 年是民主恳谈正式冠名的一年,“恳谈”两字准确地表达了有意见需要沟通、有想法需要交流、有问题需要协商、有分歧需要博弈等丰富的意思;2002 年是民主恳谈转型的一年,这一年恳谈的议题从一家一户的具体问题转入公共事务;2003、2004 年是民主恳谈从发生到巩固的两年——2003 年工资集体协商起步,2004 年党内民主恳谈开局;2005 年至今,以参与式预算为标志,民主恳谈进入了完善、深化和提高的新阶段——2008 年的部门预算民主恳谈,2011 年通过人大代表工作站的预算征询恳谈,2013 年的市福利中心项目选址民主听证会和票决部门预算,2014 年的选民议政会和预算绩效评价,2015 年市级人代会上的预算修正议案。十多年来,践行者栉风沐雨、砥砺前行,民主恳谈精彩不断、亮点纷呈。随着议题的不断拓展,范围的不断扩大,方法的不断创新,重大公共事务决策经过民主恳谈,已成为当地的一个“规定动作”和“前置条件”。

民主恳谈寓于民主治理的形成过程中,是党的群众路线的制度化和程序化,体现了社会权利或公民权利空间的扩展过程。民主恳谈极大地丰富了基层协商民主的想象力,体现了人民群众政治参与的真实性、广泛性和有效性,必将载入中国民主政治建设的发展史册。

目 录

第 1 章 Hadoop 基础	1
1.1 Hadoop 简介	1
1.1.1 Hadoop 概述	1
1.1.2 Hadoop 项目及其结构	3
1.1.3 Hadoop 体系结构	5
1.1.4 Hadoop 数据管理	7
1.1.5 Hadoop 分布式开发	8
1.1.6 Hadoop 计算模型-MapReduce	11
1.1.7 Hadoop 集群安全策略	11
1.2 Hadoop 的安装与配置	13
1.2.1 在 Linux 上安装与配置 Hadoop	14
1.2.2 在 Windows 上安装与配置 Hadoop	20
1.2.3 安装和配置 Hadoop 集群	26
1.3 Hadoop 的常用插件与开发	40
1.3.1 Hadoop Eclipse 的安装环境	40
1.3.2 Hadoop Eclipse 的编译步骤	40
1.3.3 Hadoop Eclipse 的安装步骤	44
1.3.4 Hadoop Eclipse 的使用举例	46
1.4 思考题	47
第 2 章 MapReduce 开发	49
2.1 MapReduce 计算模型	49
2.1.1 MapReduce 计算模型	49
2.1.2 MapReduce 任务的优化	62
2.1.3 Hadoop 流	64
2.1.4 Hadoop Pipes	71
2.2 开发 MapReduce 应用程序	75
2.2.1 系统参数的配置	75
2.2.2 配置开发环境	77
2.2.3 编写 MapReduce 程序	78
2.2.4 本地测试	82

2.2.5 网络用户界面	83
2.2.6 性能调优	86
2.2.7 MapReduce 工作流	90
2.3 MapReduce 应用案例	95
2.3.1 单词计数	96
2.3.2 数据去重	103
2.3.3 排序	106
2.3.4 单表关联	110
2.3.5 多表关联	114
2.4 MapReduce 工作机制	119
2.4.1 MapReduce 作业的执行流程	119
2.4.2 错误处理机制	128
2.4.3 作业调度机制	129
2.4.4 Shuffle 和排序	130
2.4.5 任务执行	134
2.5 思考题	137
第3章 Hadoop 进阶	138
3.1 Hadoop I/O 操作	138
3.1.1 I/O 操作中的数据检查	138
3.1.2 数据的压缩	146
3.1.3 数据的 I/O 序列化操作	148
3.1.4 针对 Mapreduce 的文件类	162
3.2 Hadoop 的管理	174
3.2.1 HDFS 文件结构	174
3.2.2 Hadoop 的状态监视和管理工具	178
3.2.3 Hadoop 集群的维护	196
3.3 下一代 MapReduce: YARN	202
3.3.1 MapReduce V2 设计需求	202
3.3.2 MapReduce V2 主要思想和架构	203
3.3.3 MapReduce V2 设计细节	205
3.3.4 MapReduce V2 优势	208
3.4 思考题	209
第4章 Hadoop 实战	210
4.1 实战一 MapReduce 实现推荐系统	210
4.1.1 作业描述	210
4.1.2 作业分析	210

4.1.3 程序代码	213
4.1.4 准备输入数据	221
4.1.5 运行程序	221
4.1.6 代码结果	221
4.2 实战二 使用 MapReduce 求每年最低温度	223
4.2.1 作业描述	223
4.2.2 程序代码	223
4.2.3 准备输入数据	225
4.2.4 运行程序	226
4.2.5 代码结果	226
参考文献	227

第1章 Hadoop 基础

1.1 Hadoop 简介

1.1.1 Hadoop 概述

Hadoop 是 Apache 软件基金会旗下的一个开源分布式计算平台，以 Hadoop 分布式文件系统(Hadoop Distributed File System, HDFS)和 MapReduce(Google MapReduce 的开源实现)为核心的 Hadoop 为用户提供了系统底层细节透明的分布式基础架构。HDFS 的高容错性、高伸缩性等优点允许用户将 Hadoop 部署在低廉的硬件上，形成分布式系统；MapReduce 分布式编程模型允许用户在不了解分布式系统底层细节的情况下开发并行应用程序。所以用户可以利用 Hadoop 轻松地组织计算机资源，从而搭建自己的分布式计算平台，并且可以充分利用集群的计算和存储能力，完成海量数据的处理。经过业界和学术界长达十余年的锤炼，目前的 Hadoop 2.7.3 已经趋于完善，在实际的数据处理和分析任务中担当着不可替代的角色。

(1) Hadoop 的历史

Hadoop 起源于 Apache Nutch 项目，始于 2002 年，是 Apache Lucene 的子项目之一。2004 年，Google 在“操作系统设计与实现”(Operating System Design and Implementation, OSDI)会议上公开发表了题为 MapReduce: Simplified Data Processing on Large Clusters (MapReduce: 简化大规模集群上的数据处理)的论文之后，受到启发的 Doug Cutting 等人开始尝试实现 MapReduce 计算框架，并将它与 NDFS(Nutch Distributed File System)结合，用以支持 Nutch 引擎的主要算法。由于 NDFS 和 MapReduce 在 Nutch 引擎中有着良好的应用，所以它们于 2006 年 2 月被分离出来，成为一套完整而独立的软件，并被命名为 Hadoop。到了 2008 年年初，hadoop 已成为 Apache 的顶级项目，包含众多子项目，被应用到包括 Yahoo! 在内的很多互联网公司。现在的 Hadoop 2.7.3 版本已经发展成为包含 HDFS、MapReduce 子项目，与 Pig、ZooKeeper、Hive、HBase 等项目相关的大型应用工程。

(2) Hadoop 的功能与作用

众所周知，现代社会信息增长速度很快，这些信息中又积累着大量数据，其中包括个人数据和工业数据。预计到 2020 年，每年产生的数字信息中将会有超过 1/3 的内容驻留在云平台或借助云平台处理。我们需要对这些数据进行分析处理，以获取更多有价值的信息。那么我们如何高效地存储管理这些数据、如何分析这些数据呢？这时可以选用

Hadoop 系统。在处理这类问题时，它采用分布式存储方式来提高读写速度和扩大存储容量；采用 MapReduce 整合分布式文件系统上的数据，保证高速分析处理数据；与此同时还采用存储冗余数据来保证数据的安全性。

Hadoop 中的 HDFS 具有高容错性，并且是基于 Java 语言开发的，这使得 Hadoop 可以部署在低廉的计算机集群中，同时不限于某个操作系统。Hadoop 中 HDFS 的数据管理能力、MapReduce 处理任务时的高效率以及它的开源特性，使其在同类分布式系统中大放异彩，并在众多行业和科研领域中被广泛应用。

(3) Hadoop 的优势

Hadoop 是一个能够让用户轻松架构和使用的分布式计算平台。用户可以轻松地在 Hadoop 上开发运行处理海量数据的应用程序。它主要有以下几个优点：

高可靠性：Hadoop 按位存储和处理数据的能力值得人们信赖。

高扩展性：Hadoop 是在可用的计算机集簇间分配数据完成计算任务的，这些集簇可以方便地扩展到数以千计的节点中。

高效性：Hadoop 能够在节点之间动态地移动数据，以保证各个节点的动态平衡，因此其处理速度非常快。

高容错性：Hadoop 能够自动保存数据的多份副本，并且能够自动将失败的任务重新分配。

(4) Hadoop 应用现状和发展趋势

由于 Hadoop 优势突出，基于 Hadoop 的应用已经遍地开花，尤其是在互联网领域。Yahoo! 通过集群运行 Hadoop，用于支持广告系统和 Web 搜索的研究；Facebook 借助集群运行 Hadoop 来支持其数据分析和机器学习；搜索引擎公司百度则使用 Hadoop 进行搜索日志分析和网页数据挖掘工作；淘宝的 Hadoop 系统用于存储并处理电子商务交易的相关数据；中国移动研究院基于 Hadoop 的“大云”(BigCloud)系统对数据进行分析并对外提供服务。

2008 年 2 月，作为 Hadoop 最大贡献者的 Yahoo! 构建了当时最大规模的 Hadoop 应用，在 2000 个节点上面执行了超过 1 万个 Hadoop 虚拟机器来处理超过 5PB 的网页内容，分析大约 1 兆个网络连接之间的网页索引资料。这些网页索引资料压缩后超过 300TB。Yahoo! 正是基于这些为用户提供了高质量的搜索服务。

Hadoop 开源社区于 2011 年 10 月推出了基于新一代构架的 Hadoop0.23.0 测试版，该版本后演化为 Hadoop2.0 版本，即新一代的 Hadoop 系统 YARN。YARN 构架将主控节点的资源管理和作业管理功能分离设置，引入了全局资源管理器(Resource Manager)和针对每个作业的应用主控管理器(Application Master)，以此减轻原主控节点的负担，并可基于 Zookeeper 实现资源管理器的失效恢复，以此提高了 Hadoop 系统的高可用性(High Availability, HA)。YARN 还引入了资源容器(Resource Container)的概念，将系统计算资源统一划分和封装为很多资源单元，以此提高计算资源的利用率。此外，YARN 还能容纳 MapReduce 以外的其他多种并行计算模型和框架，提高了 Hadoop 框架并行化编程的灵活性。

Hadoop 目前已经取得了非常突出的成绩。随着互联网的发展，新的业务模式还将不

断涌现，Hadoop 的应用也会从互联网领域向电信、电子商务、银行、生物制药等领域拓展。相信在未来，Hadoop 将会在更多的领域中扮演幕后英雄，为我们提供更加快捷优质的服务。

1.1.2 Hadoop 项目及其结构

现在 Hadoop 已经发展成为包含很多项目的集合。虽然其核心内容是 MapReduce 和 Hadoop 分布式文件系统，但与 Hadoop 相关的 Common、Avro、Chukwa、Hive、HBase 等项目也是不可或缺的。它们提供了互补性服务或在核心层上提供了更高层的服务。图 1.1.1 是 Hadoop 的项目结构图。

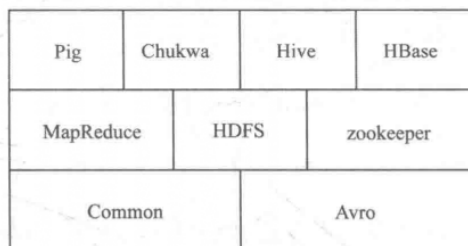


图 1.1.1 Hadoop 的项目结构图

下面将对 Hadoop 的各个关联项目进行更详细的介绍。

(1) Common

Common 是为 Hadoop 其他子项目提供支持的常用工具，它主要包括 FileSystem、RPC 和串行化库。它们为在廉价硬件上搭建云计算环境提供基本的服务，并且会为运行在该平台上的软件开发提供所需的 API。

(2) Avro

Avro 是用于数据序列化的系统。它提供了丰富的数据结构类型、快速可压缩的二进制数据格式、存储持久性数据的文件集、远程调用 RPC 的功能和简单的动态语言集成功能。其中代码生成器既不需要读写文件数据，也不需要实现或实现 RPC 协议，它只是一个可选的对静态类型语言的实现。

Avro 系统依赖于模式(Schema)，数据的读和写是在模式之下完成的。这样可以减少写入数据的开销，提高序列化的速度并缩减其大小；同时，也可以方便动态脚本语言的使用，因为数据连同其模式都是自描述的。

在 RPC 中，Avro 系统的客户端和服务端通过握手协议进行模式的交换，因此当客户端和服务端拥有彼此全部的模式时，不同模式下相同命名字段、丢失字段和附加字段等信息的一致性问题的得到了很好的解决。

(3) MapReduce

MapReduce 是一种编程模式，用于大规模数据集(大于 1TB)的并行运算。映射(Map)、化简(Reduce)的概念和它们的主要思想都是从函数式编程语言中借鉴而来的。

它极大地方便了编程人员——即使在不了解分布式并行编程的情况下，也可以将自己的程序运行在分布式系统上。MapReduce 在执行时先指定一个 Map(映射)函数，把输入键值对映射成一组新的键值对，经过一定处理后交给 Reduce，Reduce 对相同 key 下的所有 value 进行处理后再输出键值对作为最终的结果。

图 1.1.2 是 MapReduce 的任务处理流程图，它展示了 MapReduce 程序将输入划分到不同的 Map 上，再将 Map 的结果合并到 Reduce，然后进行处理的输出过程。

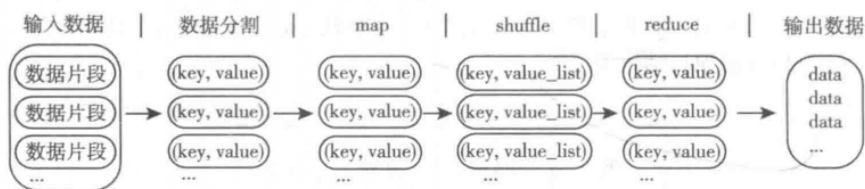


图 1.1.2 MapReduce 的任务处理流程图

(4) HDFS

HDFS 是一个分布式文件系统。因为 HDFS 具有高容错性(fault-tolerant)的特点，所以它可以设计部署在低廉(low-cost)的硬件上。它可以通过提供高吞吐率(high throughput)来访问应用程序的数据，适合那些有着超大数据集的应用程序。HDFS 放宽了对可移植操作系统接口(Portable Operating System Interface, POSIX)的要求，这样可以实现以流的形式访问文件系统中的数据。HDFS 原本是开源的 Apache 项目 Nutch 的基础结构，最后它却成为 Hadoop 基础架构之一。

以下几个方面是 HDFS 的设计目标：

检测和快速恢复硬件故障。硬件故障是计算机常见的问题。整个 HDFS 系统由数百甚至数千个存储着数据文件的服务器组成。而如此多的服务器则意味着高故障率，因此，故障的检测和快速自动恢复是 HDFS 的一个核心目标。

流式的数据访问。HDFS 使应用程序流式地访问它们的数据集。HDFS 被设计成适合进行批量处理，而不是用户交互式处理。所以它重视数据吞吐量，而不是数据访问的反应速度。

简化一致性模型。大部分的 HDFS 程序对文件的操作需要一次写入，多次读取。一个文件一旦经过创建、写入、关闭就不需要修改了。这个假设简化了数据一致性问题和高吞吐量的数据访问问题。

通信协议。所有的通信协议都是在 TCP/IP 协议之上的。一个客户端和明确配置了端口的名字节点(NameNode)建立连接之后，它和名字节点的协议便是客户端协议(Client Protocol)。数据节点(DataNode)和名字节点之间则用数据节点协议(DataNode Protocol)。

关于 HDFS 的具体介绍请参考本章 1.1.3 节。

(5) Chukwa

Chukwa 是开源的数据收集系统，用于监控和分析大型分布式系统的数据。Chukwa 是