

TURING

图灵程序设计丛书

[PACKT]  
PUBLISHING

Machine Learning with Spark Second Edition

# Spark机器学习

## (第2版)

[印] 拉结帝普·杜瓦 [印] 曼普利特·辛格·古特拉 [南非] 尼克·彭特里思 著  
蔡立宇 黄章帅 周济民 译

- Spark项目管理委员会成员作品
- 注重技术实践，通过大量实例演示如何创建有用的机器学习系统



中国工信出版集团

人民邮电出版社  
POSTS & TELECOM PRESS

设计丛书

Machine Learning with Spark Second Edition

# Spark机器学习

(第2版)

[印] 拉结帝普·杜瓦 [印] 曼普利特·辛格·古特拉 [南非] 尼克·彭特里思 著  
蔡立宇 黄章帅 周济民 译

人民邮电出版社  
北 京

## 图书在版编目 (C I P) 数据

Spark机器学习: 第2版 / (印) 拉结帝普·杜瓦,  
(印) 曼普利特·辛格·古特拉, (南非) 尼克·彭特里思  
著; 蔡立宇, 黄章帅, 周济民译. — 北京: 人民邮电  
出版社, 2018. 11

(图灵程序设计丛书)

ISBN 978-7-115-49783-3

I. ①S… II. ①拉… ②曼… ③尼… ④蔡… ⑤黄…  
⑥周… III. ①数据处理软件 IV. ①TP274

中国版本图书馆CIP数据核字(2018)第243357号

## 内 容 提 要

本书结合案例研究讲解 Spark 在机器学习中的应用, 并介绍如何从各种公开渠道获取用于机器学习系统的数据。内容涵盖推荐系统、回归、聚类、降维等经典机器学习算法及其实际应用。第2版新增了有关机器学习数学基础以及 Spark ML Pipeline API 的章节, 内容更加系统、全面、与时俱进。

本书适合数据分析从业人员以及高校数据挖掘相关专业师生阅读参考。

- 
- ◆ 著 [印] 拉结帝普·杜瓦 [印] 曼普利特·辛格·古特拉  
[南非] 尼克·彭特里思  
译 蔡立宇 黄章帅 周济民  
责任编辑 温 雪  
责任印制 周昇亮
- ◆ 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>  
三河市君旺印务有限公司印刷
- ◆ 开本: 800×1000 1/16  
印张: 24.25  
字数: 573千字 2018年11月第1版  
印数: 1-3 500册 2018年11月河北第1次印刷  
著作权合同登记号 图字: 01-2017-8608号
- 

定价: 99.00元

读者服务热线: (010)51095186转600 印装质量热线: (010)81055316

反盗版热线: (010)81055315

广告经营许可证: 京东工商广登字 20170147 号

站在巨人的肩上

Standing on Shoulders of Giants



iTuring.cn

站在巨人的肩上  
**Standing on Shoulders of Giants**



iTuring.cn

# 前言

近年来，被收集、存储和分析的数据量呈爆炸式增长，特别是与网络、移动设备相关的数据，以及传感器产生的数据。大规模数据的存储、处理、分析和建模，以前只有 Google、Yahoo!、Facebook、Twitter 和 Salesforce 这样的大公司才会涉及，而现在越来越多的机构都会面对处理海量数据的挑战。

面对如此量级的数据以及常见的实时利用数据的需求，人工驱动的系统就难以应对。这就催生了所谓的大数据和机器学习系统，它们从数据中学习并可自动决策。

为了能以低成本实现对持续增长的大规模数据的支持，Google、Yahoo!、Amazon 和 Facebook 等公司推出了大量开源技术。这些技术旨在通过在计算机集群上进行分布式数据存储和计算来简化大数据处理。

这些技术中最广为人知的是 Apache Hadoop，它极大地简化了海量数据的存储（通过 HDFS，即 Hadoop distributed file system）和计算（通过 Hadoop MapReduce，一种在集群里多个节点上进行并行计算的框架）流程，并降低了相应的成本。

然而，MapReduce 有其严重的缺点，如启动任务时的高开销、对中间数据和计算结果写入磁盘的依赖。这些都使得 Hadoop 不适合迭代式或低延迟的任务。Apache Spark 是一个新的分布式计算框架，从设计开始便注重对低延迟任务的优化，并将中间数据和结果保存在内存中，从而弥补了 Hadoop 框架中的一些主要缺陷。Spark 提供简洁明了的函数式 API，并完全兼容 Hadoop 生态系统。

不止如此，Spark 还提供针对 Scala、Java、Python 和 R 语言的原生 API。通过 Scala 和 Python 的 API，Spark 应用程序可充分利用 Scala 或 Python 语言的优势。这些优势包括使用相关的解释程序进行实时交互式的程序编写。Spark 现在还自带了一个支持分布式机器学习和包含若干数据挖掘模型的工具包（1.6 版中为 Spark MLlib，2.0 版则对应 Spark ML）。该工具包正在重点开发中，但已包括多个针对常见机器学习任务的高质量、可扩展的算法。本书将会涉及部分此类任务。

在大型数据集上进行机器学习颇具挑战性。这主要是因为常见的机器学习算法并非为并行架构而设计。多数情况下，设计这样的算法并不容易。机器学习模型一般具有迭代式的特性，而这与 Spark 的设计目标一致。并行计算的框架有很多，但很少能在兼顾速度、可扩展性、内存处理

和容错性的同时，还提供灵活、表达力丰富的 API。Spark 是其中为数不多的一个。

本书将关注机器学习技术的实际应用。我们会简要介绍机器学习算法的一些理论知识，但总体来说本书注重技术实践。具体来说，我们会通过示例程序和样例代码，举例说明如何借助 Spark、MLlib 以及其他常见的免费机器学习和数据分析套件来创建一个有用的机器学习系统。

## 本书内容

第 1 章“Spark 的环境搭建与运行”会讲到如何安装和搭建 Spark 框架的本地开发环境，以及怎样使用 Amazon EC2 在云端创建 Spark 集群；然后会介绍 Spark 编程模型和 API；最后分别用 Scala、Java 和 Python 语言创建一个简单的 Spark 应用。

第 2 章“机器学习的数学基础”会提供机器学习所需的基础数学知识。要理解算法，从而获得更好的建模效果，理解数学及其技巧十分重要。

第 3 章“机器学习系统设计”会展示一个贴合实际的机器学习系统案例。随后会针对该案例设计一个基于 Spark 的智能系统所对应的高层架构。

第 4 章“Spark 上数据的获取、处理与准备”会详细介绍如何从各种免费的公开渠道获取用于机器学习系统的数据。我们将学到如何进行数据清理和清理，并通过可用的工具、库和 Spark 函数将它们转换为符合要求的数据，使之具备可用于机器学习模型的特征。

第 5 章“Spark 构建推荐引擎”展示了如何创建一个基于协同过滤的推荐模型。该模型将用于向给定用户推荐物品，以及创建与给定物品相似的物品清单。这一章还会讲到如何使用标准指标来评估推荐模型的效果。

第 6 章“Spark 构建分类模型”阐述如何创建二元分类模型，以及如何利用标准的性能评估指标来评估分类效果。

第 7 章“Spark 构建回归模型”扩展了第 6 章中的分类模型以创建一个回归模型，并详细介绍了回归模型的评估指标。

第 8 章“Spark 构建聚类模型”探索如何创建聚类模型以及相关评估方法的使用。你会学到如何分析和可视化聚类结果。

第 9 章“Spark 应用于数据降维”将通过多种方法从数据中提取其内在结构并降低其维度。你会学到一些常见的降维方法，以及如何对它们进行应用和分析。这里还会讲到如何将降维的结果作为其他机器学习模型的输入。

第 10 章“Spark 高级文本处理技术”介绍了处理大规模文本数据的方法。这包括从文本中提取特征以及处理文本数据中常见的高维特征的方法。

第 11 章“Spark Streaming 实时机器学习”对 Spark Streaming 进行了综述，并介绍它如何在流数据上的机器学习中实现对在线和增量学习方法的支持。

第 12 章“Spark ML Pipeline API”在 Data Frames 的基础上提供了一套统一的接口（API），帮助用户创建和调试机器学习流程。

## 预备知识

本书假设读者已有基本的 Scala、Java、Python 或 R 编程经验，以及机器学习、统计学和数据分析方面的基础知识。

## 目标读者

本书的预期读者是初、中级数据科学研究者、数据分析师、软件工程师和对大规模环境下的机器学习或数据挖掘感兴趣的人。读者不需要熟悉 Spark，但若具有统计、机器学习相关软件（比如 MATLAB、scikit-learn、Mahout、R 或 Weka 等）或分布式系统（如 Hadoop）的实践经验，会很有帮助。

## 排版约定

在本书中，你会发现一些不同的文本样式，用以区别不同种类的信息。下面举例说明。

代码段的格式如下：

```
val conf = new SparkConf()
    .setAppName("Test Spark App")
    .setMaster("local[4]")
val sc = new SparkContext(conf)
```

所有的命令行输入或输出的格式如下：

```
> tar xfvz spark-2.1.0-bin-hadoop2.7.tgz
> cd spark-2.1.0-bin-hadoop2.7
```

新术语和重点词汇以黑体表示。屏幕、目录或对话框上的内容这样表示：“这些信息可以从 AWS 主页上依次点击 Account | Security Credentials | Access Credentials 看到。”



这个图标表示警告或需要特别注意的内容。



这个图标表示提示或技巧。

## 读者反馈

欢迎提出反馈。如果你对本书有任何想法，喜欢它什么，不喜欢它什么，请让我们知道。要写出真正对大家有帮助的书，了解读者的反馈很重要。一般的反馈，请发送电子邮件至 [feedback@packtpub.com](mailto:feedback@packtpub.com)，并在邮件主题中包含书名。如果你有某个主题的专业知识，并且有兴趣写成或帮助促成一本书，请参考我们的作者指南 <https://www.packtpub.com/books/info/packt/authors>。

## 客户支持

现在，你是一位令人自豪的 Packt 图书的拥有者，我们会尽全力帮你充分利用你手中的书。

## 下载示例代码

你可以用你的账户从 <https://www.packtpub.com> 下载所有已购买 Packt 图书的示例代码文件。<sup>①</sup>如果你从其他地方购买本书，可以访问 <https://www.packtpub.com/books/content/support> 并注册，我们将通过电子邮件把文件发送给你。

你可通过如下步骤下载本书示例代码：

- (1) 在上述网站上，使用你自己的邮件地址和密码登录或注册；
- (2) 将鼠标移动到顶部的 SUPPORT 标签页；
- (3) 点击 Code Downloads & Errata；
- (4) 在 Search 栏中输入本书名；
- (5) 在搜索结果中选择要下载代码的书；
- (6) 从下拉菜单中选择从何处购买的本书；
- (7) 点击 Code Download 下载包含代码和部分数据的代码包。

代码包下载后，请使用如下软件的最新版本来解压缩或提取所含资料：

- ❑ Windows 上建议使用 WinRAR/7-Zip；
- ❑ Mac 上建议使用 Zipex/iZip/UnRarX；
- ❑ Linux 上建议使用 7-Zip/PeaZip。

代码包在 GitHub 上也可获取，地址为 <https://github.com/PacktPublishing/Machine-Learning-with-Spark-Second-Edition>。<https://github.com/PacktPublishing> 上也列出了我们所出版的各类图书和视频的代码包，欢迎查看。

---

<sup>①</sup> 读者也可访问本书图灵社区页面（<http://www.it-ebooks.com.cn/book/2041>）下载示例代码及彩图，并提交本书中文版勘误。——编者注

## 勘误表

虽然我们已尽力确保本书内容正确，但出错仍旧在所难免。如果你在我们的书中发现错误，不管是文本还是代码，希望能告知我们，我们不胜感激。这样做可以减少其他读者的困扰，帮助我们改进本书的后续版本。如果你发现任何错误，请访问 <https://www.packtpub.com/books/info/packt/errata-submission-form-0> 提交，选择你的书，点击勘误表提交表单的链接 (Errata Submission Form)，并输入详细说明。勘误一经核实，你的提交将被接受，此勘误将上传到本公司网站或添加到现有勘误表。

从 <https://www.packtpub.com/books/content/support> 选择书名，在图书页面的 Errata 区域就可以查看现有的勘误表。

## 侵权行为

互联网上的盗版是所有媒体都要面对的问题。Packt 非常重视保护版权和许可证。如果你发现我们的作品在互联网上被非法复制，不管以什么形式，都请立即为我们提供位置地址或网站名称，以便我们可以寻求补救。

请把可疑盗版材料的链接发到 [copyright@packtpub.com](mailto:copyright@packtpub.com)。

非常感谢你帮助我们保护作者，以及保护我们给你带来有价值内容的能力。

## 问题

如果你对本书内容存有疑问，不管是哪个方面，都可以通过 [questions@packtpub.com](mailto:questions@packtpub.com) 联系我们，我们将尽最大努力来解决。

## 电子书

扫描如下二维码，即可购买本书电子版。



# 目 录

|                                                      |    |                            |    |
|------------------------------------------------------|----|----------------------------|----|
| 第 1 章 Spark 的环境搭建与运行                                 | 1  | 第 2 章 机器学习的数学基础            | 44 |
| 1.1 Spark 的本地安装与配置                                   | 2  | 2.1 线性代数                   | 45 |
| 1.2 Spark 集群                                         | 3  | 2.1.1 配置 IntelliJ Scala 环境 | 45 |
| 1.3 Spark 编程模型                                       | 4  | 2.1.2 配置命令行 Scala 环境       | 47 |
| 1.3.1 SparkContext 类与 SparkConf 类                    | 4  | 2.1.3 域                    | 48 |
| 1.3.2 SparkSession                                   | 5  | 2.1.4 矩阵                   | 54 |
| 1.3.3 Spark shell                                    | 6  | 2.1.5 函数                   | 64 |
| 1.3.4 弹性分布式数据集                                       | 8  | 2.2 梯度下降                   | 68 |
| 1.3.5 广播变量和累加器                                       | 12 | 2.3 先验概率、似然和后验概率           | 69 |
| 1.4 SchemaRDD                                        | 13 | 2.4 微积分                    | 69 |
| 1.5 Spark data frame                                 | 13 | 2.4.1 可微微分                 | 69 |
| 1.6 Spark Scala 编程入门                                 | 14 | 2.4.2 积分                   | 70 |
| 1.7 Spark Java 编程入门                                  | 17 | 2.4.3 拉格朗日乘子               | 70 |
| 1.8 Spark Python 编程入门                                | 19 | 2.5 可视化                    | 71 |
| 1.9 Spark R 编程入门                                     | 21 | 2.6 小结                     | 72 |
| 1.10 在 Amazon EC2 上运行 Spark                          | 23 | 第 3 章 机器学习系统设计             | 73 |
| 1.11 在 Amazon Elastic Map Reduce 上配置并运行 Spark        | 28 | 3.1 机器学习是什么                | 73 |
| 1.12 Spark 用户界面                                      | 31 | 3.2 MovieStream 介绍         | 74 |
| 1.13 Spark 所支持的机器学习算法                                | 32 | 3.3 机器学习系统商业用例             | 75 |
| 1.14 Spark ML 的优势                                    | 36 | 3.3.1 个性化                  | 75 |
| 1.15 在 Google Compute Engine 上用 Dataproc 构建 Spark 集群 | 38 | 3.3.2 目标营销和客户细分            | 76 |
| 1.15.1 Hadoop 和 Spark 版本                             | 38 | 3.3.3 预测建模与分析              | 76 |
| 1.15.2 创建集群                                          | 38 | 3.4 机器学习模型的种类              | 76 |
| 1.15.3 提交任务                                          | 41 | 3.5 数据驱动的机器学习系统的组成         | 77 |
| 1.16 小结                                              | 43 | 3.5.1 数据获取与存储              | 77 |
|                                                      |    | 3.5.2 数据清理与转换              | 78 |
|                                                      |    | 3.5.3 模型训练与测试循环            | 79 |
|                                                      |    | 3.5.4 模型部署与整合              | 79 |
|                                                      |    | 3.5.5 模型监控与反馈              | 80 |

|                                     |     |                                                  |     |
|-------------------------------------|-----|--------------------------------------------------|-----|
| 3.5.6 批处理或实时方案的选择                   | 80  | 5.3.2 使用隐式反馈数据训练模型                               | 143 |
| 3.5.7 Spark 数据管道                    | 81  | 5.4 使用推荐模型                                       | 143 |
| 3.6 机器学习系统架构                        | 82  | 5.4.1 ALS 模型推荐                                   | 144 |
| 3.7 Spark MLlib                     | 83  | 5.4.2 用户推荐                                       | 145 |
| 3.8 Spark ML 的性能提升                  | 83  | 5.4.3 物品推荐                                       | 148 |
| 3.9 MLlib 支持算法的比较                   | 85  | 5.5 推荐模型效果的评估                                    | 152 |
| 3.9.1 分类                            | 85  | 5.5.1 ALS 模型评估                                   | 152 |
| 3.9.2 聚类                            | 85  | 5.5.2 均方差                                        | 154 |
| 3.9.3 回归                            | 85  | 5.5.3 $K$ 值平均准确率                                 | 156 |
| 3.10 MLlib 支持的函数和开发者 API            | 86  | 5.5.4 使用 MLlib 内置的评估函数                           | 159 |
| 3.11 MLlib 愿景                       | 87  | 5.6 FP-Growth 算法                                 | 161 |
| 3.12 MLlib 版本的变迁                    | 87  | 5.6.1 FP-Growth 的基本例子                            | 161 |
| 3.13 小结                             | 88  | 5.6.2 FP-Growth 在 MovieLens<br>数据集上的实践           | 163 |
| 第 4 章 Spark 上数据的获取、处理<br>与准备        | 89  | 5.7 小结                                           | 164 |
| 4.1 获取公开数据集                         | 90  | 第 6 章 Spark 构建分类模型                               | 165 |
| 4.2 探索与可视化数据                        | 92  | 6.1 分类模型的种类                                      | 167 |
| 4.2.1 探索用户数据                        | 94  | 6.1.1 线性模型                                       | 167 |
| 4.2.2 探索电影数据                        | 102 | 6.1.2 朴素贝叶斯模型                                    | 177 |
| 4.2.3 探索评级数据                        | 104 | 6.1.3 决策树                                        | 180 |
| 4.3 数据的处理与转换                        | 109 | 6.1.4 树集成模型                                      | 183 |
| 4.4 从数据中提取有用特征                      | 112 | 6.2 从数据中抽取合适的特征                                  | 188 |
| 4.4.1 数值特征                          | 112 | 6.3 训练分类模型                                       | 189 |
| 4.4.2 类别特征                          | 113 | 6.4 使用分类模型                                       | 190 |
| 4.4.3 派生特征                          | 114 | 6.4.1 在 Kaggle/StumbleUpon<br>evergreen 数据集上进行预测 | 191 |
| 4.4.4 文本特征                          | 116 | 6.4.2 评估分类模型的性能                                  | 191 |
| 4.4.5 正则化特征                         | 121 | 6.4.3 预测的正确率和错误率                                 | 191 |
| 4.4.6 用软件包提取特征                      | 123 | 6.4.4 准确率和召回率                                    | 193 |
| 4.5 小结                              | 126 | 6.4.5 ROC 曲线和 AUC                                | 194 |
| 第 5 章 Spark 构建推荐引擎                  | 127 | 6.5 改进模型性能以及参数调优                                 | 196 |
| 5.1 推荐模型的分类                         | 128 | 6.5.1 特征标准化                                      | 197 |
| 5.1.1 基于内容的过滤                       | 128 | 6.5.2 其他特征                                       | 199 |
| 5.1.2 协同过滤                          | 128 | 6.5.3 使用正确的数据格式                                  | 202 |
| 5.1.3 矩阵分解                          | 130 | 6.5.4 模型参数调优                                     | 203 |
| 5.2 提取有效特征                          | 139 | 6.6 小结                                           | 211 |
| 5.3 训练推荐模型                          | 140 |                                                  |     |
| 5.3.1 使用 MovieLens 100k 数据<br>集训练模型 | 141 |                                                  |     |

|                                                      |     |
|------------------------------------------------------|-----|
| 第 7 章 Spark 构建回归模型                                   | 212 |
| 7.1 回归模型的种类                                          | 212 |
| 7.1.1 最小二乘回归                                         | 213 |
| 7.1.2 决策树回归                                          | 214 |
| 7.2 评估回归模型的性能                                        | 215 |
| 7.2.1 均方误差和均方根误差                                     | 215 |
| 7.2.2 平均绝对误差                                         | 215 |
| 7.2.3 均方根对数误差                                        | 216 |
| 7.2.4 $R$ -平方系数                                      | 216 |
| 7.3 从数据中抽取合适的特征                                      | 216 |
| 7.4 回归模型的训练和应用                                       | 220 |
| 7.4.1 BikeSharingExecutor                            | 220 |
| 7.4.2 在 bike sharing 数据集上训练<br>回归模型                  | 221 |
| 7.4.3 决策树集成                                          | 229 |
| 7.5 改进模型性能和参数调优                                      | 235 |
| 7.5.1 变换目标变量                                         | 235 |
| 7.5.2 模型参数调优                                         | 242 |
| 7.6 小结                                               | 256 |
| 第 8 章 Spark 构建聚类模型                                   | 257 |
| 8.1 聚类模型的类型                                          | 258 |
| 8.1.1 $K$ -均值聚类                                      | 258 |
| 8.1.2 混合模型                                           | 262 |
| 8.1.3 层次聚类                                           | 262 |
| 8.2 从数据中提取正确的特征                                      | 262 |
| 8.3 $K$ -均值训练聚类模型                                    | 265 |
| 8.3.1 训练 $K$ -均值聚类模型                                 | 266 |
| 8.3.2 用聚类模型来预测                                       | 267 |
| 8.3.3 解读预测结果                                         | 267 |
| 8.4 评估聚类模型的性能                                        | 271 |
| 8.4.1 内部评估指标                                         | 271 |
| 8.4.2 外部评估指标                                         | 272 |
| 8.4.3 在 MovieLens 数据集上计算<br>性能指标                     | 272 |
| 8.4.4 迭代次数对 WSSSE 的影响                                | 272 |
| 8.5 二分 $K$ -均值                                       | 275 |
| 8.5.1 二分 $K$ -均值——训练一个<br>聚类模型                       | 276 |
| 8.5.2 WSSSE 和迭代次数                                    | 280 |
| 8.6 高斯混合模型                                           | 283 |
| 8.6.1 GMM 聚类分析                                       | 283 |
| 8.6.2 可视化 GMM 类簇分布                                   | 285 |
| 8.6.3 迭代次数对类簇边界的影响                                   | 286 |
| 8.7 小结                                               | 287 |
| 第 9 章 Spark 应用于数据降维                                  | 288 |
| 9.1 降维方法的种类                                          | 289 |
| 9.1.1 主成分分析                                          | 289 |
| 9.1.2 奇异值分解                                          | 289 |
| 9.1.3 和矩阵分解的关系                                       | 290 |
| 9.1.4 聚类作为降维的方法                                      | 290 |
| 9.2 从数据中抽取合适的特征                                      | 291 |
| 9.3 训练降维模型                                           | 299 |
| 9.4 使用降维模型                                           | 302 |
| 9.4.1 在 LFW 数据集上使用 PCA<br>投影数据                       | 302 |
| 9.4.2 PCA 和 SVD 模型的关系                                | 303 |
| 9.5 评价降维模型                                           | 304 |
| 9.6 小结                                               | 307 |
| 第 10 章 Spark 高级文本处理技术                                | 308 |
| 10.1 文本数据处理的特别之处                                     | 308 |
| 10.2 从数据中抽取合适的特征                                     | 309 |
| 10.2.1 词加权表示                                         | 309 |
| 10.2.2 特征散列                                          | 310 |
| 10.2.3 从 20 Newsgroups 数据集<br>中提取 TF-IDF 特征          | 311 |
| 10.3 使用 TF-IDF 模型                                    | 324 |
| 10.3.1 20 Newsgroups 数据集的文<br>本相似度和 TF-IDF 特征        | 324 |
| 10.3.2 基于 20 Newsgroups 数据<br>集使用 TF-IDF 训练文本<br>分类器 | 326 |
| 10.4 评估文本处理技术的作用                                     | 328 |
| 10.5 Spark 2.0 上的文本分类                                | 329 |
| 10.6 Word2Vec 模型                                     | 331 |
| 10.6.1 借助 Spark MLlib 训练<br>Word2Vec 模型              | 331 |

|                                            |            |
|--------------------------------------------|------------|
| 10.6.2 借助 Spark ML 训练<br>Word2Vec 模型 ..... | 332        |
| 10.7 小结 .....                              | 334        |
| <b>第 11 章 Spark Streaming 实时机器学习 .....</b> | <b>335</b> |
| 11.1 在线学习 .....                            | 335        |
| 11.2 流处理 .....                             | 336        |
| 11.2.1 Spark Streaming 介绍 .....            | 337        |
| 11.2.2 Spark Streaming 缓存和<br>容错机制 .....   | 339        |
| 11.3 创建 Spark Streaming 应用 .....           | 340        |
| 11.3.1 消息生成器 .....                         | 341        |
| 11.3.2 创建简单的流处理程序 .....                    | 343        |
| 11.3.3 流式分析 .....                          | 346        |
| 11.3.4 有状态的流计算 .....                       | 348        |
| 11.4 使用 Spark Streaming 进行在线学习 .....       | 349        |
| 11.4.1 流回归 .....                           | 350        |
| 11.4.2 一个简单的流回归程序 .....                    | 350        |
| 11.4.3 流式 $K$ -均值 .....                    | 354        |
| 11.5 在线模型评估 .....                          | 355        |
| 11.6 结构化流 .....                            | 358        |
| 11.7 小结 .....                              | 359        |
| <b>第 12 章 Spark ML Pipeline API .....</b>  | <b>360</b> |
| 12.1 Pipeline 简介 .....                     | 360        |
| 12.1.1 DataFrame .....                     | 360        |
| 12.1.2 Pipeline 组件 .....                   | 360        |
| 12.1.3 转换器 .....                           | 361        |
| 12.1.4 评估器 .....                           | 361        |
| 12.2 Pipeline 工作原理 .....                   | 363        |
| 12.3 Pipeline 机器学习示例 .....                 | 367        |
| 12.4 小结 .....                              | 375        |

## 第 1 章

## Spark 的环境搭建与运行

Apache Spark 是一个分布式计算框架，旨在简化运行于计算机集群或虚拟机上的并行程序的编写。该框架对资源调度，任务的提交、执行和跟踪，节点间的通信以及数据并行处理的内在底层操作都进行了抽象。它提供了一个更高级别的 API 来处理分布式数据。从这方面说，它与 Apache Hadoop 等分布式处理框架类似。但在底层架构上，Spark 与它们有所不同。

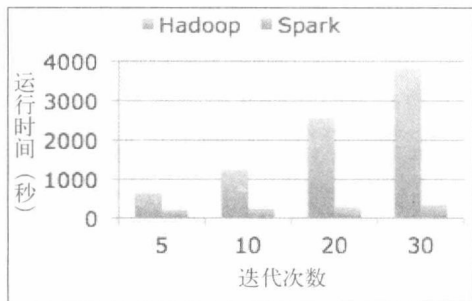
Spark 起源于加州大学伯克利分校 AMP 实验室的一个研究项目。该高校当时关注分布式机器学习算法的应用情况。因此，Spark 从一开始便是为应对迭代式应用的高性能需求而设计的。在这类应用中，相同的数据会被多次访问。该设计主要通过通过在内存中缓存数据集以及启动并行计算任务时的低延迟和低系统开销来实现高性能。再加上其容错性、灵活的分布式数据结构和强大的函数式编程接口，Spark 在各类基于机器学习和迭代分析的大规模数据处理任务上有广泛的应用，这也表明了其实用性。

关于 Spark 项目的更多信息，请参见：



- <http://spark.apache.org/community.html>
- <http://spark.apache.org/community.html#history>

从性能上说，Spark 在不同工作负载下的运行速度明显高于 Hadoop，如下图所示。



来源：<https://amplab.cs.berkeley.edu/wp-content/uploads/2011/11/spark-lr.png>

Spark 支持 4 种运行模式。

- ❑ 本地单机模式：所有 Spark 进程都运行在同一个 Java 虚拟机（JVM, Java virtual machine）进程中。
- ❑ 集群单机模式：使用 Spark 内置的任务调度框架。
- ❑ 基于 Mesos：Mesos 是一个流行的开源集群计算框架。
- ❑ 基于 Hadoop YARN：YARN 常被称作 NextGen MapReduce。

本章主要包括以下内容。

- ❑ 下载 Spark 二进制版本，并搭建一个在本地单机模式下运行的开发环境。本书各章的代码示例都在该环境下运行。
- ❑ 通过 Spark 的交互式终端来了解它的编程模型及 API。
- ❑ 分别用 Scala、Java、R 和 Python 语言来编写第一个 Spark 程序。
- ❑ 在 Amazon 的 EC2（Elastic Cloud Compute）平台上架设一个 Spark 集群。相比本地模式，该集群可以应对数据量更大、计算更复杂的任务。
- ❑ 借助 Amazon Elastic Map Reduce 服务来构建一个 Spark 集群。

如果读者曾构建过 Spark 环境并熟悉有关 Spark 程序编写的基础知识，可以跳过本章。

## 1.1 Spark 的本地安装与配置

Spark 能通过内置的单机集群调度器在本地模式下运行。此时，所有的 Spark 进程都高效地运行在同一个 Java 虚拟机中。这实际上构造了一个独立、多线程版本的 Spark 环境。本地模式常用于原型设计、开发、调试及测试。同样，它也适应于在单机上进行多核并行计算的实际场景。

Spark 的本地模式与集群模式完全兼容，在本地编写和测试过的程序仅需增加少许设置便能在集群上运行。

本地构建 Spark 环境的第一步是下载其最新的版本包。各版本的版本包及源代码的 GitHub 地址可在 Spark 项目的下载页面找到：<http://spark.apache.org/downloads.html>。



Spark 的在线文档（<http://spark.apache.org/docs/latest/>）涵盖了进一步学习 Spark 所需的各种资料。强烈推荐读者浏览查阅。

为了访问 Hadoop 分布式文件系统（HDFS）以及标准或定制的 Hadoop 输入源，Spark 的编译需要与 Hadoop 的版本对应。上述下载页面提供了针对 Cloudera 的 Hadoop 发行版（CHD）、MapR 的 Hadoop 发行版和 Hadoop 2（YARN）的预编译二进制包。除非你想构建针对特定版本 Hadoop 的 Spark，否则建议你通过如下链接从 Apache 镜像下载 Hadoop 2.7 预编译版本：<http://d3kbcqa49mib13.cloudfront.net/spark-2.0.2-bin-hadoop2.7.tgz>。

Spark 的运行依赖 Scala 编程语言（写作本书时为 2.10.x 或 2.11.x 版）。好在预编译的二进制

包中已包含 Scala 运行环境，我们不需要另外安装 Scala 便可运行 Spark。但是，你需要先安装好 Java 运行时环境（JRE）或 Java 开发工具包（JDK）。



相应的安装指南可参见本书代码包中的软硬件列表。推荐使用 R 3.1 或以上版本。

下载完上述版本包后，在终端输入如下指令解压软件包并进入解压出的文件夹：

```
$ tar xfvz spark-2.0.0-bin-hadoop2.7.tgz
$ cd spark-2.0.0-bin-hadoop2.7
```

用户启动 Spark 所用的脚本在该目录的 bin 文件夹下。可通过如下命令运行 Spark 附带的一个示例程序来测试是否一切正常：

```
$ ./bin/run-example SparkPi 100
```

该命令将在本地单机模式下运行 SparkPi 这个示例。在该模式下，所有的 Spark 进程均运行于同一个 JVM 中，而并行处理则通过多线程来实现。默认情况下，该示例会启用的线程数与本地系统的 CPU 核心数目相同。示例运行完后，应可在输出的结尾看到如下的提示：

```
...
16/11/24 14:41:58 INFO Executor: Finished task 99.0 in stage 0.0
(TID 99). 872 bytes result sent to driver
16/11/24 14:41:58 INFO TaskSetManager: Finished task 99.0 in stage
0.0 (TID 99) in 59 ms on localhost (100/100)
16/11/24 14:41:58 INFO DAGScheduler: ResultStage 0 (reduce at
SparkPi.scala:38) finished in 1.988 s
16/11/24 14:41:58 INFO TaskSchedulerImpl: Removed TaskSet 0.0,
whose tasks have all completed, from pool
16/11/24 14:41:58 INFO DAGScheduler: Job 0 finished: reduce at
SparkPi.scala:38, took 2.235920 s
Pi is roughly 3.1409527140952713
```

上述命令调用了 `org.apache.spark.examples.SparkPi` 类。

该类以 `local[N]` 格式来接受输入参数，其中 `N` 表示要启用的线程数目。比如只使用两个线程时，便可使用如下命令：

```
$ ./bin/spark-submit --class org.apache.spark.examples.SparkPi
--master local[2] ./examples/jars/spark-examples_2.11-2.0.0.jar 100
```

按惯例，将命令中的 `local[2]` 改为 `local[*]` 则会使用本机所有可用的核心。

## 1.2 Spark 集群

Spark 集群由两类进程构成：一个驱动程序和多个执行程序。在本地模式下，所有的进程都运行在同一个 JVM 内，而在集群模式下时，它们通常运行在不同的节点上。