



Spark Streaming

技术内幕及源码剖析

王家林 夏 阳 编著

- ◆ 彻底解密Spark Streaming架构设计
- ◆ 深入挖掘Spark Streaming运行机制
- ◆ 逐行剖析Spark Streaming框架源码
- ◆ 手把手讲解Spark Streaming性能调优



Spark Streaming

技术内幕及源码剖析



王家林 夏 阳 编著

清华大学出版社
北 京

内 容 简 介

本书以大数据处理引擎 Spark 的稳定版本 1.6.x 为基础,从应用案例、原理、源码、流程、调优等多个角度剖析 Spark 上的实时计算框架 Spark Streaming。在勾勒出 Spark Streaming 架构轮廓的基础上,从基本源码开始进行剖析,由浅入深地引导已具有 Spark 和 Spark Streaming 基础技术知识的读者进行 Spark Streaming 的进阶学习,理解 Spark Streaming 的原理和运行机制,为流数据处理的决策和应用提供了技术参考;结合 Spark Streaming 的深入应用的需要,对 Spark Streaming 的性能调优进行了分析,也对 Spark Streaming 功能的改造和扩展提供了指导。

本书适合大数据领域 CTO、架构师、高级软件工程师,尤其是 Spark 领域已有 Spark Streaming 基础知识的从业人员阅读,也可供需要深入学习 Spark、Spark Streaming 的高校研究生和高年级本科生参考。

本书封面贴有清华大学出版社防伪标签,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

Spark Streaming技术内幕及源码剖析/王家林,夏阳编著. —北京:清华大学出版社,2017
ISBN 978-7-302-46491-4

I. ①S… II. ①王… ②夏… III. ①数据处理软件 IV. ①TP274

中国版本图书馆 CIP 数据核字(2017)第 025588 号

责任编辑:袁金敏 战晓雷

封面设计:刘新新

责任校对:徐俊伟

责任印制:李红英

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×230mm 印 张:16.25 字 数:264 千字

版 次:2017 年 5 月第 1 版 印 次:2017 年 5 月第 1 次印刷

印 数:1~3000

定 价:49.00 元

产品编号:072312-01

前言

大数据浪潮汹涌来袭，这绝不仅仅是信息技术领域的革命，更是在全球范围引领社会变革的机遇。大数据的集群计算开源软件Spark在大数据计算平台应用领域日益凸显其重要地位。如果大数据技术领域从业人员的技术水平仍停留在只知使用开源软件，而不从开源软件的原理、架构上去理解，不到源码中去体会细节，则难以从根本上彻底解决现实工作中遇到的技术问题，更难以胜任大数据领域的技术创新工作。

可以预见，大数据的处理将越来越强调实时处理。Spark Streaming是建立在Spark上的实时计算框架，在Spark的各子框架中处于举足轻重的地位。彻底掌握 Spark Streaming的同时，也能加深对Spark Core技术的理解和掌握，还能具备开发同样高端的Spark应用程序的实力。对于有志向的Spark学习进阶者来说，深入了解Spark Streaming的源码是提高核心竞争力的捷径。

本书不仅对Spark Streaming的API做总结性介绍，而且重点针对Spark 1.6.x的Spark Streaming进行源码剖析。该书的开始部分对Spark的基本原理有一些阐述，但主要是彻底深入剖析Spark Streaming的内部原理。

读源码的人都怕自己走进大量源码的迷宫。为了提高源码学习效率，本书在剖析源码前，会对源码实现的功能的大致原理和流程轮廓进行介绍。书中有方便源码剖析的流程图，这对于理解和掌握Spark Streaming的各个功能非常重要。读者看到复杂的流程图时不一定要立刻全部理解掌握，但可以在源码学习过程中经常回过头来对照流程图以加深印象。

为了在书的页面内清晰展示复杂的流程图，书中绝大多数流程图采取了从上至下的树状结构来体现调用关系。每个方框中注明了类和方法，被其调用的类的方法会在下一行从左至右依次显示，调用和被调用的类方法间用有向线连接。有些方框上部会给出类的成员

变量，其类型就是方框中指明的类。粗箭头不是表示调用关系，而是表示传递消息。

源码剖析过程中，源码中关键的类名、方法名、注释会以粗体显示，使读者清楚重点。读者应以粗体部分为重点进行阅读，其他部分可以粗略浏览。有些源码篇幅过大，可能会省略其中的部分代码，以突出当前读者需要阅读的源码主体。

王家林 夏 阳

2017年2月27日于北京

目 录

第1章 Spark Streaming应用概述	1
1.1 Spark Streaming应用案例	2
1.2 Spark Streaming应用剖析	13
第2章 Spark Streaming基本原理	15
2.1 Spark Core简介	16
2.2 Spark Streaming设计思想	26
2.3 Spark Streaming整体架构	30
2.4 编程接口	33
第3章 Spark Streaming运行流程详解	39
3.1 从StreamingContext的初始化到启动	40
3.2 数据接收	54
3.3 数据处理	91
3.4 数据清理	115
3.5 容错机制	127
3.5.1 容错原理	128
3.5.2 Driver容错机制	152
3.5.3 Executor容错机制	161
3.6 No Receiver方式	167
3.7 输出不重复	175
3.8 消费速率的动态控制	176

3.9 状态操作	189
3.10 窗口操作	212
3.11 页面展示	216
3.12 Spark Streaming应用程序的停止	227

第4章 Spark Streaming性能调优机制

237

4.1 并行度解析	238
4.1.1 数据接收的并行度	238
4.1.2 数据处理的并行度	240
4.2 内存	240
4.3 序列化	240
4.4 Batch Interval	241
4.5 Task	242
4.6 JVM GC	242

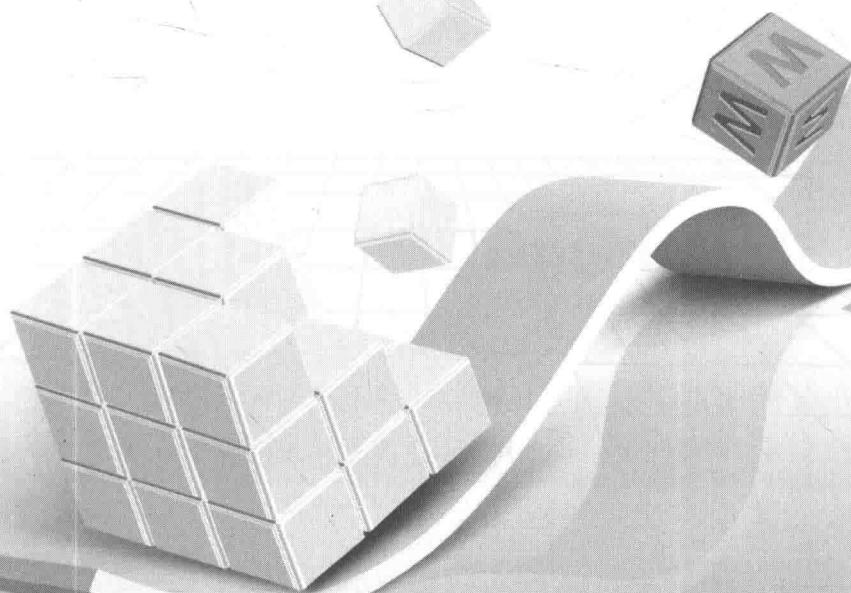
第5章 Spark 2.0中的流计算

245

5.1 连续应用程序	246
5.2 无边界表unbounded table	248
5.3 增量输出模式	249
5.4 API简化	250
5.5 其他改进	250

第1章

Spark Streaming应用概述



Spark是一个类似于Hadoop的MapReduce的分布式计算框架，其核心是弹性分布式数据集（Resilient Distributed Dataset，RDD），提供了比MapReduce更丰富的模型，可以在内存中快速对数据集进行多次迭代，以支持复杂的数据挖掘算法和图形计算算法。

Spark包含Spark Core、Spark SQL、Spark Streaming、MLlib、GraphX等主要部件，如图1-1所示。

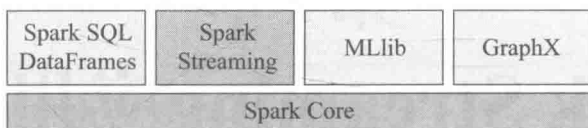


图1-1 Spark的主要部件

其中的Spark Core包含Spark的基本功能，是核心的计算引擎，并定义了RDD的API。其他的Spark部件都是构建在Spark Core之上的。

Spark Streaming是建立在Spark上的实时计算框架，提供了丰富的API，以支持Spark处理大规模流式数据。

这是一个流处理的时代。Spark十分强大的一个重要原因在于，它的流式处理可以在线使用图计算、机器学习或者SparkR的成果，这得益于Spark一体化、多元化的基础架构设计。在Spark Streaming中可以调用Spark SQL、MLlib、GraphX等其他子框架，无须任何设置。这是Spark无可匹敌之处。但Spark的应用中，Spark Streaming也是很容易出问题的。

Spark Streaming与其他子框架不同之处在于，它更像是Spark Core之上的一个复杂应用程序。如果要做Spark的定制开发，Spark Streaming的剖析也会提供最好的参考。

1.1 Spark Streaming应用案例

Spark Streaming应用程序运行的时候，往往在短时间内会产生大量日志信息，不利于研究分析。可以通过加大批处理时间间隔（batch interval）来降低批处理频率，减少日志

信息量，以便看清楚各个环节。

下面从一个Spark Streaming应用程序的开发入手，观察运行过程，以增强感性认识。以下是一个广告点击的在线黑名单过滤的Spark Streaming应用程序，程序中有详细注释，以方便初次接触Spark Streaming的读者理解。

源码1-1 OnlineBlackListFilter

```
package com.dt.spark.streaming

import org.apache.spark.SparkConf
import org.apache.spark.streaming.StreamingContext
import org.apache.spark.streaming.Seconds

object OnlineBlackListFilter {

  def main(args: Array[String]){

    /**
     * 第1步: 创建Spark的配置对象SparkConf, 设置Spark程序的运行时的配置信息,
     * 例如, 通过setMaster来设置程序要链接的Spark集群的Master的URL, 如果设置
     * 为local, 则代表Spark程序在本地运行, 特别适合于机器配置条件非常差(例如
     * 只有1GB内存)的初学者
     */

    // 创建SparkConf对象

    val conf = new SparkConf()

    // 设置应用程序的名称, 在程序运行的监控界面可以看到名称

    conf.setAppName("OnlineBlackListFilter")

    // 此时, 程序在Spark集群

    conf.setMaster("spark://Master:7077")
```

```
val ssc = new StreamingContext(conf, Seconds(30))

/**
 * 黑名单数据准备。实际上黑名单一般都是动态的，例如在Redis或者数据库中
 * 黑名单的生成往往有复杂的业务逻辑，具体情况算法不同
 * 但是在Spark Streaming进行处理的时候每次都能够访问完整的信息
 */

val blacklist = Array(("Spy", true), ("Cheater", true))

val blacklistRDD = ssc.sparkContext.parallelize(blacklist, 8)

val adsClickStream = ssc.socketTextStream("Master", 9999)

/**
 * 此处模拟的广告点击的每条数据的格式为: time、name
 * 此处map操作的结果是name、(time, name) 的格式
 */

val adsClickStreamFormatted = adsClickStream.map { ads => (ads.split(" ")(1), ads) }

adsClickStreamFormatted.transform(userClickRDD => {

    // 通过leftOuterJoin操作既保留了左侧用户广告点击内容的RDD的所有内容，
    // 又获得了相应点击内容是否在黑名单中

    val joinedBlackListRDD = userClickRDD.leftOuterJoin(blacklistRDD)

    /**
     * 进行filter过滤的时候，其输入元素是一个元组: (name, ((time, name), boolean))
     * 其中，第一个元素是黑名单的名称
     */
}
```

```

* 第二元素的第二个元素boolean是进行leftOuterJoin的时候是否存在的值。
* 如果存在，表明当前广告点击是黑名单，需要过滤掉，否则是有效点击内容
*/

val validClicked = joinedBlackListRDD.filter(joinedItem => {
    if(joinedItem._2._2.getOrElse(false))
    {
        false
    } else {
        true
    }
})

validClicked.map(validClick => {validClick._2._1})
}).print

ssc.start()

ssc.awaitTermination()

}
}

```

此程序接收Socket信息，过滤掉其中名称为Spy、Cheater的信息，并打印输出。

把程序的批处理时间间隔设置从30s改成300s:

```
val ssc = new StreamingContext(conf, Seconds(300))
```

然后重新生成一下jar包。

Spark集群有5台机器：Master、Worker1、Worker2、Worker3、Worker4。

启动Spark的History Server。

打开数据发送的端口：

```
nc -lk 9999
```

用spark-submit运行前面生成的jar包。

在数据发送端口输入若干数据，例如：

```
1375864674543 Tom
1375864674553 Spy
1375864674571 Andy
1375864688436 Cheater
1375864784240 Kelvin
1375864853892 Steven
1375864979347 John
```

每行第一项为时刻的毫秒数，第二项是程序中要过滤的名称。

打开浏览器，看History Server的日志信息，如图1-2所示。

图中按时间顺序显示了曾经运行过的应用程序，第一列是App ID，有各应用程序执行信息的链接。

单击最新的应用，看目前运行的应用程序中有什么Job，如图1-3所示。

这样一个Spark Streaming应用程序运行时总共有5个Job。

观察这些Job的内容，可以揭示一些现象。

Job 0不体现应用程序的业务逻辑代码，如图1-4所示。其实此Job是Spark Streaming出于对后面计算的负载均衡的考虑而产生的。



Spark 1.5.1 History Server

Event log directory: hdfs://Master:9000/historyserverforSpark

Showing 1-16 of 16

App ID	App Name	Started	Completed	Duration	Spark User	Last Up
app-2016043021417-0000	OnlineBlackListFilter	2016/04/30 21:14:13	2016/04/30 21:15:56	1.7 min	root	2016/04
app-20160430211216-0002	OnlineBlackListFilter	2016/04/30 21:12:12	2016/04/30 21:12:51	39 s	root	2016/04
app-20160430193930-0001	TransformBlackList	2016/04/30 19:39:56	2016/04/30 19:44:32	5.6 min	root	2016/04
app-20160430193930-0000	TransformBlackList	2016/04/30 19:39:53	2016/04/30 19:39:17	1.4 min	root	2016/04
local-1461736034129	NetworkWordCount	2016/04/27 13:52:06	2016/04/27 13:52:48	42 s	root	2016/04
local-1461736912866	NetworkWordCount	2016/04/27 13:50:05	2016/04/27 13:50:53	48 s	root	2016/04
local-14617366150514	NetworkWordCount	2016/04/27 13:48:24	2016/04/27 13:49:04	40 s	root	2016/04
app-20160427111949-0000	SparkSQL_192.168.112.235	2016/04/27 11:19:44	2016/04/27 12:47:06	2.5 h	root	2016/04
app-20160426140749-0000	Spark shell	2016/04/26 14:07:42	2016/04/26 14:09:03	1.3 min	root	2016/04
app-20160423132822-0001	MovieLensALS	2016/04/23 13:38:18	2016/04/23 13:51:24	13 min	root	2016/04
app-20160423133522-0002	MovieLensALS	2016/04/23 13:35:17	2016/04/23 13:35:40	23 s	root	2016/04
app-20160423132612-0001	MovieLensALS	2016/04/23 13:26:05	2016/04/23 13:28:15	10 s	root	2016/04
app-20160423132949-0000	Spark shell	2016/04/22 13:29:47	2016/04/22 13:30:34	47 s	root	2016/04
app-20160421072917-0001	FlumePushWordCount	2016/04/21 07:29:11	2016/04/21 07:29:54	43 s	root	2016/04
app-20160421072759-0000	FlumePushWordCount	2016/04/21 07:27:45	2016/04/21 07:29:05	1.3 min	root	2016/04
local-1461194188915	FlumePushWordCount	2016/04/21 07:16:19	2016/04/21 07:27:36	11 min	root	2016/04

Show incomplete applications.

图1-2 History Server日志信息示例



Spark 1.5.1 Jobs Stages Storage Environment Executors OnlineBlackListFilter application LR

Spark Jobs (7)

Total Uptime: 2.6 min
Scheduling Mode: FIFO
Completed Jobs: 5

Completed Jobs (5)

Job ID	Description	Submitted	Duration	Stages: Succeeded/Total	Tasks (for all stages): Succeeded/Total
1	Streaming job running monitor 0 start at OnlineBlackListFilter scala 72	2016/05/01 06:33:36	1.5 min	1/1	1/1
4	Streaming job from (input: operation C, batch time 06:35:06) start at OnlineBlackListFilter scala 67	2016/05/01 06:35:06	6.6 s	1/1 (2 skipped)	2/3 (15 skipped)
5	Streaming job from (input: operation C, batch time 06:36:02) start at OnlineBlackListFilter scala 67	2016/05/01 06:36:04	1 s	1/1 (2 skipped)	6/4 (15 skipped)
2	Streaming job from (input: operation C, batch time 06:35:06) start at OnlineBlackListFilter scala 57	2016/05/01 06:35:06	3 s	3/3	1/17
0	start at OnlineBlackListFilter scala 72	2016/05/01 06:33:20	18 s	2/2	26/75

图1-3 Spark Jobs页面示例



Details for Job 0

Status: SUCCEEDED
Completed Stages: 2

Completed Stages (2)

Stage ID	Description	Submitted	Duration	Tasks: Succeeded/Total	Input	Output	Shuffle Read	Shuffle Write
5	start at OnlineBlackListFilter scala 72	2016/05/01 06:33:34	2 s	20/22			1300.0 B	
0	start at OnlineBlackListFilter scala 72	2016/05/01 06:33:21	10 s	55/56				1300.0 B

图1-4 Details for Job 0示例

Job 0包含Stage 0、Stage 1。随便看一个Stage，比如Stage 0，看看其中的Aggregated Metrics by Executor部分，如图1-5所示。

Aggregated Metrics by Executor

Executor ID	Address	Task Time	Total Tasks	Failed Tasks	Succeeded Tasks	Shuffle Write Size / Records
0	Worker1.52115	45 s	8	0	8	206.0 B / 8
1	Worker4.34159	45 s	8	0	8	206.0 B / 8
2	Worker2.38749	52 s	26	0	26	676.0 B / 26
3	Worker3.50215	1.1 min	8	0	8	206.0 B / 8

图1-5 Aggregated Metrics by Executor页面示例

因为是分布式环境做负载均衡，所以Job 0 的Stage 1是在4个Worker的Executor上运行。

Job 1的运行时间比较长，耗时1.5min，如图1-6所示。

Details for Job 1

Status: **SUCCEEDED**
Completed Stages: 1
► Full Timeline
◄ DAG Visualization



Completed Stages (1)

Stage Id	Description	Submitted	Duration	Tasks: Succeeded/Total	Input	Output	Shuffle Read	Shuffle Write
2	Streaming job running receiver 0 start of DirectedAcyclicGraphJob PG	2016/05/01 08:33:36	1.5 min	1/1				

图1-6 页面示例：Details for Job 1

单击Stage 2的链接，看看Aggregated Metrics By Executor部分，如图1-7所示。

Aggregated Metrics by Executor

Executor ID	Address	Task Time	Total Tasks	Failed Tasks	Succeeded Tasks
1	Worker4.34159	1.5 min	1	0	1

图1-7 页面示例：Aggregated Metrics by Executor

可以知道，Stage 2只在Worker4上的一个Executor执行，而且执行了1.5min。

从业务处理的角度看，此前发送了很少的数据，这里却显示有一个运行1.5min的任务。这个任务是做什么呢？

从DAG Visualization部分可以知道此Job实际就是启动了一个接收数据的接收器（Receiver），如图1-8所示。

原来Receiver是通过一个Job来启动的。

Tasks部分如图1-9所示。

Details for Stage 2 (Attempt 0)

Total Time Across All Tasks: 1.5 min

Locality Level Summary: Process local: 1

▼ DAG Visualization

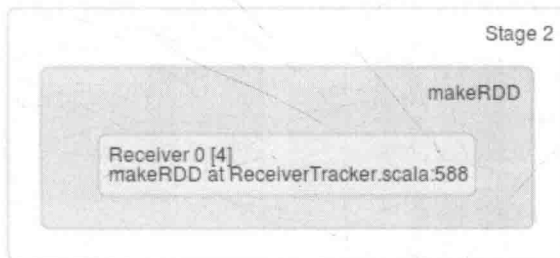


图1-8 页面示例: Details for Stage 2

Tasks														
Index	ID	Attempt	Status	Locality Level	Executor ID Host	Launch Time	Duration	Scheduler Delay	Task Deserialization Time	GC Time	Result Serialization Time	Getting Result Time	Peak Execution Memory	Errors
0	76	0	SUCCESS	PROCESS_LOCAL	1-1 Worker4	2016/05/01 05:22:38	1.5 min	25 ms	0.2 s	0 ms	0 ms	0 ms	0.0 B	

图1-9 页面示例: Tasks

只有一个Worker运行此Job，用于接收数据。

Spark Streaming应用程序启动后，自己会启动一些Job。默认情况是启动一个Job接收数据，为后续处理做准备。

Locality Level是PROCESS_LOCAL。所以，默认情况下，数据接收不会使用磁盘，而是直接使用内存中的数据。

从Job 2的Details可以发现程序的主要业务逻辑，体现在Stage 3、Stage 4、Stage 5中，如图1-10所示。

仔细观察Stage 3、Stage 4，可以知道这两个Stage都是用4个Executor执行的，所有数据处理是在4台机器上进行的，如图1-11所示。

Details for Job 2

Status: SUCCEEDED

Completed Stages: 3

▸ Event Timeline

▼ DAG Visualization

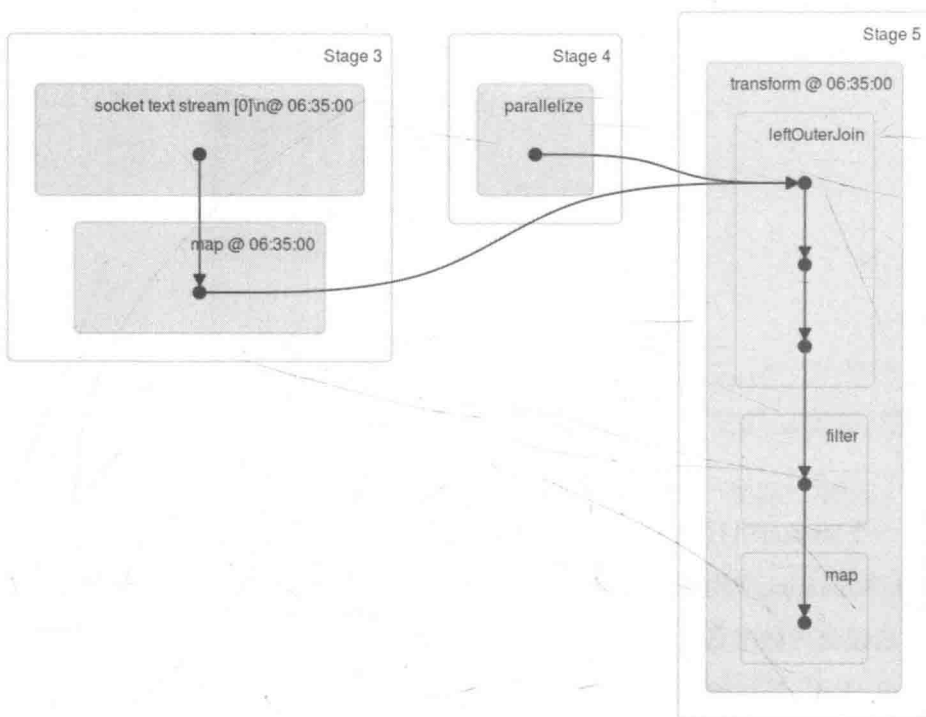


图1-10 页面示例：Details for Job 2

Aggregated Metrics by Executor

Executor ID ▲	Address	Task Time	Total Tasks	Failed Tasks	Succeeded Tasks	Shuffle Write Size / Records
0	Worker1:53116	2 s	1	0	1	152.0 B / 1
1	Worker4:34159	8 s	3	0	3	466.0 B / 3
2	Worker2:39749	4 s	2	0	2	309.0 B / 2
3	Worker3:50215	6 s	2	0	2	314.0 B / 2

图1-11 页面示例：Aggregated Metrics by Executor

Stage 5只在Worker4上，这是因为这个Stage有Shuffle操作。

Job 3有Stage 6、Stage 7、Stage 8，其中Stage 6、Stage 7显示灰色，说明被跳过，如图1-12所示。