

Apache Kafka 实战

胡夕 著



- 基于Apache Kafka 1.0.0版本进行介绍，Kafka Contributor执笔。
- 从Kafka基本概念与特性开始，详细介绍了Kafka的部署、开发、运营、监控、调试、优化以及重要组件的设计原理，并给出了翔实的案例。
- 本书既适合作为Kafka的入门书籍，也适合系统架构师和一线开发工程师参考阅读。



中国工信出版集团



电子工业出版社
PUBLISHING HOUSE OF ELECTRONICS INDUSTRY
<http://www.phei.com.cn>

Apache Kafka 实战

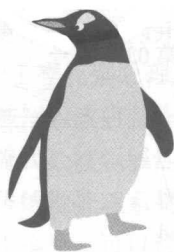
——



本书是一本关于 Apache Kafka 的实战指南，旨在帮助读者快速上手并深入理解 Kafka 的核心概念和最佳实践。本书从 Kafka 的架构、安装、配置、使用到高级特性，进行了全面的讲解。通过丰富的代码示例和实战案例，读者可以掌握 Kafka 在实际项目中的应用。本书适合对 Kafka 感兴趣的开发者和运维人员阅读。

Apache Kafka 实战

胡夕 著



電子工業出版社

Publishing House of Electronics Industry

北京•BEIJING

内 容 简 介

本书是涵盖 Apache Kafka 各方面的具有实践指导意义的工具书和参考书。作者结合典型的使用场景,对 Kafka 整个技术体系进行了较为全面的讲解,以便读者能够举一反三,直接应用于实践。同时,本书还对 Kafka 的设计原理及其流式处理组件进行了较深入的探讨,并给出了翔实的案例。

本书共分为 10 章:第 1 章全面介绍消息引擎系统以及 Kafka 的基本概念与特性,快速带领读者走进 Kafka 的世界;第 2 章简要回顾了 Apache Kafka 的发展历史;第 3 章详细介绍了 Kafka 集群环境的搭建;第 4、5 章深入探讨了 Kafka 客户端的使用方法;第 6 章带领读者一览 Kafka 内部设计原理;第 7~9 章以实例的方式讲解了 Kafka 集群的管理、监控与调优;第 10 章介绍了 Kafka 新引入的流式处理组件。

本书适合所有对云计算、大数据处理感兴趣的技术人员阅读,尤其适合对消息引擎、流式处理技术及框架感兴趣的技术人员参考阅读。

未经许可,不得以任何方式复制或抄袭本书之部分或全部内容。
版权所有,侵权必究。

图书在版编目(CIP)数据

Apache Kafka 实战 / 胡夕著. —北京:电子工业出版社, 2018.5
ISBN 978-7-121-33776-5

I. ①A... II. ①胡... III. ①分布式操作系统 IV. ①TP316.4

中国版本图书馆 CIP 数据核字 (2018) 第 037942 号

责任编辑:付 睿

印 刷:三河市双峰印刷装订有限公司

装 订:三河市双峰印刷装订有限公司

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:25

字数:557 千字

版 次:2018 年 5 月第 1 版

印 次:2018 年 5 月第 1 次印刷

定 价:89.00 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, 88258888。

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

本书咨询联系方式:010-51260888-819, faq@phei.com.cn。

前言

2011 年年初，美国领英公司（LinkedIn）开源了一款基础架构软件，以奥地利作家弗兰兹·卡夫卡（Franz Kafka）的名字命名，之后 LinkedIn 将其贡献给 Apache 基金会，随后该软件于 2012 年 10 月成功完成孵化并顺利晋升为 Apache 顶级项目——这便是大名鼎鼎的 Apache Kafka。历经 7 年发展，2017 年 11 月，Apache Kafka 正式演进到 1.0 时代，本书就是基于 1.0.0 版本来展开介绍 Kafka 的设计原理与实战的。

背景

这是一个最好的大数据时代，这是一个最坏的大数据时代！

很抱歉，我使用了这句改编后的狄更斯名言作为开头，我想没有谁会质疑“当今是大数据时代”这个论点。今年（2018 年）两会上李克强总理所做的政府工作报告中多次提及大数据等关键词，这已然是“大数据”第 5 次被写入政府工作报告了。具体到大数据行业内，各种各样的大数据产业方兴未艾，其中在实时流式处理领域涌现出大量的技术与框架，令技术人员们应接不暇。实时流式处理系统在克服了传统批处理系统延时方面的固有缺陷的同时，还摆脱了设计上的桎梏，实现了“梦寐以求”的正确性。可以说，对于流式处理从业人员来说，这正是摩拳擦掌、大展宏图的最好时代。

与此同时，我们也清醒地意识到当今大数据领域内的细分越来越精细化。不必说日渐火爆的人工智能和机器学习潮流引诱着我们改弦易辙，也不必说那些纷繁复杂的技术框架令人眼花缭乱，单是静下心来沉淀所学、思考方向的片刻时光于我们这些从业者来说都已显得弥足珍贵。我们仿佛在黑暗密林中徘徊，试图找出那条通往光明的“康庄大道”。每当发现了一条羊肠小路都好似救命稻草一般紧紧抓住。多年后我们回望那只不过是不断追逐热点罢了，在技术的海洋中我们迷失了前进的方向。从这个意义上说，这实在是一个糟糕的时代。

时光切回到 4 年前的某个下午，那时我正在做着 Kafka 的大数据项目。我突然发现与其盲目跟风各种技术趋势，何不精进手头的工作，把当前工作中用到的技术搞明白，于是我萌发了

研究 Kafka 的想法。直到今天，我都无比庆幸那个午后做出的冲动决定，正如 Adam Grant 在《离经叛道》一书中所说：最正确的决定都是在冲动之下做出的。诚不欺我！

想要深入学习 Kafka，不掌握 Scala 语言是不行的，毕竟 Kafka 就是使用 Scala 语言编写的。苦于当时没有合适的 Scala 中文书籍，我依稀记得找到了一本 600 多页的 Scala 原书（*Programming Scala Edition 2*）进行学习。那段时间实在是难熬！不得不说，英文版书籍虽然内容翔实，但在表述上实在晦涩难懂，比如 `partially applied function` 和 `partial function` 两者的区别直至今天我都不是特别清晰，还是要不断地翻阅资料才能隐约记得它们之间的不同。庆幸的是，我没有半途而废，600 多页的英文文档硬是啃了下来。对于 Scala 的初步掌握也让我觉得研究 Kafka 的时机到了。有意思的是，在之后通读 Kafka 的源码时我不禁大呼上当，Kafka 的源码中只使用了最简单的函数式编程，我有些后悔自己花了那么多时间去学习 Scala 的函数式编程，当然这是后话。

既然是研究 Kafka，那么研读源码是必不可少的步骤。如果不分析源码，我们就无法定位问题发生的根本原因。实话实说，阅读别人源码的过程是痛苦的，因而在理解的过程中我走了不少弯路。为了记录阅读 Kafka 源码的心得，我努力为每个 Kafka 源码包撰写博客。现在翻看我之前的博客，大家还能看到那好似流水账一般的 Kafka 源码分析系列文章。

随着对源码的不断熟悉，我加入了 Apache Kafka 社区，希望贡献自己的微薄之力。时至今日，我依然记得当初发送邮件要求加入开发组时的惶恐，也记得第一次贡献代码时的惴惴不安；我记得为了研究某个 Kafka bug，自己曾忘记吃中饭的执着，也记得自己被标记为“Kafka contributor”时的喜悦。在混迹社区的日子里，我逐渐认识了一些 Kafka 的 committer 们，比如 Kafka PMC 成员王国璋，国璋兄对于网上 Kafka 问题的权威解答令我受教良多，同时我也很感激他于百忙之中为本书写推荐语。还有 Kafka 的三位原作者之一的饶军（Rao Jun），几次问题交流让我看到了他霸气的决断能力以及对于疑难问题原因的毒辣分析。当然还有非常敬业的 Ijuma，他是我见过的最勤劳的 Kafka committer，没有之一。在编写本书的过程中，我都或多或少地得到过他们的帮助，再次表示衷心感谢。

由于对 Kafka 研究的日益深入，我终于有了写书的冲动。我希望通过把学到的知识和原理集中整理并书写成文字来帮助那些尚未接触 Kafka 的广大读者快速上手，降低他们学习使用 Kafka 的成本，于是有了今天这本《Apache Kafka 实战》。借着写作本书的契机，我本人对 Kafka 的方方面面做了梳理，自觉收获良多。每当搞懂了一个以前未了解的机制时，心中的那种满足感和兴奋感至今都令人神往。在此，我深深地希望读者在阅读完本书后也能有这样的体会。

面向的读者

我衷心希望本书可以成为各行各业的大数据从业者使用消息队列甚至是进入流式处理领域

内的“敲门砖”，也希望各大公司能够充分利用 Kafka 来实现自己的业务目标。

在编写本书的过程中，我阅读了大量的英文资料和源代码，试图通过自己的理解将 Kafka 的使用实战技巧深入浅出地呈现给广大读者。没错，我希望这本书给人的感觉是通俗易懂、深入浅出，从而方便引领读者快速进入 Kafka 学习的大门。

我本人维护了一个微信公众号（名为“大数据 Kafka 技术分享”），希望在该公众号中我能和读者朋友们一起深入交流和探讨 Kafka 学习过程中碰到的各种问题，同时我也会及时分享和推送各种最新的 Kafka 使用心得。

致谢

非常感谢 Kafka PMC 成员、Kafka Committer 王国璋对本书的大力支持。自开始编写本书之日起，国璋兄就给予我很大的鼓励与帮助，这也让我坚定了传播 Kafka 实战心得的决心。

感谢腾讯 AI 平台助理总经理王迪先生和我的好友贾兴华，你们对本书的评价之高实在是过誉了，但也令本人倍感振奋。

感谢我的前同事、新浪微博技术专家付稳。付总对本书整体结构和具体知识点的建议发人深省，其独到的行业见解令人佩服。

非常感谢电子工业出版社的编辑付睿女士。她细致、专业、严谨的工作作风深深地感染了我，在本书编写过程中她总是能及时地就书中的内容给出合理的建议和指导。

另外，我还想感谢一下我的家人，特别是我的妻子刘丹女士。过去一年中正是你坚定的支持和默默的付出才成就我撰写本书。对于你偶尔在学术上给予的提点我既感到惊讶，同时也欣慰不已。这为我漫长枯燥的写书过程平添了很多温暖。

最后，非常感谢本书的每一位读者。本人已经在写作过程中收获良多，我衷心希望你们在阅读本书时也有大呼过瘾的感觉。另外，我在“知乎”（ID: huxihx）的 Kafka 专栏以及 StackOverflow 网站上也会尽力回答关于 Kafka 的各类问题，希望通过这些途径可以和读者进行更加深入的交流。

由于本人水平有限，书中难免有遗漏和疏忽，也恳请各位读者多多指正。

胡夕

2018 年 3 月 15 日于北京

个人博客: <https://www.cnblogs.com/huxi2b/>

微信公众号: 大数据 Kafka 技术分享

电子邮箱: huxi_2b@hotmail.com

目录

第 1 章 认识 Apache Kafka.....	1
1.1 Kafka 快速入门.....	1
1.1.1 下载并解压缩 Kafka 二进制代码压缩包文件.....	2
1.1.2 启动服务器.....	3
1.1.3 创建 topic.....	3
1.1.4 发送消息.....	4
1.1.5 消费消息.....	4
1.2 消息引擎系统.....	5
1.2.1 消息设计.....	6
1.2.2 传输协议设计.....	6
1.2.3 消息引擎范型.....	6
1.2.4 Java 消息服务.....	8
1.3 Kafka 概要设计.....	8
1.3.1 吞吐量/延时.....	8
1.3.2 消息持久化.....	11
1.3.3 负载均衡和故障转移.....	12
1.3.4 伸缩性.....	13
1.4 Kafka 基本概念与术语.....	13
1.4.1 消息.....	14
1.4.2 topic 和 partition.....	16
1.4.3 offset.....	17
1.4.4 replica.....	18

1.4.5	leader 和 follower.....	18
1.4.6	ISR.....	19
1.5	Kafka 使用场景.....	20
1.5.1	消息传输.....	20
1.5.2	网站行为日志追踪.....	20
1.5.3	审计数据收集.....	20
1.5.4	日志收集.....	20
1.5.5	Event Sourcing.....	21
1.5.6	流式处理.....	21
1.6	本章小结.....	21
第 2 章	Kafka 发展历史.....	22
2.1	Kafka 的历史.....	22
2.1.1	背景.....	22
2.1.2	Kafka 横空出世.....	23
2.1.3	Kafka 开源.....	24
2.2	Kafka 版本变迁.....	25
2.2.1	Kafka 的版本演进.....	25
2.2.2	Kafka 的版本格式.....	26
2.2.3	新版本功能简介.....	26
2.2.4	旧版本功能简介.....	31
2.3	如何选择 Kafka 版本.....	35
2.3.1	根据功能场景.....	35
2.3.2	根据客户端使用场景.....	35
2.4	Kafka 与 Confluent.....	36
2.5	本章小结.....	37
第 3 章	Kafka 线上环境部署.....	38
3.1	集群环境规划.....	38
3.1.1	操作系统的选型.....	38
3.1.2	磁盘规划.....	40
3.1.3	磁盘容量规划.....	42

3.1.4	内存规划	43
3.1.5	CPU 规划	43
3.1.6	带宽规划	44
3.1.7	典型线上环境配置	45
3.2	伪分布式环境安装	45
3.2.1	安装 Java	46
3.2.2	安装 ZooKeeper	47
3.2.3	安装单节点 Kafka 集群	48
3.3	多节点环境安装	49
3.3.1	安装多节点 ZooKeeper 集群	50
3.3.2	安装多节点 Kafka	54
3.4	验证部署	55
3.4.1	测试 topic 创建与删除	55
3.4.2	测试消息发送与消费	57
3.4.3	生产者吞吐量测试	58
3.4.4	消费者吞吐量测试	58
3.5	参数设置	59
3.5.1	broker 端参数	59
3.5.2	topic 级别参数	62
3.5.3	GC 参数	63
3.5.4	JVM 参数	64
3.5.5	OS 参数	64
3.6	本章小结	65
第 4 章	producer 开发	66
4.1	producer 概览	66
4.2	构造 producer	69
4.2.1	producer 程序实例	69
4.2.2	producer 主要参数	75
4.3	消息分区机制	80
4.3.1	分区策略	80
4.3.2	自定义分区机制	80

4.4	消息序列化	83
4.4.1	默认序列化	83
4.4.2	自定义序列化	84
4.5	producer 拦截器	87
4.6	无消息丢失配置	90
4.6.1	producer 端配置	91
4.6.2	broker 端配置	92
4.7	消息压缩	92
4.7.1	Kafka 支持的压缩算法	93
4.7.2	算法性能比较与调优	93
4.8	多线程处理	95
4.9	旧版本 producer	96
4.10	本章小结	98
第 5 章	consumer 开发	99
5.1	consumer 概览	99
5.1.1	消费者 (consumer)	99
5.1.2	消费者组 (consumer group)	101
5.1.3	位移 (offset)	102
5.1.4	位移提交	103
5.1.5	__consumer_offsets	104
5.1.6	消费者组重平衡 (consumer group rebalance)	106
5.2	构建 consumer	106
5.2.1	consumer 程序实例	106
5.2.2	consumer 脚本命令	111
5.2.3	consumer 主要参数	112
5.3	订阅 topic	115
5.3.1	订阅 topic 列表	115
5.3.2	基于正则表达式订阅 topic	115
5.4	消息轮询	115
5.4.1	poll 内部原理	115
5.4.2	poll 使用方法	116

5.5	位移管理	118
5.5.1	consumer 位移	119
5.5.2	新版本 consumer 位移管理	120
5.5.3	自动提交与手动提交	121
5.5.4	旧版本 consumer 位移管理	123
5.6	重平衡 (rebalance)	123
5.6.1	rebalance 概览	123
5.6.2	rebalance 触发条件	124
5.6.3	rebalance 分区分配	124
5.6.4	rebalance generation	126
5.6.5	rebalance 协议	126
5.6.6	rebalance 流程	127
5.6.7	rebalance 监听器	128
5.7	解序列化	130
5.7.1	默认解序列化器	130
5.7.2	自定义解序列化器	131
5.8	多线程消费实例	132
5.8.1	每个线程维护一个 KafkaConsumer	133
5.8.2	单 KafkaConsumer 实例+多 worker 线程	135
5.8.3	两种方法对比	140
5.9	独立 consumer	141
5.10	旧版本 consumer	142
5.10.1	概览	142
5.10.2	high-level consumer	143
5.10.3	low-level consumer	147
5.11	本章小结	153
第 6 章	Kafka 设计原理	154
6.1	broker 端设计架构	154
6.1.1	消息设计	155
6.1.2	集群管理	166
6.1.3	副本与 ISR 设计	169

6.1.4	水印 (watermark) 和 leader epoch	174
6.1.5	日志存储设计	185
6.1.6	通信协议 (wire protocol)	194
6.1.7	controller 设计	205
6.1.8	broker 请求处理	216
6.2	producer 端设计	219
6.2.1	producer 端基本数据结构	219
6.2.2	工作流程	220
6.3	consumer 端设计	223
6.3.1	consumer group 状态机	223
6.3.2	group 管理协议	226
6.3.3	rebalance 场景剖析	227
6.4	实现精确一次处理语义	230
6.4.1	消息交付语义	230
6.4.2	幂等性 producer (idempotent producer)	231
6.4.3	事务 (transaction)	232
6.5	本章小结	234
第 7 章	管理 Kafka 集群	235
7.1	集群管理	235
7.1.1	启动 broker	235
7.1.2	关闭 broker	236
7.1.3	设置 JMX 端口	237
7.1.4	增加 broker	238
7.1.5	升级 broker 版本	238
7.2	topic 管理	241
7.2.1	创建 topic	241
7.2.2	删除 topic	243
7.2.3	查询 topic 列表	244
7.2.4	查询 topic 详情	244
7.2.5	修改 topic	245
7.3	topic 动态配置管理	246

7.3.1	增加 topic 配置	246
7.3.2	查看 topic 配置	247
7.3.3	删除 topic 配置	248
7.4	consumer 相关管理	248
7.4.1	查询消费者组	248
7.4.2	重设消费者组位移	251
7.4.3	删除消费者组	256
7.4.4	kafka-consumer-offset-checker	257
7.5	topic 分区管理	258
7.5.1	preferred leader 选举	258
7.5.2	分区重分配	260
7.5.3	增加副本因子	263
7.6	Kafka 常见脚本工具	264
7.6.1	kafka-console-producer 脚本	264
7.6.2	kafka-console-consumer 脚本	265
7.6.3	kafka-run-class 脚本	267
7.6.4	查看消息元数据	268
7.6.5	获取 topic 当前消息数	270
7.6.6	查询 consumer_offsets	271
7.7	API 方式管理集群	273
7.7.1	服务器端 API 管理 topic	273
7.7.2	服务器端 API 管理位移	275
7.7.3	客户端 API 管理 topic	276
7.7.4	客户端 API 查看位移	280
7.7.5	0.11.0.0 版本客户端 API	281
7.8	MirrorMaker	285
7.8.1	概要介绍	285
7.8.2	主要参数	286
7.8.3	使用实例	287
7.9	Kafka 安全	288
7.9.1	SASL+ACL	289
7.9.2	SSL 加密	297

7.10 常见问题	301
7.11 本章小结	304
第 8 章 监控 Kafka 集群	305
8.1 集群健康度检查	305
8.2 MBean 监控	306
8.2.1 监控指标	306
8.2.2 指标分类	308
8.2.3 定义和查询 JMX 端口	309
8.3 broker 端 JMX 监控	310
8.3.1 消息入站/出站速率	310
8.3.2 controller 存活 JMX 指标	311
8.3.3 备份不足的分区数	312
8.3.4 leader 分区数	312
8.3.5 ISR 变化速率	313
8.3.6 broker I/O 工作处理线程空闲率	313
8.3.7 broker 网络处理线程空闲率	314
8.3.8 单个 topic 总字节数	314
8.4 clients 端 JMX 监控	314
8.4.1 producer 端 JMX 监控	314
8.4.2 consumer 端 JMX 监控	316
8.5 JVM 监控	317
8.5.1 进程状态	318
8.5.2 GC 性能	318
8.6 OS 监控	318
8.7 主流监控框架	319
8.7.1 JmxTool	320
8.7.2 kafka-manager	320
8.7.3 Kafka Monitor	325
8.7.4 Kafka Offset Monitor	327
8.7.5 CruiseControl	329
8.8 本章小结	330

第 9 章 调优 Kafka 集群.....	331
9.1 引言	331
9.2 确定调优目标	333
9.3 集群基础调优	334
9.3.1 禁止 atime 更新	335
9.3.2 文件系统选择	335
9.3.3 设置 swapiness	336
9.3.4 JVM 设置	337
9.3.5 其他调优	337
9.4 调优吞吐量	338
9.5 调优延时	342
9.6 调优持久性	343
9.7 调优可用性	347
9.8 本章小结	349
第 10 章 Kafka Connect 与 Kafka Streams	350
10.1 引言	350
10.2 Kafka Connect	351
10.2.1 概要介绍	351
10.2.2 standalone Connect	353
10.2.3 distributed Connect	356
10.2.4 开发 connector	359
10.3 Kafka Streams	362
10.3.1 流处理	362
10.3.2 Kafka Streams 核心概念	364
10.3.3 Kafka Streams 与其他框架的异同	368
10.3.4 Word Count 实例	369
10.3.5 Kafka Streams 应用开发	372
10.3.6 Kafka Streams 状态查询	382
10.4 本章小结	386

第 1 章

认识 Apache Kafka

随着大数据时代的到来，数据中蕴含的价值日益得到展现，仿佛一座待人挖掘的金矿，引来无数的掘金者。但随着数据量越来越大，如何实时准确地收集并分析数据成为摆在所有从业人员面前的难题。

本章作为本书的第 1 章，将带领读者对 Apache Kafka 系统及其生态圈建立一个宏观的概念和认识。同时，本章将结合消息引擎系统的相关知识与设计理念，循序渐进地对 Kafka 系统的设计架构和相关概念进行展开，并给出简单示例以快速上手 Kafka。

学习本章，你将了解到以下内容。

- 消息引擎系统的定义与特点。
- Apache Kafka 系统概要设计及术语。
- Apache Kafka 快速入门。

1.1 Kafka 快速入门

Kafka 的核心功能是什么？一言以蔽之，高性能的消息发送与高性能的消息消费。接下来，本书打算先给出 Kafka 的快速入门教程，即 Kafka 的“Hello, world”示例。这么做与其他书不同，其目的就是可以让各位读者快速体验 Kafka 的核心功能——消息的发送与消费。这对于很多不太熟悉 Kafka 的初学者来说，可以让他们瞬间提升学习的快感，从而增强对于掌握与理解 Kafka 系统的信心。

作为消息引擎系统中的佼佼者，Kafka 诸多独特的设计理念及强大的功能特性非常值得学习，但现在我们首先跑通一个 Kafka 的简单示例，切身感受一下 Kafka 这个目前消息引擎领域