

华晟经世ICT专业群系列教材

Hadoop

大数据平台

集群部署与开发

罗文浪 邱波 郭炳宇 姜善永 主编



中国工信出版集团



人民邮电出版社
POSTS & TELECOM PRESS

华晟经世ICT专业群系列教材

Hadoop

大数据平台

集群部署与开发

罗文浪 邱波 郭炳宇 姜善永 主编

人民邮电出版社
北京

图书在版编目 (C I P) 数据

Hadoop大数据平台集群部署与开发 / 罗文浪等主编

— 北京 : 人民邮电出版社, 2018. 12

华晟经世ICT专业群系列教材

ISBN 978-7-115-49417-7

I. ①H… II. ①罗… III. ①数据处理软件—程序设计—教材 IV. ①TP274

中国版本图书馆CIP数据核字(2018)第224492号

内 容 提 要

本教材一共6个项目,项目1为Hadoop导入,主要介绍了Hadoop的作用、特点、发展情况,并详细介绍了Hadoop伪分布式搭建及使用方法;项目2主要对Hadoop的核心元素、接口操作进行了细致讲解;项目3对为实现Hadoop HA所需的Zookeeper的架构、部署等进行了解释;项目4至项目6详细介绍了Hadoop生态圈中的几个核心组件——分布式存储数据库(HBase)、数据迁移神器(Sqoop)、数据采集神器(Flume)以及数据仓库(Hive),在介绍这几个核心组件的同时也融入了对于大数据综合分析的分析。本教材具有较强实用性,教材内容以“学”和“导学”交织呈现,十分适合学习者使用。

◆ 主 编 罗文浪 邱 波 郭炳宇 姜善永

责任编辑 李 静

责任印制 彭志环

◆ 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>

固安县铭成印刷有限公司印刷

◆ 开本: 787×1092 1/16

印张: 12.5

2018年12月第1版

字数: 292千字

2018年12月河北第1次印刷

定价: 49.00元

读者服务热线: (010)81055488 印装质量热线: (010)81055316

反盗版热线: (010)81055315

前言

在这样一个数据信息时代，以云计算、大数据、物联网为代表的新一代信息技术已经受到空前的关注，教育战略服务国家战略，相关的职业教育急需升级以顺应和助推产业发展。

在这本教材的编写中，我们在内容上贯穿以“学习者”为中心的设计理念——教学目标以任务驱动，教材内容以“学”和“导学”交织呈现，项目引入以情景化的职业元素构成，学习足迹借助图谱得以可视化，学习效果通过最终的创新项目得以校验。

本教材一共分为6个项目，第1个项目为Hadoop导入，主要是关于Hadoop作用、特点、发展的介绍，详细介绍了Hadoop伪分布式搭建及使用。第2个项目主要对Hadoop的核心元素、接口操作进行讲解。第3个项目对为实现Hadoop HA所需的Zookeeper的架构、部署等进行介绍。第4至第6个项目详细介绍了Hadoop生态圈中的几个核心组件：分布式存储数据库（HBase）、数据迁移神器（Sqoop）、数据采集神器（Flume）以及数据仓库（Hive），介绍这几个核心组件的同时也夹杂了对大数据的综合实验分析。

本教材具有以下特点。

1. 教材内容的组织强调以学习行为为主线，构建了“学”与“导学”的内容逻辑。“学”是主体内容，包括项目描述、任务解决及项目总结；“导学”是引导学生自主学习、独立实践的部分，包括项目引入、交互窗口、思考练习、拓展训练及双创项目。

2. 情景剧式的项目引入。教材中每个项目都模拟了一个完整的项目团队，以情景剧的形式开篇，并融入职业元素，让内容更加接近行业、企业和生产实际。项目还原工作场景，展示项目进程，嵌入岗位、行业认知，融入工作的方法和技巧，更多地传递一种解决问题的思路 and 理念。

3. 项目篇章以项目为核心载体，强调知识输入，经过任务的解决与训练，再到技能输出，采用“两点（知识点、技能点）”“两图（知识图谱、技能图谱）”的方式梳理知识、技能，项目开篇清晰地描绘出该项目所覆盖的和需要的知识点，项目最后总结出经过任

务训练学生所能获得的技能图谱。

4. 教材强调动手和实操，以解决任务为驱动，实现“做中学、学中做”的目标。任务驱动式的学习可以让学生遵循一般的学习规律，由简到难，循环往复，融会贯通；加强实践、动手训练，在实操中学习更加直观和深刻地了解知识内容；融入最新技术应用，结合真实应用场景，解决现实性客户需求。

5. 教材中有创新的双创项目设计。教材结尾设计双创项目与其他教材形成呼应，体现了项目的完整性、创新性和挑战性，既能培养学生面对困难勇于挑战的创业意识，又能培养学生使用新技术解决问题的创新精神。

本教材由罗文浪、邱波、郭炳宇、姜善永老师主编。主编除了参与编写外，还负责拟定大纲并总体编纂。本教材执笔人依次是：项目1 罗文浪，项目2 邱波，项目3至项目5 朱胜，项目6 黎正林。本教材初稿完结后，由郭炳宇、姜善永、王田甜、苏尚停、王雪松、刘静、张瑞元、朱胜、李慧蕾、杨慧东、唐斌、何勇、李文强、范雪梅、冉芬、曹利洁、张静、蒋平新、赵艳慧、杨晓蕊、刘红申、黎正林、李想组成的编审委员会的相关成员对教材内容进行审核和修订。

整本教材从开发和总体设计到每个细节都包含了我们整个团队的协作和细心打磨过程，我们希望以专业的精神尽量克服知识和经验的不足，并以此书飨慰读者。

本教材配套代码链接：<http://114.115.179.78/teaching-resources/Hadoop.zip>

本教材配套 PPT 链接：<http://114.115.179.78/teaching-resources/PPT-Hadoop.zip>

编 者

2018 年 7 月

■ ■ ■ 目 录 |

项目 1 搭建 Hadoop 开发环境	1
1.1 任务一：Hadoop 简介	2
1.1.1 Hadoop 介绍	2
1.1.2 Hadoop 的发展历史及现状	3
1.1.3 任务回顾	5
1.2 任务二：搭建 Hadoop 伪分布式环境	6
1.2.1 准备工作	6
1.2.2 搭建伪分布式环境	13
1.2.3 Hadoop 测试	23
1.2.4 任务回顾	26
1.3 项目总结	27
1.4 拓展训练	28
项目 2 Hadoop 入门及实战	29
2.1 任务一：HDFS 体系结构与基本原理	30
2.1.1 HDFS 概述	30
2.1.2 HDFS 核心元素及其原理	32
2.1.3 任务回顾	38
2.2 任务二：HDFS 接口操作	39
2.2.1 Shell 接口操作	39
2.2.2 Java 接口操作	41

2.2.3	任务回顾	47
2.3	任务三：MapReduce 开发实战	48
2.3.1	MapReduce 工作机制	48
2.3.2	MapReduce 开发实战	54
2.3.3	任务回顾	63
2.4	项目总结	64
2.5	拓展训练	65
项目 3	搭建 Zookeeper 运行环境	67
3.1	任务一：Zookeeper 概述	68
3.1.1	Zookeeper 原理	68
3.1.2	Zookeeper 系统架构	70
3.1.3	任务回顾	71
3.2	任务二：Zookeeper 集群搭建	72
3.2.1	集群规划	72
3.2.2	安装 Zookeeper 集群	74
3.2.3	任务回顾	79
3.3	任务三：使用 Zookeeper 来实现 Hadoop 的高可用性	79
3.3.1	Zookeeper 集群与 Hadoop 高可用性	79
3.3.2	Hadoop 高可用性集群部署	81
3.3.3	任务回顾	92
3.4	项目总结	93
3.5	拓展训练	93
项目 4	分布式存储数据库	95
4.1	任务一：HBase 概述	96
4.1.1	HBase 简介	96
4.1.2	HBase 表结构	97
4.1.3	HBase 核心进程	100

4.1.4	HBase 系统架构	103
4.1.5	任务回顾	105
4.2	任务二：HBase 集群部署	106
4.2.1	HBase 单节点部署	106
4.2.2	HBase 集群部署	108
4.2.3	任务回顾	112
4.3	任务三：HBase 实战	112
4.3.1	HBase Shell	112
4.3.2	HBase Java	116
4.3.3	任务回顾	130
4.4	项目总结	131
4.5	拓展训练	132
项目 5	数据迁移和数据采集	133
5.1	任务一：数据迁移神器——Sqoop	134
5.1.1	Sqoop 概述	134
5.1.2	Sqoop 部署	135
5.1.3	Sqoop 实战	136
5.1.4	任务回顾	142
5.2	任务二：数据采集神器——Flume	143
5.2.1	Flume 概述	143
5.2.2	Flume 部署	148
5.2.3	Flume 实战	150
5.2.4	任务回顾	156
5.3	项目总结	157
5.4	拓展训练	157
项目 6	数据分析	159
6.1	任务一：Hive 概述	160

6.1.1	Hive 介绍	160
6.1.2	Hive 架构及原理分析	161
6.1.3	Hive 数据类型	163
6.1.4	Hive 表类型	165
6.1.5	任务回顾	168
6.2	任务二：Hive 部署与实战	168
6.2.1	Hive 部署	169
6.2.2	Hive 表操作	175
6.2.3	Hive 数据分析	187
6.2.4	任务回顾	190
6.3	项目总结	191
6.4	拓展训练	192

项目 1

搭建 Hadoop 开发环境

项目引入

我叫 Snkey，是一名大数据分析师，在一家电商公司工作，每天和大量的数据打交道；我的同事 Windy 负责公司项目的系统运维，监管公司各种网络问题；而 Suzan 则是一名资深的 Java 开发工程师。我们分管不同的部门，但是因为一个项目我们彼此有了交集。

一次会议上，大 Boss 提出公司要成立一个专门处理大数据的团队。

Boss：现在大数据越来越流行了，可谓遍地开花，随着公司业务的蒸蒸日上，公司的后台数据日益增长，在大数据时代，数据就是黄金，我们也不能落伍，我们是否考虑把大数据技术引入进来。

我：目前积累的数据已经很庞大了，随着时间的推移，数据量只会越来越庞大，传统的数据架构和计算框架在处理一些复杂的业务上已经捉襟见肘，要想挖掘出大量数据中的有用价值，必须运用大数据技术，例如 Hadoop、MapReduce、Hive 等。

Boss 对这个想法大加赞赏，当场就拍板由我们三个人成立一个大数据团队。

知识图谱

项目 1 的知识图谱如图 1-1 所示。

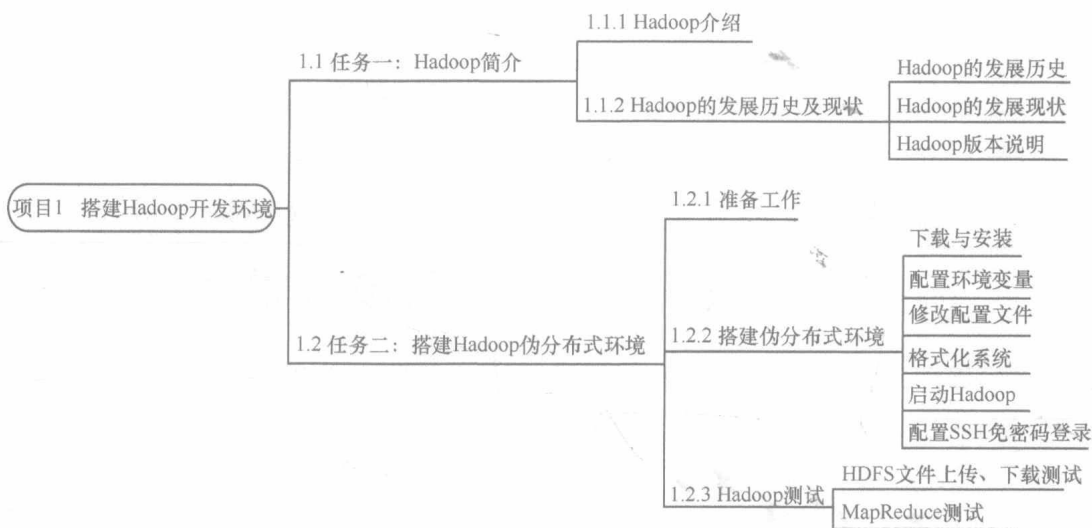


图1-1 项目1知识图谱

1.1 任务一：Hadoop 简介

【任务描述】

经过前期的调研，我们知道大数据涉及 Linux、Java 和数据分析 3 方面的内容，为了方便学习大数据技术，我们决定在公司的服务器上部署 Hadoop 集群，并为每个人搭建一个 Hadoop 伪分布式环境，方便进行本地测试。

为了方便大家了解大数据，所以在第一个任务中，我们不会涉及非常复杂的内容，仅对大数据的整体进行介绍。

1.1.1 Hadoop介绍

谈到大数据，就不得不提与它关系紧密的分布式系统基础架构——Hadoop，它的 Logo 是一头大象，那么 Hadoop 究竟是怎么产生的呢？它又为什么会被命名为 Hadoop 呢？下面为大家娓娓道来。

Hadoop 的创作者是 Doug Cutting，他受 Google 三篇论文（GFS、MapReduce、BigTable）的启发而开创了 Hadoop 项目，Hadoop 最初是 Doug 的女儿给玩具起的名字，后来，Doug Cutting 将其用作自己项目的名称，目前 Hadoop 项目属于 Apache 基金会的一个顶级开源项目。

Hadoop 主要用于解决两个问题：海量数据存储和海量数据分析。

Hadoop 就是为处理海量数据而生，它本质上是一个能够对大量数据进行分布式处理的软件框架，并且以一种可靠、高效、可伸缩的方式进行数据处理。因此它具有以下几个方面的特性。

① 高可扩展性：Hadoop 可以使用大量的普通计算机来完成专业服务器才能完成的计算工作，我们可以很方便地根据实际业务需求，横向扩展集群机器数量。

② 高容错性：Hadoop 可以将数据按照 Block 进行存储，而且每一个 Block 都自动保存多个副本，保证数据不会丢失。对于执行失败的任务能够进行重新分配执行。

③ 扩容能力强：Hadoop 能够可靠地存储和处理十亿兆字节（PB）的数据。

④ 成本低：可以通过总计数千个节点的普通机器组成的服务器群来分发以及处理数据。

⑤ 高效率：Hadoop 通过分发数据，可以在数据所在的节点上并行地进行处理，处理速度非常快。

⑥ 高可靠性：Hadoop 能自动地维护数据的多份副本，并且在任务失败后能自动地重新部署计算任务。

1.1.2 Hadoop的发展历史及现状

1. Hadoop 的发展历史

Hadoop 源自 2002 年的一个开源项目 Apache Nutch。其最初的雏形是由 Apache Lucene 项目的创始人 Doug Cutting 开发的文本搜索库。2004 年，Nutch 项目模仿 Google 文件系统（Google File System，GFS）开发了自己的分布式文件系统（Nutch Distributed File System，NDFS），也就是 HDFS 的前身。同年，谷歌公司发表了另一篇具有深远影响的论文，阐述了 MapReduce 分布式编程思想。

2005 年，Nutch 项目团队参考 MapReduce 分布式编程思想开发了 MapReduce 分布式处理框架。

2006 年 2 月，NDFS 和 MapReduce 从 Nutch 项目独立出来，成为 Lucene 项目的一个子项目，被命名为 Hadoop。

2008 年 1 月，Hadoop 正式成为 Apache 顶级项目，并逐渐被雅虎、FaceBook 等大公司采用。

2008 年 4 月，Hadoop 打破世界纪录，成为排序 1TB 数据最快的系统，它采用一个由 910 个节点构成的集群进行运算，排序时间只用了 209 s。到 2009 年，这个排序时间缩短到了 62 s。由此，Hadoop 迅速跃升为最具影响力的开源分布式开发平台，在大数据时代获得大量拥趸。

2. Hadoop 的发展现状

Hadoop 凭借其实用性、易用性，自推出以来在几年间就满足了大部分工业界的应用需求，还引起了学术界对其的广泛关注和研究。Hadoop 已然成为目前大数据处理主流技术和系统平台。不夸张地说，Hadoop 现在已经成为大数据处理的潜在标准，并在工业界，尤其是互联网行业中得到大量频繁的进一步开发和改进。

Yahoo 作为 Hadoop 曾经的最大的支持者，截至 2012 年，其 Hadoop 机器总节点数目超过 420000 个，有超过 10 万的核心 CPU 在运行 Hadoop。最大的一个单 Master 节点集群有 4500 个节点，总的集群存储容量大于 350PB，每月提交的作业数目超过 1000 万个。Facebook 也使用 Hadoop 存储内部日志与多维数据，并以此作为报告、分析和机器学习的数据源。Facebook 不仅是 Hadoop 的忠实用户，同时它还在 Hadoop 基础上建立了一个名为 Hive 的高级数据仓库框架，用来进行数据清洗、处理等工作，目前 Hive 已经成为基于 Hadoop 的 Apache 一级项目。

在国内，很多大型互联网企业都逐渐开始使用 Hadoop 来处理离线数据，例如阿里巴巴、百度、淘宝、网易等。阿里巴巴的 Hadoop 集群数据覆盖了它的诸多业务线，非常庞大，它需要为淘宝、支付宝、聚划算等提供底层的存储和基础计算服务。仅看 2012 年的数据，其集群已经有超过 3200 台服务器，总的存储容量超过 60PB，每天的作业数目超过 1500000 个，到今天，这些数字只会越来越大。腾讯也是使用 Hadoop 最早的中国互联网公司之一，由于腾讯的用户量庞大，因此其集群数量也是非常庞大的。腾讯的社交广告平台、腾讯微博、QQ、财付通、微信、QQ 音乐等平台都要依靠 Hadoop 进行存储和计算。除此之外，腾讯还利用 Hadoop-Hive 构建了一套自己的数据仓库系统，取名为“TDW”。

随着互联网行业的发展，Hadoop 也在不断地被应用和升级，相信在未来，它的应用领域和范围还会逐渐增大。

3. Hadoop 版本说明

Hadoop 发展至今，主要的版本有两代，我们习惯将第一代 Hadoop 称为 Hadoop 1.0，第二代 Hadoop 称为 Hadoop 2.0，其版本演变如图 1-2 所示。

第一代 Hadoop 包含 3 个大版本，分别是 0.20.x，0.21.x 和 0.22.x，其中，0.20.x 最后演化成 1.0.x，变成了稳定版，而 0.21.x 和 0.22.x 则增加了 NameNode HA 等新的重大特性。

第二代 Hadoop 包含两个版本，分别是 0.23.x 和 2.x，它们完全不同于 Hadoop 1.0，是一套全新的架构，包含 HDFS Federation 和 YARN 两个系统。0.23.x 和 2.x 两个版本在结构上也有重大的区别，相比于 0.23.x，2.x 增加了 NameNode HA 和 Wire-compatibility 两个重大特性。

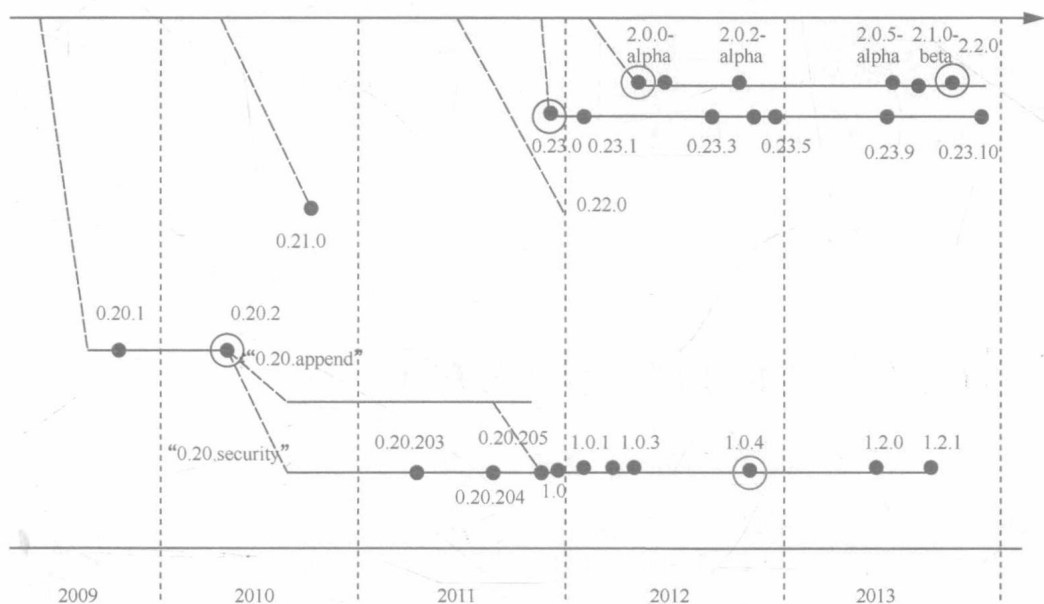


图1-2 Hadoop版本演变

1.1.3 任务回顾

知识点总结

1. Hadoop 起源及其作用。
2. Hadoop 的优势。
3. Hadoop 的发展历史及版本说明。
4. Hadoop 的发展现状。

学习足迹

项目1 任务一的学习足迹如图1-3所示。

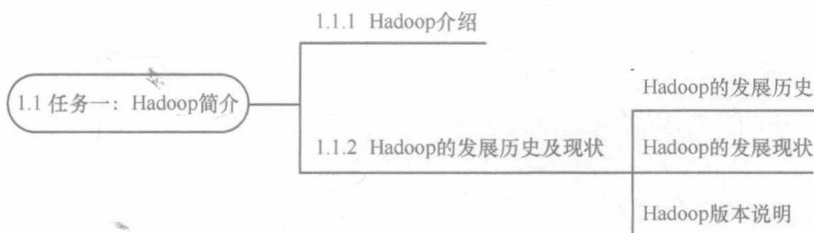


图1-3 项目1任务一学习足迹



思考与练习

1. Hadoop 主要解决什么问题？
2. 简述 Hadoop 的几大优势。
3. Hadoop 有几个版本，主要区别是什么？

1.2 任务二：搭建 Hadoop 伪分布式环境

【任务描述】

在认知了大数据之后，我们是不是也来体验一下 Hadoop？但是我们还缺少环境条件，如果直接搭建 Hadoop 生产环境（高可用集群）就稍显复杂了，不如搭建一个 Hadoop 伪分布式环境，这也不妨碍我们的体验。而且搭建 Hadoop 伪分布式环境和搭建生产环境有很多相通之处，学习了搭建 Hadoop 伪分布式环境同样有利于我们后续搭建 Hadoop 生产环境。

1.2.1 准备工作

搭建 Hadoop 伪分布式环境，需要在单个节点上进行部署。在安装 Hadoop 之前，我们需要安装 Hadoop 的运行环境——Linux 系统，本教材中选择安装的是 CentOS7 mini server 版本。我们可以通过 VMWare、VirtualBox 等虚拟化软件来创建部署所需要的虚拟机，安装过程略。需要注意的是，在安装时需要配置虚拟机的网卡信息，我们选择桥接网卡，这样虚拟机与虚拟机、虚拟机与主机之间都可以进行通信，同时也方便在虚拟机中下载安装所需要的资源。

配置好网卡后，建议测试一下网络环境是否存在问题，测试代码如下：

【代码 1-1】 测试网络环境

```
[root@huatec01 ~]# ping baidu.com
PING baidu.com (220.181.57.217) 56(84) Byte of data.
64 Byte from 220.181.57.217: icmp_seq=1 ttl=57 time=5.10 ms
64 Byte from 220.181.57.217: icmp_seq=2 ttl=57 time=4.28 ms
...
[root@huatec01 ~]# ping 192.168.14.103
PING 192.168.14.103 (192.168.14.103) 56(84) Byte of data.
```

```
64 Byte from 192.168.14.103: icmp_seq=1 ttl=64 time=0.929 ms
64 Byte from 192.168.14.103: icmp_seq=2 ttl=64 time=0.256 ms
...
```

其中, PING baidu.com 用于测试外网环境, PING 192.168.14.103 用于测试内部网络环境, 通过上面的代码我们看出, 测试通过了, 这表明我们的虚拟机的网卡配置信息是正确的。

同时, 我们还需要关闭系统的防火墙, CentOS 7 默认关闭 iptables, 关闭 firewalld 防火墙和 selinux 防火墙即可, 代码如下所示:

【代码 1-2】 关闭并查看 firewalld 防火墙

```
[root@huatec01 ~]# systemctl stop firewalld
[root@huatec01 ~]# systemctl disable firewalld
[root@huatec01 ~]# systemctl status firewalld
firewalld.service - firewalld - dynamic firewall daemon
    Loaded: loaded (/usr/lib/systemd/system/firewalld.service;
disabled)
    Active: inactive (dead)
    Oct 12 22:54:57 huatec01 systemd[1]: Stopped firewalld - dynamic
firewall daemon.
```

在关闭 firewalld 防火墙之前, 需要先执行 “systemctl stop firewalld” 命令来停止防火墙, 避免其已经处于运行状态, 导致关闭失败, 然后调用 systemctl disable firewalld 让其彻底不可用。最后, 执行 systemctl status firewalld 指令查看防火墙是否关闭成功。从上述的代码中可以看到最后的防火墙状态是 inactive (dead), 说明操作成功。

关闭 selinux 防火墙, 代码如下:

【代码 1-3】 关闭 selinux 防火墙

```
[root@huatec01 ~]# vi /etc/sysconfig/selinux
# This file controls the state of SELinux on the system.
# SELINUX= can take one of these three values:
#     enforcing - SELinux security policy is enforced.
#     permissive - SELinux prints warnings instead of enforcing.
#     disabled - No SELinux policy is loaded.
SELINUX=disabled
```

```
# SELINUXTYPE= can take one of three two values:
#     targeted - Targeted processes are protected,
#     minimum - Modification of targeted policy. Only selected
processes are protected.
#     mls - Multi Level Security protection.
SELINUXTYPE=targeted
```

我们修改其中的一行，将 SELINUX 的值改为 disabled 即可。从上面的注释中可以看到 SELINUX 的取值有 3 个，分别为 enforcing、permissive 和 disabled。enforcing 表示 SELINUX 安全策略是强制性的；permissive 表示 SELINUX 安全策略将会提示权限问题，输出提示信息；disabled 是直接让其不可用。

【知识引申】

如果是 CentOS 6.0，参考如下方式关闭防火墙。

第一步：关闭 iptables 和 ip6tables

查看防火墙状态

```
service iptables status
```

```
service ip6tables status
```

关闭防火墙

```
service iptables stop
```

```
service ip6tables stop
```

查看防火墙开机启动状态

```
chkconfig iptables - list
```

```
chkconfig ip6tables --list
```

关闭防火墙开机启动

```
chkconfig iptables off
```

```
chkconfig ip6tables off
```

在生产环境中，需要为每个虚拟机设置固定的 ip 和主机名，即使是搭建伪分布式环境，也建议这么做。执行 `vim /etc/sysconfig/network-scripts/ifcfg-eth0` 指令，打开 ip 配置文件进行编辑，编辑后的文件结果如下：