

狂神聊ElasticSearch

聊聊Doug Cutting

1998年9月4日，Google公司在美国硅谷成立。正如大家所知，它是一家做搜索引擎起家的公司。



无独有偶，一位名叫**Doug Cutting**的美国工程师，也迷上了搜索引擎。他做了一个用于文本搜索的函数库（姑且理解为软件的功能组件），命名为**Lucene**。



左为Doug Cutting，右为Lucene的LOGO

Lucene是用JAVA写成的，目标是为各种中小型应用软件加入全文检索功能。因为好用而且开源（代码公开），非常受程序员们的欢迎。

早期的时候，这个项目被发布在Doug Cutting的个人网站和SourceForge（一个开源软件网站）。后来，2001年底，Lucene成为**Apache软件基金会**jakarta项目的一个子项目。



Apache软件基金会，搞IT的应该都认识

2004年，Doug Cutting再接再厉，在Lucene的基础上，和Apache开源伙伴Mike Cafarella合作，开发了一款可以代替当时的主流搜索的开源搜索引擎，命名为**Nutch**。



Nutch是一个建立在Lucene核心之上的网页搜索应用程序，可以下载下来直接使用。它在Lucene的基础上加了网络爬虫和一些网页相关的功能，目的就是从一个简单的站内检索推广到全球网络的搜索上，就像Google一样。

Nutch在业界的影响力比Lucene更大。

大批网站采用了Nutch平台，大大降低了技术门槛，使低成本的普通计算机取代高价的Web服务器成为可能。甚至有一段时间，在硅谷有了一股用Nutch低成本创业的潮流。

随着时间的推移，无论是Google还是Nutch，都面临搜索对象“体积”不断增大的问题。

尤其是Google，作为互联网搜索引擎，需要存储大量的网页，并不断优化自己的搜索算法，提升搜索效率。



Google搜索栏

在这个过程中，Google确实找到了不少好办法，并且无私地分享了出来。

2003年，Google发表了一篇技术学术论文，公开介绍了自己的谷歌文件系统**GFS (Google File System)**。这是Google公司为了存储海量搜索数据而设计的专用文件系统。

第二年，也就是2004年，Doug Cutting基于Google的GFS论文，实现了**分布式文件存储系统**，并将它命名为**NDFS (Nutch Distributed File System)**。



还是2004年，Google又发表了一篇技术学术论文，介绍自己的**MapReduce编程模型**。这个编程模型，用于大规模数据集（大于1TB）的并行分析运算。

第二年（2005年），Doug Cutting又基于MapReduce，在Nutch搜索引擎实现了该功能。



2006年，当时依然很厉害的**Yahoo (雅虎) 公司**，招安了Doug Cutting。



这里要补充说明一下雅虎招安Doug的背景：2004年之前，作为互联网开拓者的雅虎，是使用Google搜索引擎作为自家搜索服务的。在2004年开始，雅虎放弃了Google，开始自己研发搜索引擎。所以。。。

加盟Yahoo之后，Doug Cutting将NDFS和MapReduce进行了升级改造，并重新命名为**Hadoop**（NDFS也改名为HDFS，Hadoop Distributed File System）。

这个，就是后来大名鼎鼎的大数据框架系统——Hadoop的由来。而Doug Cutting，则被人们称为**Hadoop之父**。



Hadoop这个名字，实际上是Doug Cutting他儿子的黄色玩具大象的名字。所以，Hadoop的Logo，就是一只奔跑的黄色大象。



我们继续往下说。

还是2006年，Google又发论文了。

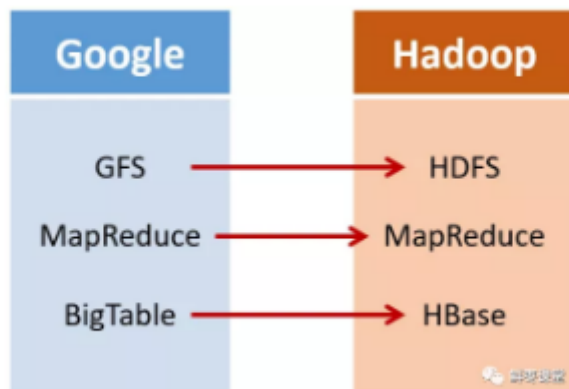
这次，它们介绍了自己的**BigTable**。这是一种分布式数据存储系统，一种用来处理海量数据的非关系型数据库。

Doug Cutting当然没有放过，在自己的hadoop系统里面，引入了BigTable，并命名为**HBase**。



好吧，反正就是紧跟Google时代步伐，你出什么，我学什么。

所以，Hadoop的核心部分，基本上都有Google的影子。



2008年1月，Hadoop成功上位，正式成为Apache基金会的顶级项目。

同年2月，Yahoo宣布建成了一个拥有1万个内核的Hadoop集群，并将自己的搜索引擎产品部署在上面。

7月，Hadoop打破世界纪录，成为最快排序1TB数据的系统，用时209秒。

[回到主题](#)

Lucene是一套信息检索工具包，并不包含搜索引擎系统，它包含了索引结构、读写索引工具、相关性工具、排序等功能，因此在使用Lucene时仍需要关注搜索引擎系统，例如数据获取、解析、分词等方面的东西。

为什么要给大家介绍下Lucene呢，因为我们学过的solr 和 即将要学习的elasticsearch都是基于该工具包做的一些封装和增强罢了~

按照我的风格，一般在开始都会给大家聊聊历史！学习更重要的是培养编程的兴趣，而不是只会使用一些基本的API！

狂神的教学风格，免费，授人以渔，还有我那句每个课程都会说的至理名言：只要学不死，就往死里学！

ElasticSearch概述

Elasticsearch，简称为es，es是一个开源的高扩展的分布式全文检索引擎，它可以近乎实时的存储、检索数据；本身扩展性很好，可以扩展到上百台服务器，处理PB级别的数据。es也使用Java开发并使用Lucene作为其核心来实现所有索引和搜索的功能，但是它的目的是通过简单的RESTful API来隐藏Lucene的复杂性，从而让全文搜索变得简单。

据国际权威的数据库产品评测机构DB Engines的统计，在2016年1月，ElasticSearch已超过Solr等，成为排名第一的搜索引擎类应用。

[历史](#)

多年前，一个叫做Shay Banon的刚结婚不久的失业开发者，由于妻子要去伦敦学习厨师，他便跟着也去了。在他找工作的过程中，为了给妻子构建一个食谱的搜索引擎，他开始构建一个早期版本的Lucene。

直接基于Lucene工作会比较困难，所以Shay开始抽象Lucene代码以便Java程序员可以在应用中添加搜索功能。他发布了他的第一个开源项目，叫做“Compass”。

后来Shay找到一份工作，这份工作处在高性能和内存数据网格的分布式环境中，因此高性能的、实时的、分布式的搜索引擎也是理所当然需要的。然后他决定重写Compass库使其成为一个独立的服务叫做Elasticsearch。

第一个公开版本出现在2010年2月，在那之后Elasticsearch已经成为Github上最受欢迎的项目之一，代码贡献者超过300人。一家主营Elasticsearch的公司就此成立，他们一边提供商业支持一边开发新功能，不过Elasticsearch将永远开源且对所有人可用。

Shay的妻子依旧等待着她的食谱搜索.....

谁在使用：

- 1、维基百科，类似百度百科，全文检索，高亮，搜索推荐/2
 - 2、The Guardian（国外新闻网站），类似搜狐新闻，用户行为日志（点击，浏览，收藏，评论）+社交网络数据（对某某新闻的相关看法），数据分析，给到每篇新闻文章的作者，让他知道他的文章的公众反馈（好，坏，热门，垃圾，鄙视，崇拜）
 - 3、Stack Overflow（国外的程序异常讨论论坛），IT问题，程序的报错，提交上去，有人会跟你讨论和回答，全文检索，搜索相关问题和答案，程序报错了，就会将报错信息粘贴到里面去，搜索有没有对应的答案
 - 4、GitHub（开源代码管理），搜索上千亿行代码
 - 5、电商网站，检索商品
 - 6、日志数据分析，logstash采集日志，ES进行复杂的数据分析，ELK技术，elasticsearch+logstash+kibana
 - 7、商品价格监控网站，用户设定某商品的价格阈值，当低于该阈值的时候，发送通知消息给用户，比如说订阅牙膏的监控，如果高露洁牙膏的家庭套装低于50块钱，就通知我，我就去买
 - 8、BI系统，商业智能，Business Intelligence。比如说有个大型商场集团，BI，分析一下某某区域最近3年的用户消费金额的趋势以及用户群体的组成构成，产出相关的数张报表，**区，最近3年，每年消费金额呈现100%的增长，而且用户群体85%是高级白领，开一个新商场。ES执行数据分析和挖掘，Kibana进行数据可视化
 - 9、国内：站内搜索（电商，招聘，门户，等等），IT系统搜索（OA，CRM，ERP，等等），数据分析（ES热门
- 的一个使用场景）

ES和solr的差别

Elasticsearch简介

Elasticsearch是一个实时分布式搜索和分析引擎。它让你以前所未有的速度处理大数据成为可能。

它用于全文搜索、结构化搜索、分析以及将这三者混合使用：

维基百科使用Elasticsearch提供全文搜索并高亮关键字，以及输入实时搜索(search-as-you-type)和搜索纠错(did-you-mean)等搜索建议功能。

英国卫报使用Elasticsearch结合用户日志和社交网络数据提供给他们编辑以实时的反馈，以便及时了解公众对新发表的文章的回应。

StackOverflow结合全文搜索与地理位置查询，以及more-like-this功能来找到相关的问题和答案。

Github使用Elasticsearch检索1300亿行的代码。

但是Elasticsearch不仅用于大型企业，它还让像DataDog以及Klout这样的创业公司将最初的想法变成可

扩展的解决方案。Elasticsearch可以在你的笔记本上运行，也可以在数以百计的服务器上处理PB级别的数据。

Elasticsearch是一个基于Apache Lucene(TM)的开源搜索引擎。无论在开源还是专有领域，Lucene可以被认为是迄今为止最先进、性能最好的、功能最全的搜索引擎库。

但是，Lucene只是一个库。想要使用它，你必须使用Java来作为开发语言并将其直接集成到你的应用中，更糟糕的是，Lucene非常复杂，你需要深入了解检索的相关知识来理解它是如何工作的。

Elasticsearch也使用Java开发并使用Lucene作为其核心来实现所有索引和搜索的功能，但是它的目的是通过简单的 RESTful API 来隐藏Lucene的复杂性，从而让全文搜索变得简单。

Solr简介

Solr 是Apache下的一个顶级开源项目，采用Java开发，它是基于Lucene的全文搜索服务器。Solr提供了比Lucene更为丰富的查询语言，同时实现了可配置、可扩展，并对索引、搜索性能进行了优化

Solr可以独立运行，运行在Jetty、Tomcat等这些Servlet容器中，Solr 索引的实现方法很简单，用 POST 方法向 Solr 服务器发送一个描述 Field 及其内容的 XML 文档，Solr根据xml文档添加、删除、更新索引。Solr 搜索只需要发送 HTTP GET 请求，然后对 Solr 返回Xml、json等格式的查询结果进行解析，组织页面布局。Solr不提供构建UI的功能，Solr提供了一个管理界面，通过管理界面可以查询Solr的配置和运行情况。

solr是基于lucene开发企业级搜索服务器，实际上就是封装了lucene。

Solr是一个独立的企业级搜索应用服务器，它对外提供类似于Web-service的API接口。用户可以通过http请求，向搜索引擎服务器提交一定格式的文件，生成索引；也可以通过提出查找请求，并得到返回结果。

Lucene简介

Lucene是apache软件基金会4 jakarta项目组的一个子项目，是一个开放源代码的全文检索引擎工具包，但它不是一个完整的全文检索引擎，而是一个全文检索引擎的架构，提供了完整的查询引擎和索引引擎，部分文本分析引擎（英文与德文两种西方语言）。Lucene的目的是为软件开发人员提供一个简单易用的工具包，以方便的在目标系统中实现全文检索的功能，或者是以此为基础建立起完整的全文检索引擎。Lucene是一套用于全文检索和搜寻的开源程式库，由Apache软件基金会支持和提供。Lucene提供了一个简单却强大的应用程序接口，能够做全文索引和搜寻。在Java开发环境里Lucene是一个成熟的免费开源工具。就其本身而言，Lucene是当前以及最近几年最受欢迎的免费Java信息检索程序库。人们经常提到信息检索程序库，虽然与搜索引擎有关，但不应该将信息检索程序库与搜索引擎相混淆。

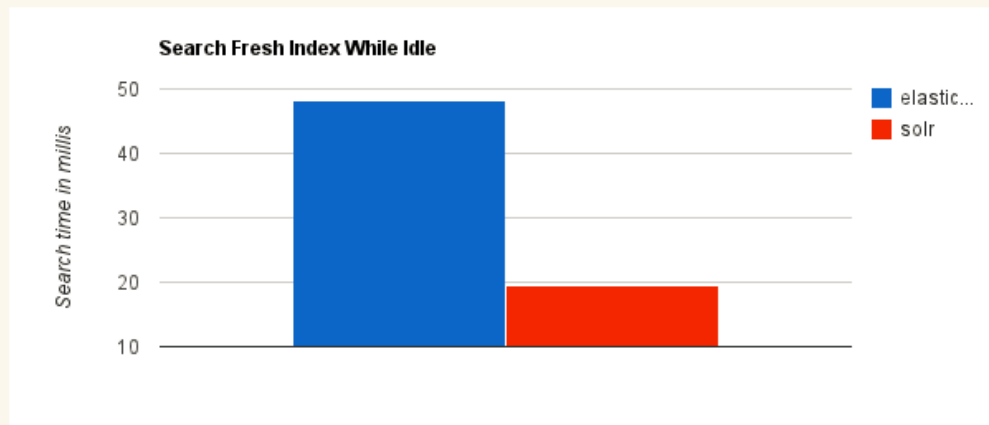
Lucene是一个全文检索引擎的架构。那什么是全文搜索引擎？

全文搜索引擎是名副其实的搜索引擎，国外具代表性的有Google、Fast/AllTheWeb、AltaVista、Inktomi、Teoma、WiseNut等，国内著名的有百度（Baidu）。它们都是通过从互联网上提取的各个网站的信息（以网页文字为主）而建立的数据库中，检索与用户查询条件匹配的相关记录，然后按一定的排列顺序将结果返回给用户，因此他们是真正的搜索引擎。

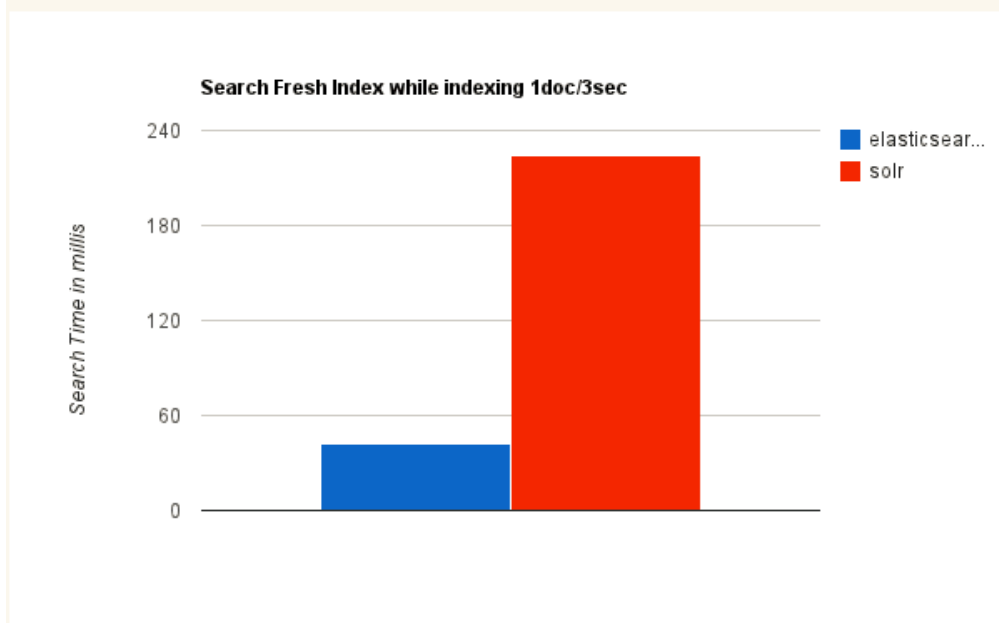
从搜索结果来源的角度，全文搜索引擎又可细分为两种，一种是拥有自己的检索程序（Indexer），俗称“蜘蛛”（Spider）程序或“机器人”（Robot）程序，并自建网页数据库，搜索结果直接从自身的数据库中调用，如上面提到的7家引擎；另一种则是租用其他引擎的数据库，并按自定的格式排列搜索结果，如Lycos引擎。

Elasticsearch和Solr比较

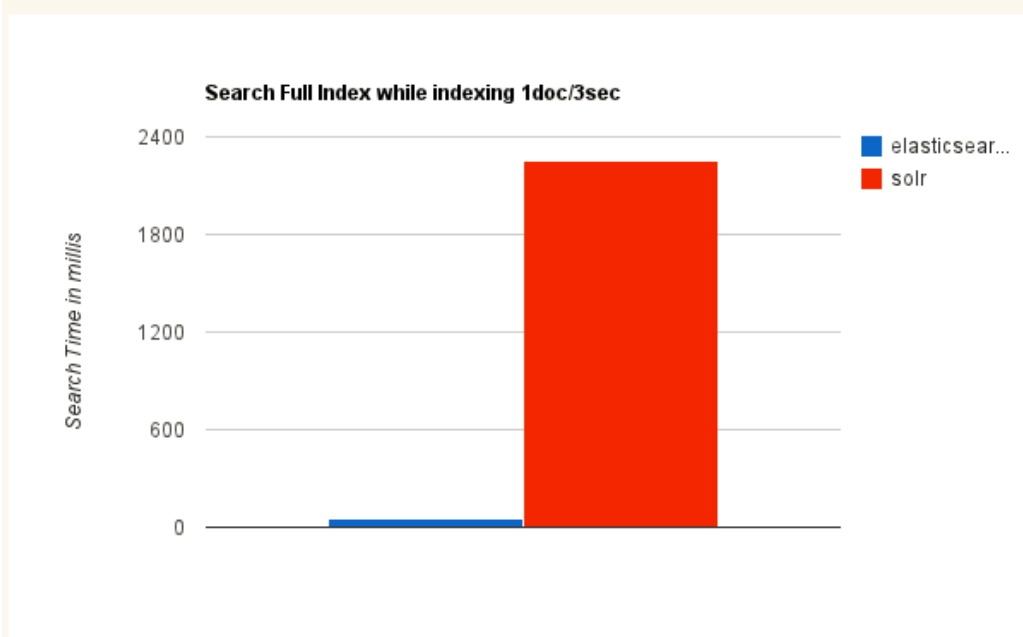
当单纯的对已有数据进行搜索时，Solr更快。



当实时建立索引时，Solr会产生io阻塞，查询性能较差，Elasticsearch具有明显的优势。



随着数据量的增加，Solr的搜索效率会变得更低，而Elasticsearch却没有明显的变化。



转变我们的搜索基础设施后从Solr Elasticsearch,我们看见一个即时~ 50 x提高搜索性能!



ElasticSearch vs Solr 总结

- 1、es基本是开箱即用，非常简单。Solr安装略微复杂一丢丢！
- 2、Solr 利用 Zookeeper 进行分布式管理，而 Elasticsearch 自身带有分布式协调管理功能。
 - (3) Solr 支持更多格式的数据，比如JSON、XML、CSV，而 Elasticsearch 仅支持json文件格式。
- 4、Solr 官方提供的功能更多，而 Elasticsearch 本身更侧重于核心功能，高级功能多有第三方插件提供，例如图形化界面需要kibana友好支撑
- 5、Solr 查询快，但更新索引时慢（即插入删除慢），用于电商等查询多的应用；
 - ES建立索引快（即查询慢），即实时性查询快，用于facebook新浪等搜索。
 - Solr 是传统搜索应用的有力解决方案，但 Elasticsearch 更适用于新兴的实时搜索应用。
- 6、Solr比较成熟，有一个更大，更成熟的用户、开发和贡献者社区，而 Elasticsearch相对开发维护者较少，更新太快，学习使用成本较高。

ElasticSearch安装

说明

我们需要下载和安装ElasticSearch的服务端和客户端！

注意：

ElasticSearch是使用java开发的，且本版本的es需要的jdk版本要是1.8以上，所以安装ElasticSearch 之前保证JDK1.8+安装完毕，并正确的配置好JDK环境变量，否则启动ElasticSearch失败。

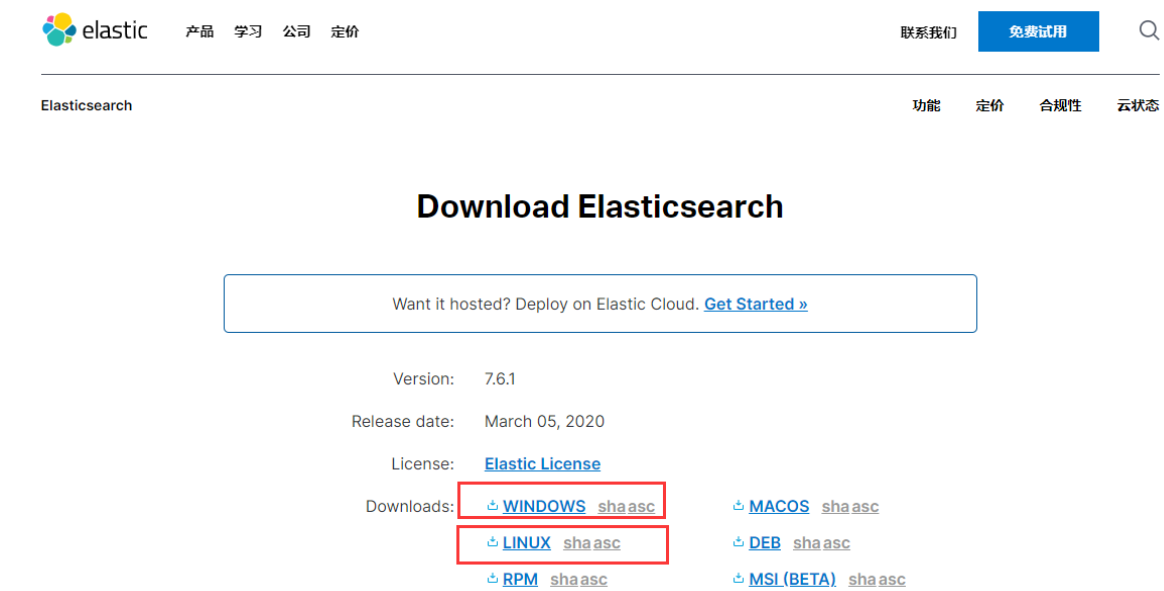
下载

ElasticSearch的官方地址：<https://www.elastic.co/products/elasticsearch>



官方下载地址: <https://www.elastic.co/cn/downloads/elasticsearch> (很慢, 可以翻墙下载!)

win下载: <https://artifacts.elastic.co/downloads/elasticsearch/elasticsearch-7.6.1.zip>



我这里已经帮大家下载好了, Linux 和 window 版的!

名称	修改日期	类型	大小
 elasticsearch-7.6.1-windows-x86_64.zip	2020/3/25 星期...	WinRAR ZIP 压缩...	286,253 KB
 elasticsearch-7.6.1-linux-x86_64.tar.gz	2020/3/25 星期...	WinRAR 压缩文件	289,507 KB

我们学习的话使用 window 或者 linux 都是可以的, 对于我们 java 开发来说没有区别, 只是连接的问题!

Windows更加方便一点! 所以我们前期都是用Window安装使用! 后面我们再真正的安装到Linux服务器上跑项目!

window 下安装使用

1、解压window的压缩包!

软件 (D:) > Environment > elasticsearch-7.6.1

名称	修改日期	类型	大小
bin	2020/3/25 星期...	文件夹	
config	2020/3/25 星期...	文件夹	
jdk	2020/3/25 星期...	文件夹	
lib	2020/3/25 星期...	文件夹	
logs	2020/2/29 星期...	文件夹	
modules	2020/3/25 星期...	文件夹	
plugins	2020/2/29 星期...	文件夹	
LICENSE.txt	2020/2/29 星期...	TXT 文件	14 KB
NOTICE.txt	2020/2/29 星期...	TXT 文件	511 KB
README.asciidoc	2020/2/29 星期...	ASCIIDOC 文件	8 KB

- 1 bin: 启动文件
- 2 config: 配置文件
 - 3 log4j2.properties: 日志配置文件
 - 4 jvm.options: java虚拟机的配置
 - 5 elasticsearch.yml: es的配置文件
- 6 data: 索引数据目录
- 7 lib: 相关类库Jar包
- 8 logs: 日志目录
- 9 modules: 功能模块
- 10 plugins: 插件

2、双击ElasticSearch下的bin目录中的elasticsearch.bat启动，控制台显示的日志（等待启动完毕!）：

```

[2020-03-25T17:12:45,966][INFO ][o.e.g.GatewayService] [KUANGSHEN] recovered [0] indices into cluster_state
[2020-03-25T17:12:46,493][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.watches] for index patterns [.watches*]
[2020-03-25T17:12:46,793][INFO ][o.e.h.AbstractHttpServerTransport] [KUANGSHEN] publish_address {127.0.0.1:9200}, bound_addresses {127.0.0.1:9200}, {:::1:9200}
[2020-03-25T17:12:46,796][INFO ][o.e.n.Node] [KUANGSHEN] started
[2020-03-25T17:12:49,123][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.watch-history-10] for index patterns [.watcher-history-10*]
[2020-03-25T17:12:49,217][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.triggered_watches] for index patterns [.triggered_watches*]
[2020-03-25T17:12:49,310][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [ilm-history] for index patterns [ilm-history-1*]
[2020-03-25T17:12:49,392][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.slm-history] for index patterns [.slm-history-1*]
[2020-03-25T17:12:49,481][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.monitoring-logstash-7] for index patterns [.monitoring-logstash-7*]
[2020-03-25T17:12:49,584][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.monitoring-es] for index patterns [.monitoring-es-7*]
[2020-03-25T17:12:49,679][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.monitoring-beats] for index patterns [.monitoring-beats-7*]
[2020-03-25T17:12:49,780][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.monitoring-alerts-7] for index patterns [.monitoring-alerts-7*]
[2020-03-25T17:12:49,876][INFO ][o.e.c.m.MetaDataIndexTemplateService] [KUANGSHEN] adding template [.monitoring-kibana-7] for index patterns [.monitoring-kibana-7*]
[2020-03-25T17:12:49,965][INFO ][o.e.x.i.a.TransportPutLifecycleAction] [KUANGSHEN] adding index lifecycle policy [watch-history-ilm-policy]
[2020-03-25T17:12:50,056][INFO ][o.e.x.i.a.TransportPutLifecycleAction] [KUANGSHEN] adding index lifecycle policy [ilm-history-ilm-policy]
[2020-03-25T17:12:50,169][INFO ][o.e.x.i.a.TransportPutLifecycleAction] [KUANGSHEN] adding index lifecycle policy [slm-history-ilm-policy]
[2020-03-25T17:12:50,375][INFO ][o.e.l.LicenseService] [KUANGSHEN] license [c4f6e661-7773-4196-8416-591729054545] mode [basic] - valid
[2020-03-25T17:12:50,377][INFO ][o.e.x.s.s.SecurityStatusChangeListener] [KUANGSHEN] Active license is now [BASIC]; Security is disabled
  
```

3、然后在浏览器访问：<http://localhost:9200> 得到如下信息，说明安装成功了：



安装ES的图形化界面插件客户端

注意：需要NodeJS的环境，我们讲解大前端进阶已经安装过了，没安装的需要安装！

Head是elasticsearch的集群管理工具，可以用于数据的浏览查询！被托管在github上面！

地址：<https://github.com/mobz/elasticsearch-head/>

1、下载 elasticsearch-head-master.zip

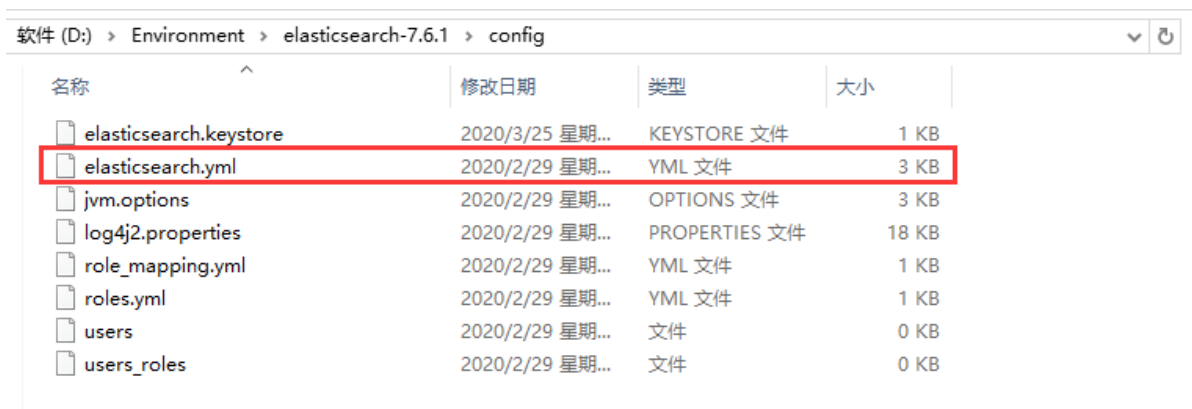
2、解压之后安装依赖！

```
1 cnpm install
2 npm run start
```

这将启动在端口9100上运行的本地web服务器，为elasticsearch-head服务！访问测试：



3、由于ES进程和客户端进程端口号不同，存在跨域问题，所以我们要在ES的配置文件中配置下跨域问题：



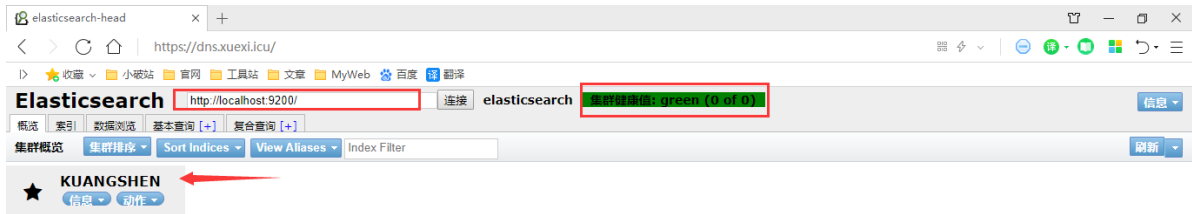
```
1 # 跨域配置:
2 http.cors.enabled: true
3 http.cors.allow-origin: "*"
```

```

84 # ----- Various -----
85 #
86 # Require explicit names when deleting indices:
87 #
88 #action.destructive_requires_name: true
89 # 跨域配置:
90 http.cors.enabled: true
91 http.cors.allow-origin: "*"

```

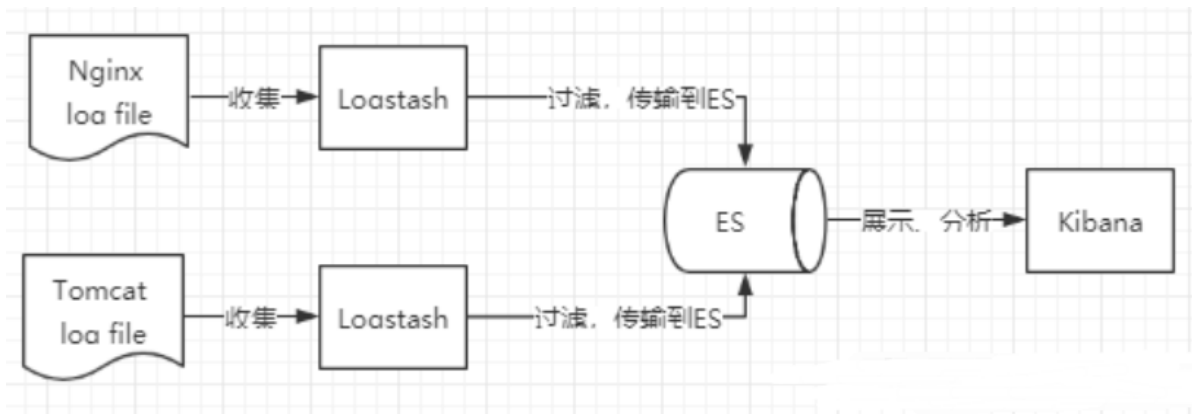
4、启动ElasticSearch，使用head工具进行连接测试！



了解 ELK

ELK是Elasticsearch、Logstash、Kibana三大开源框架首字母大写简称。市面上也被成为Elastic Stack。其中Elasticsearch是一个基于Lucene、分布式、通过Restful方式进行交互的近实时搜索平台框架。像类似百度、谷歌这种大数据全文搜索引擎的场景都可以使用Elasticsearch作为底层支持框架，可见Elasticsearch提供的搜索能力确实强大,市面上很多时候我们简称Elasticsearch为es。Logstash是ELK的中央数据流引擎，用于从不同目标（文件/数据存储/MQ）收集的不同格式数据，经过过滤后支持输出到不同目的地（文件/MQ/redis/elasticsearch/kafka等）。Kibana可以将elasticsearch的数据通过友好的页面展示出来，提供实时分析的功能。

市面上很多开发只要提到ELK能够一致说出它是一个日志分析架构技术栈总称，但实际上ELK不仅仅适用于日志分析，它还可以支持其它任何数据分析和收集的场景，日志分析和收集只是更具有代表性。并非唯一性。



安装Kibana

Kibana是一个针对Elasticsearch的开源分析及可视化平台，用来搜索、查看交互存储在Elasticsearch索引中的数据。使用Kibana，可以通过各种图表进行高级数据分析及展示。Kibana让海量数据更容易理解。它操作简单，基于浏览器的用户界面可以快速创建仪表盘（dashboard）实时显示Elasticsearch查询动态。设置Kibana非常简单。无需编码或者额外的基础架构，几分钟内就可以完成Kibana安装并启动Elasticsearch索引监测。

官网: <https://www.elastic.co/cn/kibana>

1、下载Kibana <https://www.elastic.co/cn/downloads/kibana>（注意版本对应关系）

Download Kibana

Want it hosted? Deploy on Elastic Cloud. [Get Started »](#)

Version: 7.6.1

Release date: March 05, 2020

License: [Elastic License](#)

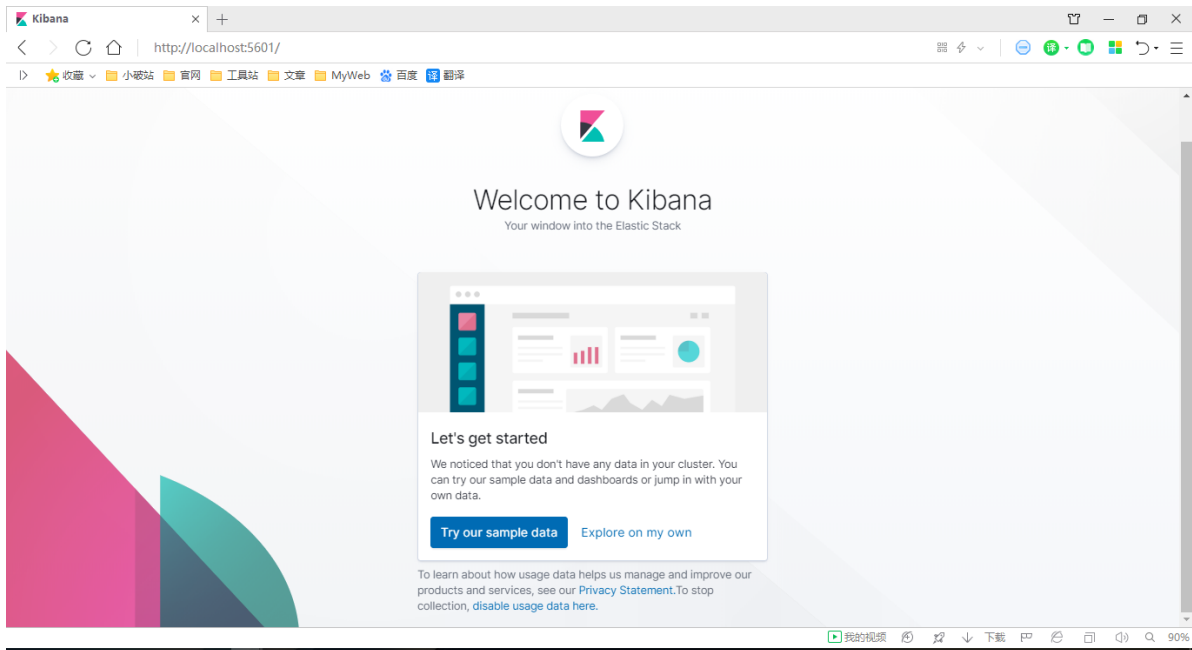
Downloads: [WINDOWS sha asc](#) [MAC sha asc](#)
[LINUX 64-BIT sha asc](#) [RPM 64-BIT sha asc](#)
[DEB 64-BIT sha asc](#)

2、将压缩包解压即可（需要一些时间）！

3、然后进入到bin目录下，启动服务就可以了（需要等待启动完成），ELK基本上都是拆箱即用的

```
C:\Windows\system32\cmd.exe
log [05:24:07.928] [info] [status] [plugin:transform@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:07.935] [info] [status] [plugin:encryptedSavedObjects@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:07.955] [info] [status] [plugin:actions@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:07.979] [info] [status] [plugin:alerting@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.077] [info] [status] [plugin:siem@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.088] [info] [status] [plugin:remote_clusters@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.094] [info] [status] [plugin:cross_cluster_replication@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.121] [info] [status] [plugin:upgrade_assistant@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.155] [info] [status] [plugin:uptime@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.165] [info] [status] [plugin:oss_telemetry@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.175] [info] [status] [plugin:file_upload@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.181] [info] [status] [plugin:data@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.189] [info] [status] [plugin:lens@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.219] [info] [status] [plugin:snapshot_restore@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.233] [info] [status] [plugin:input_control_vis@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.240] [info] [status] [plugin:kibana_react@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.245] [info] [status] [plugin:management@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.251] [info] [status] [plugin:navigation@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.259] [info] [status] [plugin:region_map@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.275] [info] [status] [plugin:telemetry@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.282] [info] [status] [plugin:ui_metric@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.510] [info] [status] [plugin:timelion@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.517] [info] [status] [plugin:markdown_vis@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.523] [info] [status] [plugin:metric_vis@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.530] [info] [status] [plugin:table_vis@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.538] [info] [status] [plugin:tagcloud@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:08.545] [info] [status] [plugin:vega@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:24.779] [warning] [reporting] Generating a random key for xpack.reporting.encryptionKey. To prevent pending reports from failing on restart, please set xpack.reporting.encryptionKey in kibana.yml
log [05:24:24.808] [info] [status] [plugin:reporting@7.6.1] Status changed from uninitialized to green - Ready
log [05:24:24.947] [info] [listening] Server running at http://localhost:5601
log [05:24:25.670] [info] [server] [Kibana] [http] http server running at http://localhost:5601
```

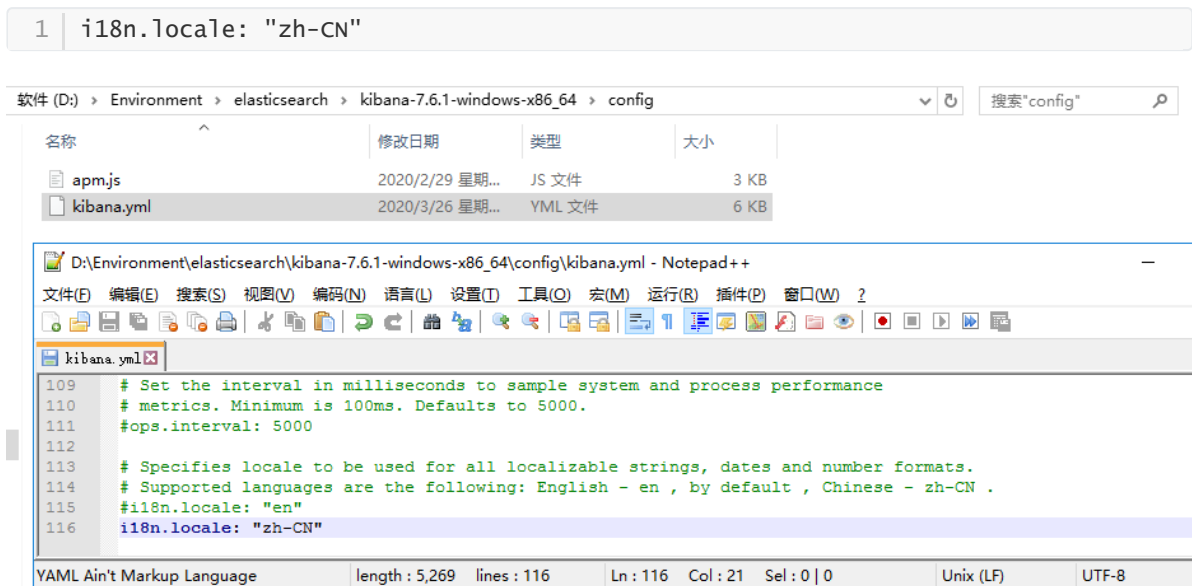
4、然后访问IP:5601，kibana会自动去访问9200，也就是elasticsearch的端口号（当然elasticsearch这个时候必须启动着），然后就可以使用kibana了！



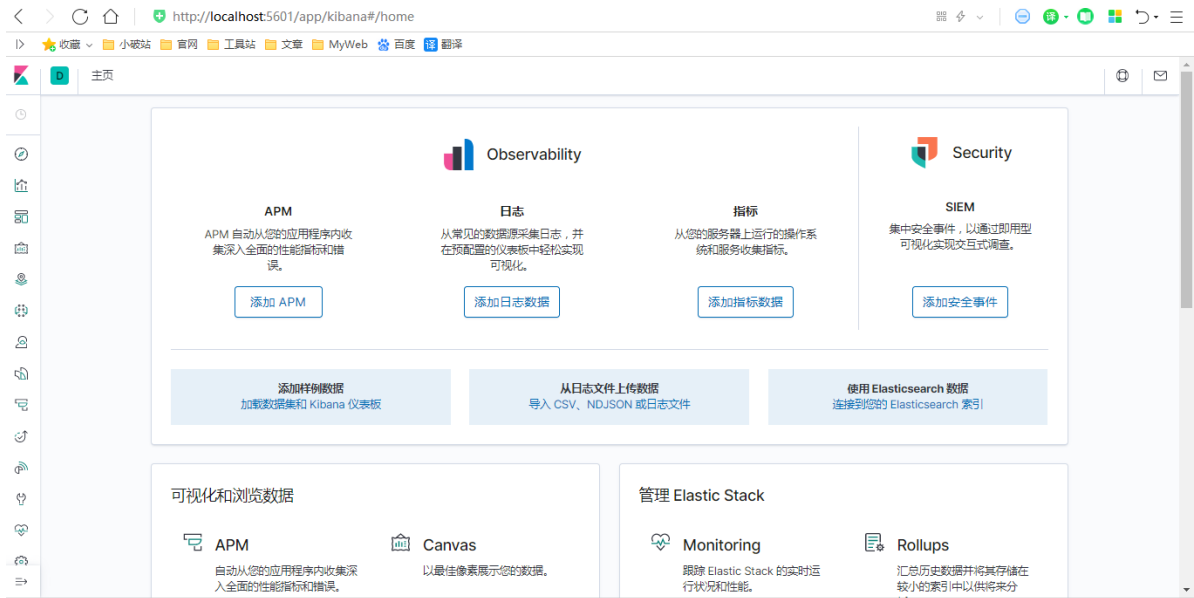
5、现在是英文的，看着有些吃力，我们配置为中文的！

中文包在 `kibana\x-pack\plugins\translations\translations\zh-CN.json`

只需要在配置文件 `kibana.yml` 中加入



6、重启查看效果！成功切换为中文的了！



ES核心概念

概述

在前面的学习中，我们已经掌握了es是什么，同时也把es的服务已经安装启动，那么es是如何去存储数据，数据结构是什么，又是如何实现搜索的呢？我们先来聊聊ElasticSearch的相关概念吧！

集群，节点，索引，类型，文档，分片，映射是什么？

elasticsearch是面向文档，关系行数据库 和 elasticsearch 客观的对比！

Relational DB	Elasticsearch
数据库(database)	索引 (indices)
表(tables)	types
行(rows)	documents
字段(columns)	fields

elasticsearch(集群)中可以包含多个索引(数据库)，每个索引中可以包含多个类型(表)，每个类型下又包含多个文档(行)，每个文档中又包含多个字段(列)。

物理设计：

elasticsearch 在后台把每个索引划分成多个分片，每分分片可以在集群中的不同服务器间迁移

逻辑设计：

一个索引类型中，包含多个文档，比如说文档1，文档2。当我们索引一篇文档时，可以通过这样的一各顺序找到 它: 索引 > 类型 > 文档ID，通过这个组合我们就能索引到某个具体的文档。注意:ID不必是整数，实际上它是个字符串。

文档

之前说elasticsearch是面向文档的，那么就意味着索引和搜索数据的最小单位是文档，elasticsearch中，文档有几个重要属性：

- 自我包含，一篇文档同时包含字段和对应的值，也就是同时包含 key:value！
- 可以是层次型的，一个文档中包含自文档，复杂的逻辑实体就是这么来的！
- 灵活的结构，文档不依赖预先定义的模式，我们知道关系型数据库中，要提前定义字段才能使用，在elasticsearch中，对于字段是非常灵活的，有时候，我们可以忽略该字段，或者动态的添加一个新的字段。

尽管我们可以随意的新增或者忽略某个字段，但是，每个字段的类型非常重要，比如一个年龄字段类型，可以是字符串也可以是整形。因为elasticsearch会保存字段和类型之间的映射及其他的设置。这种映射具体到每个映射的每种类型，这也是为什么在elasticsearch中，类型有时候也称为映射类型。

类型

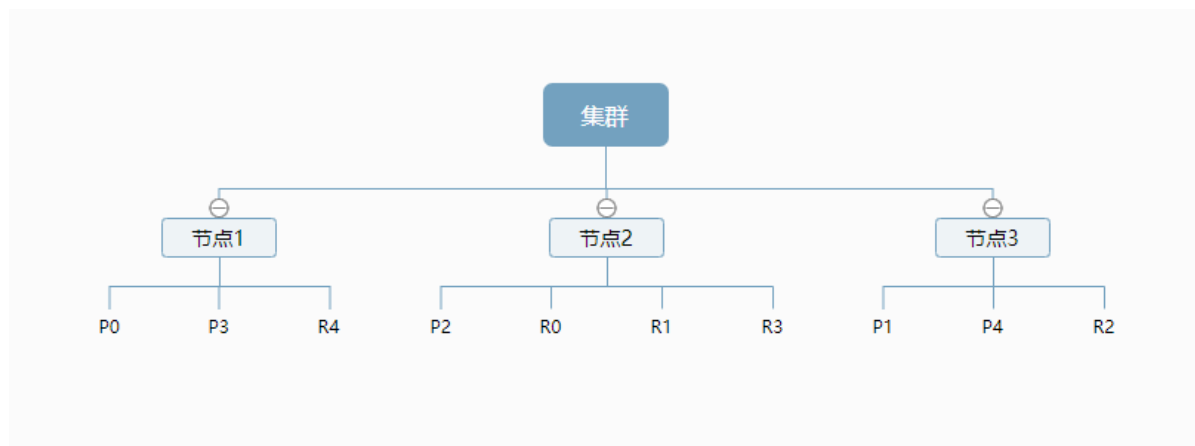
类型是文档的逻辑容器，就像关系型数据库一样，表格是行的容器。类型中对于字段的定义称为映射，比如 name 映射为字符串类型。我们说文档是无模式的，它们不需要拥有映射中所定义的所有字段，比如新增一个字段，那么elasticsearch是怎么做的呢？elasticsearch会自动的将新字段加入映射，但是这个字段的不确定它是什么类型，elasticsearch就开始猜，如果这个值是18，那么elasticsearch会认为它是整形。但是elasticsearch也可能猜不对，所以最安全的方式就是提前定义好所需要的映射，这点跟关系型数据库殊途同归了，先定义好字段，然后再使用，别整什么幺蛾子。

索引

索引是映射类型的容器，elasticsearch中的索引是一个非常大的文档集合。索引存储了映射类型的字段和其他设置。然后它们被存储到了各个分片上了。我们来研究下分片是如何工作的。

物理设计：节点和分片如何工作

一个集群至少有一个节点，而一个节点就是一个elasticsearch进程，节点可以有多个索引默认的，如果你创建索引，那么索引将会有个5个分片（primary shard ,又称主分片）构成的，每一个主分片会有一个副本（replica shard ,又称复制分片）



上图是一个有3个节点的集群，可以看到主分片和对应的复制分片都不会在同一个节点内，这样有利于某个节点挂掉了，数据也不至于丢失。实际上，一个分片是一个Lucene索引，一个包含倒排索引的文件目录，倒排索引的结构使得elasticsearch在不扫描全部文档的情况下，就能告诉你哪些文档包含特定的关键字。不过，等等，倒排索引是什么鬼？

倒排索引

elasticsearch使用的是一种称为倒排索引的结构，采用Lucene倒排索作为底层。这种结构适用于快速的全文搜索，一个索引由文档中所有不重复的列表构成，对于每一个词，都有一个包含它的文档列表。例如，现在有两个文档，每个文档包含如下内容：

1	Study every day, good good up to forever # 文档1包含的内容
2	To forever, study every day, good good up # 文档2包含的内容

为了创建倒排索引，我们首先要将每个文档拆分成独立的词(或称为词条或者tokens)，然后创建一个包含所有不重复的词条的排序列表，然后列出每个词条出现在哪个文档：

term	doc_1	doc_2
Study	√	×
To	×	×
every	√	√
forever	√	√
day	√	√
study	×	√
good	√	√
every	√	√
to	√	×
up	√	√

现在，我们试图搜索 to forever，只需要查看包含每个词条的文档

term	doc_1	doc_2
to	√	×
forever	√	√
total	2	1

两个文档都匹配，但是第一个文档比第二个匹配程度更高。如果没有别的条件，现在，这两个包含关键字的文档都将返回。

再来看一个示例，比如我们通过博客标签来搜索博客文章。那么倒排索引列表就是这样的一个结构：

博客文章(原始数据)		索引列表(倒排索引)	
博客文章ID	标签	标签	博客文章ID
1	python	python	1, 2, 3
2	python	linux	3, 4
3	linux, python		
4	linux		

如果要搜索含有 python 标签的文章，那相对于查找所有原始数据而言，查找倒排索引后的数据将会快的多。只需要 查看标签这一栏，然后获取相关的文章ID即可。

elasticsearch的索引和Lucene的索引对比

在elasticsearch中，索引这个词被频繁使用，这就是术语的使用。在elasticsearch中，索引被分为多个分片，每份分片是一个Lucene的索引。所以一个elasticsearch索引是由多个Lucene索引组成的。别问为什么，谁让elasticsearch使用Lucene作为底层呢！如无特指，说起索引都是指elasticsearch的索引。

接下来的一切操作都在kibana中Dev Tools下的Console里完成。基础操作！

ES基础操作

IK分词器插件

什么是IK分词器？

分词：即把一段中文或者别的划分成一个个的关键字，我们在搜索时候会把自己的信息进行分词，会把数据库中或者索引库中的数据进行分词，然后进行一个匹配操作，默认的中文分词是将每个字看成一个词，比如“我爱狂神”会被分为“我”，“爱”，“狂”，“神”，这显然是不符合要求的，所以我们需要安装中文分词器ik来解决这个问题。

IK提供了两个分词算法：ik_smart 和 ik_max_word，其中 ik_smart 为最少切分，ik_max_word为最细粒度划分！一会我们测试！

安装步骤

- 1、下载ik分词器的包，Github地址：<https://github.com/medcl/elasticsearch-analysis-ik/>（版本要对应）
- 2、下载后解压，并将目录拷贝到ElasticSearch根目录下的 plugins 目录中。



名称	修改日期	类型	大小
httpclient-4.5.2.jar	2018/5/6 星期日 ...	Executable Jar File	720 KB
httpcore-4.4.4.jar	2018/5/6 星期日 ...	Executable Jar File	320 KB
plugin-descriptor.properties	2020/3/7 星期六 ...	PROPERTIES 文件	2 KB
plugin-security.policy	2020/3/7 星期六 ...	POLICY 文件	1 KB
config	2019/12/25 星期...	文件夹	
commons-codec-1.9.jar	2018/5/6 星期日 ...	Executable Jar File	258 KB
commons-logging-1.2.jar	2018/5/6 星期日 ...	Executable Jar File	61 KB
elasticsearch-analysis-ik-7.6.1.jar	2020/3/7 星期六 ...	Executable Jar File	54 KB

- 3、重新启动 ElasticSearch 服务，在启动过程中，你可以看到正在加载“analysis-ik”插件的提示信息，服务启动后，在命令行运行 `elasticsearch-plugin list` 命令，确认 ik 插件安装成功。

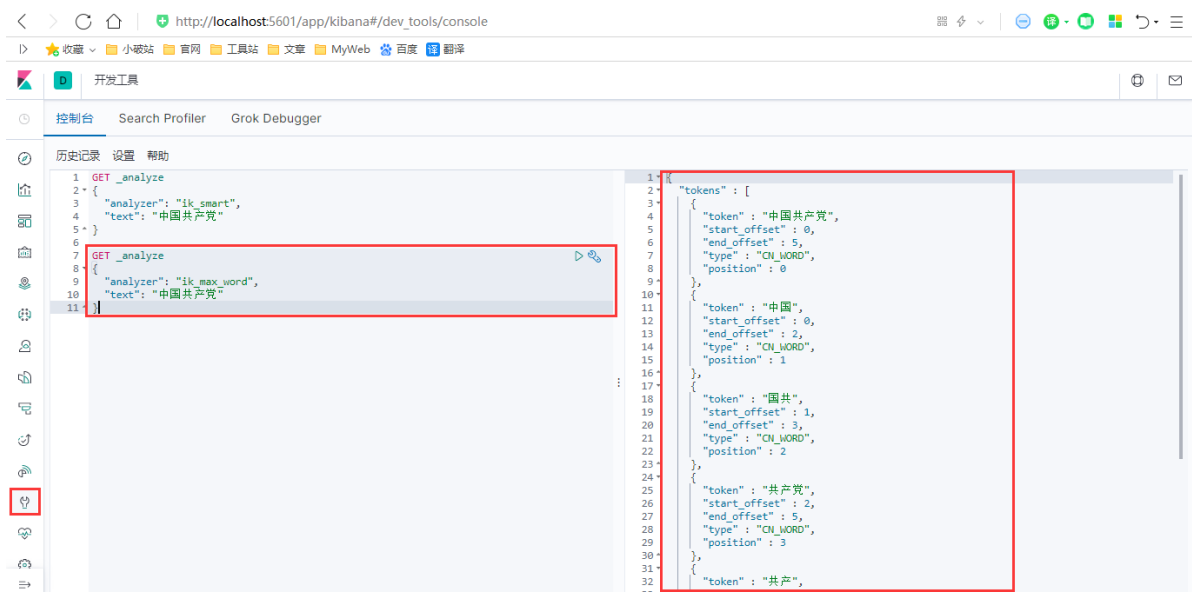
```
C:\Windows\system32\cmd.exe
[2020-03-26T14:19:26,706][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [spatial]
[2020-03-26T14:19:26,707][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [transform]
[2020-03-26T14:19:26,707][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [transport-netty4]
[2020-03-26T14:19:26,709][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [vectors]
[2020-03-26T14:19:26,712][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-analytics]
[2020-03-26T14:19:26,713][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-ccr]
[2020-03-26T14:19:26,714][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-core]
[2020-03-26T14:19:26,714][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-deprecation]
[2020-03-26T14:19:26,715][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-enrich]
[2020-03-26T14:19:26,716][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-graph]
[2020-03-26T14:19:26,716][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-ilm]
[2020-03-26T14:19:26,717][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-logstash]
[2020-03-26T14:19:26,719][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-ml]
[2020-03-26T14:19:26,719][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-monitoring]
[2020-03-26T14:19:26,721][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-rollover]
[2020-03-26T14:19:26,722][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-security]
[2020-03-26T14:19:26,725][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-sql]
[2020-03-26T14:19:26,725][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-voting-only-node]
[2020-03-26T14:19:26,726][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded module [x-pack-watcher]
[2020-03-26T14:19:26,727][INFO ][o.e.p.PluginsService ][KUANGSHEN] loaded plugin [analysis-ik]
[2020-03-26T14:19:31,684][INFO ][o.e.x.s.a.s.FileRolesStore][KUANGSHEN] parsed [0] roles from file [D:\Environment\elasticsearch\elasticsearch-7.6.1\config\rroles.yml]
[2020-03-26T14:19:32,397][INFO ][o.e.x.m.p.l.CppLogMessageHandler][KUANGSHEN] [controller/4604] [Main.cc@110] controller (64 bit): V
version 7.6.1 (Build 6eb6e036390036) Copyright (c) 2020 Elasticsearch BV
[2020-03-26T14:19:33,087][DEBUG][o.e.a.ActionModule ][KUANGSHEN] Using REST wrapper from plugin org.elasticsearch.xpack.security
ty.Security
[2020-03-26T14:19:33,222][INFO ][o.e.d.DiscoveryModule ][KUANGSHEN] using discovery type [zen] and seed hosts providers [settings
]
[2020-03-26T14:19:34,259][INFO ][o.e.n.Node ][KUANGSHEN] initialized
[2020-03-26T14:19:34,260][INFO ][o.e.n.Node ][KUANGSHEN] starting ...
[2020-03-26T14:19:35,351][INFO ][o.e.t.TransportService ][KUANGSHEN] publish_address {127.0.0.1:9300}, bound_addresses {127.0.0.1:
9300}, {:::1}:9300}
[2020-03-26T14:19:35,851][WARN ][o.e.b.BootstrapChecks ][KUANGSHEN] the default discovery settings are unsuitable for production
搜狗拼音输入法 全 :f [discovery.seed_hosts, discovery.seed_providers, cluster.initial_master_nodes] must be configured
```

```
Microsoft Windows [版本 10.0.14393]
(c) 2016 Microsoft Corporation. 保留所有权利。

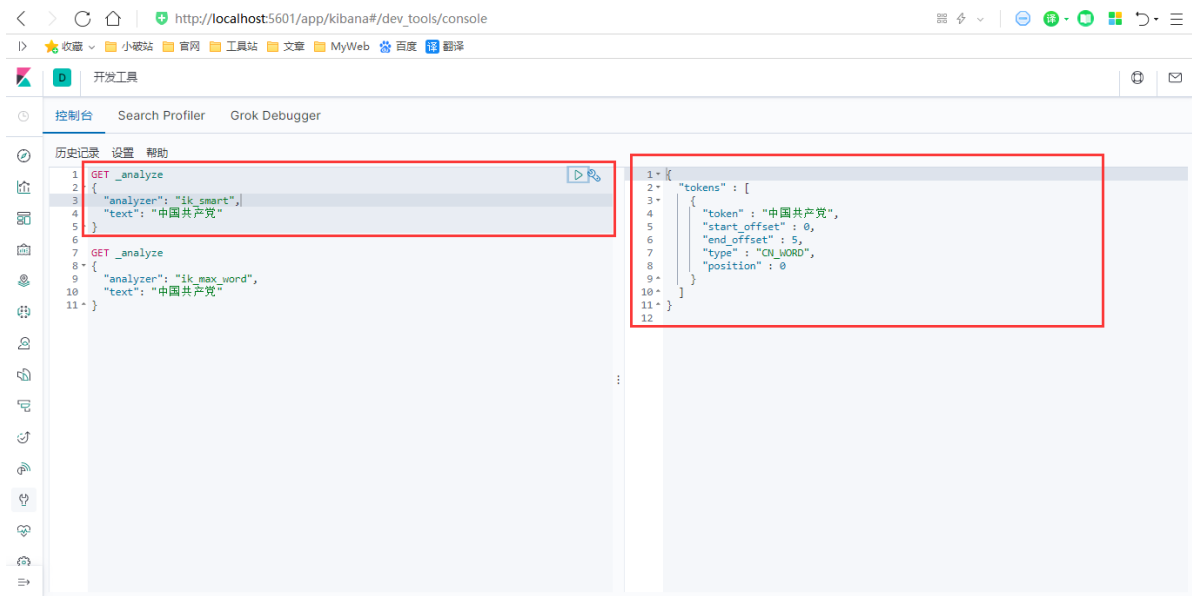
D:\Environment\elasticsearch\elasticsearch-7.6.1\bin>elasticsearch-plugin list
future versions of Elasticsearch will require Java 11; your Java version from [D:\Environment\java\jdk1.8.0_201\jre] does not meet th
is requirement
ik
D:\Environment\elasticsearch\elasticsearch-7.6.1\bin>_
```

4、在 kibana 中测试 ik 分词器，并就相关分词结果和 icu 分词器进行对比。

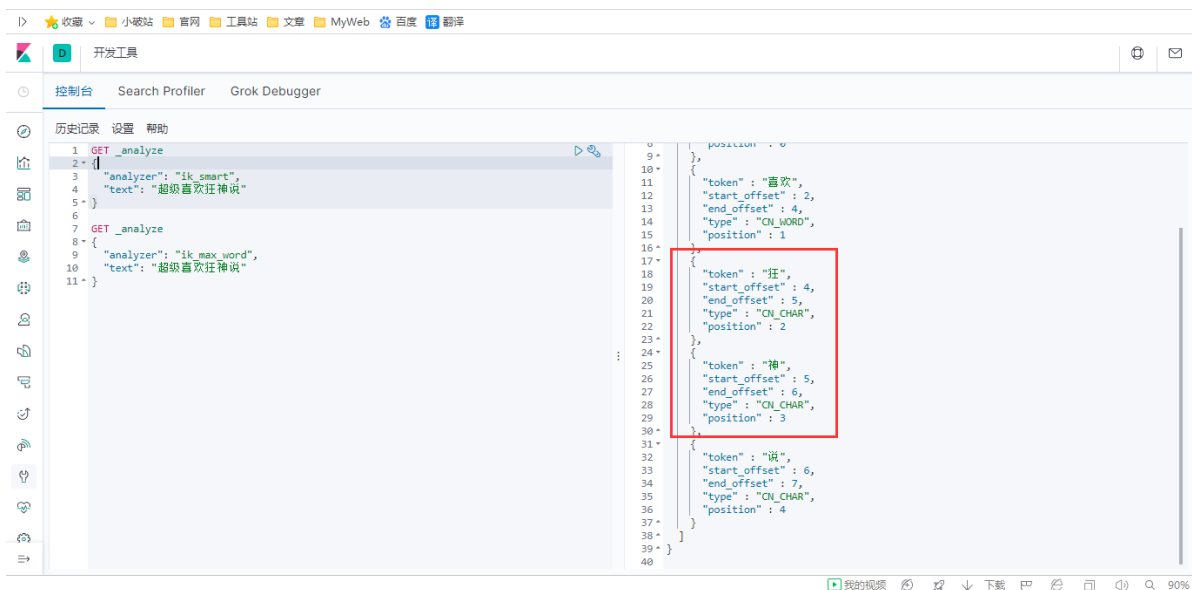
ik_max_word : 细粒度分词，会穷尽一个语句中所有分词可能，测试！



ik_smart : 粗粒度分词，优先匹配最长词，只有1个词！



5、我们输入超级喜欢狂神说！发现狂神说被切分了



如果我们想让系统识别“狂神说”是一个词，需要编辑自定义词库。

步骤：

- (1) 进入elasticsearch/plugins/ik/config目录
- (2) 新建一个my.dic文件，编辑内容：

```
1 狂神说
```

- (3) 修改IKAnalyzer.cfg.xml（在ik/config目录下）

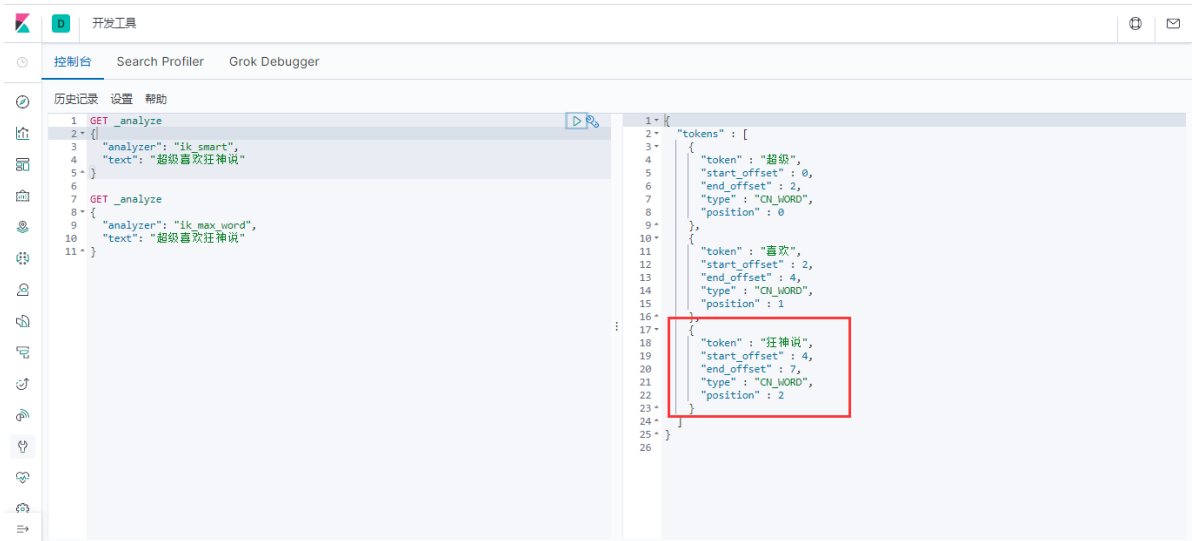
```
1 <properties>
2   <comment>IK Analyzer 扩展配置</comment>
3   <!-- 用户可以在这里配置自己的扩展字典 -->
4   <entry key="ext_dict">my.dic</entry>
5   <!-- 用户可以在这里配置自己的扩展停止词字典 -->
6   <entry key="ext_stopwords"></entry>
7 </properties>
```

修改完配置重新启动elasticsearch，再次测试！

发现监视了我们自己写的规则文件：

```
[2020-03-26T15:27:51.735][INFO ][o.e.g.GatewayService ] [KUANGSHEN] recovered [3] indices into cluster state
[2020-03-26T15:27:52.310][INFO ][o.e.h.AbstractHttpServerTransport] [KUANGSHEN] publish_address {127.0.0.1:9200}, bound_addresses [{127.0.0.1:9200}, {:::11:9200}]
[2020-03-26T15:27:52.314][INFO ][o.e.n.Node ] [KUANGSHEN] started
[2020-03-26T15:27:52.424][INFO ][o.w.a.d.Monitor ] [KUANGSHEN] try load config from D:\Environment\elasticsearch\elasticsearch-7.6.1\config\analysis-ik\IKAnalyzer.cfg.xml
[2020-03-26T15:27:52.433][INFO ][o.w.a.d.Monitor ] [KUANGSHEN] try load config from D:\Environment\elasticsearch\elasticsearch-7.6.1\plugins\ik\config\IKAnalyzer.cfg.xml
[2020-03-26T15:27:52.761][INFO ][o.w.a.d.Monitor ] [KUANGSHEN] [Dict Loading] D:\Environment\elasticsearch\elasticsearch-7.6.1\plugins\ik\config\my.dic
[2020-03-26T15:27:56.287][INFO ][o.e.c.r.a.AllocationService] [KUANGSHEN] Cluster health status changed from [RED] to [GREEN] (reason: [shards started [[.kibana_task_manager_1][0], [.kibana_1][0]])].
```

再次测试，发现狂神说变成了一个词：



到了这里，我们就明白了分词器的基本规则和使用！

Rest风格说明

一种软件架构风格，而不是标准，只是提供了一组设计原则和约束条件。它主要用于客户端和服务端交互的软件。基于这个风格设计的软件可以更简洁，更有层次，更易于实现缓存等机制。

基本Rest命令说明：

method	url地址	描述
PUT	localhost:9200/索引名称/类型名称/文档id	创建文档（指定文档id）
POST	localhost:9200/索引名称/类型名称	创建文档（随机文档id）
POST	localhost:9200/索引名称/类型名称/文档id/_update	修改文档
DELETE	localhost:9200/索引名称/类型名称/文档id	删除文档
GET	localhost:9200/索引名称/类型名称/文档id	查询文档通过文档id
POST	localhost:9200/索引名称/类型名称/_search	查询所有数据

基础测试

- 首先我们浏览器 <http://localhost:5601/> 进入 kibana里的Console
- 首先让我们在 Console 中输入：

```

1 // 命令解释
2 // PUT 创建命令 test1 索引 type1 类型 1 id
3 PUT /test1/type1/1
4 {
5     "name": "狂神说", // 属性
6     "age": 16 // 属性
7 }

```

返回结果（是以RESTful风格返回的）：

```

1 // 警告信息：不支持在文档索引请求中指定类型
2 // 而是使用无类型的端点(/{index}/_doc/{id}, /{index}/_doc, 或
  // /{index}/_create/{id})。
3 {
4     "_index" : "test1", // 索引
5     "_type" : "type1", // 类型
6     "_id" : "1", // id
7     "_version" : 1, // 版本
8     "result" : "created", // 操作类型
9     "_shards" : { // 分片信息
10         "total" : 2,
11         "successful" : 1,
12         "failed" : 0
13     },
14     "_seq_no" : 0,
15     "_primary_term" : 1
16 }

```

3、那么 name 这个字段用不用指定类型呢。毕竟我们关系型数据库 是需要指定类型的啊！

- 字符串类型
[text](#)、[keyword](#)
- 数值类型
[long](#), [integer](#), [short](#), [byte](#), [double](#), [float](#), [half float](#), [scaled float](#)
- 日期类型
[date](#)
- 布尔值类型
[boolean](#)
- 二进制类型
[binary](#)
- 等等.....

4、指定字段类型

```

1 PUT /test2
2 {
3     "mappings": {
4         "properties": {
5             "name": {
6                 "type": "text"
7             },
8             "age": {
9                 "type": "long"
10            },
11            "birthday": {
12                "type": "date"
13            }
14        }
15    }
16 }

```

```
14     }
15   }
16 }
```

输出:

```
1 {
2   "acknowledged" : true,
3   "shards_acknowledged" : true,
4   "index" : "test2"
5 }
```

5、查看一下索引字段

```
1 GET test2
```

输出:

```
1 {
2   "test2" : {
3     "aliases" : { },
4     "mappings" : {
5       "properties" : {
6         "age" : {
7           "type" : "long"
8         },
9         "birthday" : {
10          "type" : "date"
11        },
12        "name" : {
13          "type" : "text"
14        }
15      }
16    },
17    "settings" : {
18      "index" : {
19        "creation_date" : "1585384302712",
20        "number_of_shards" : "1",
21        "number_of_replicas" : "1",
22        "uuid" : "71TUZ84wRTW5P8lKeN4I4Q",
23        "version" : {
24          "created" : "7060199"
25        },
26        "provided_name" : "test2"
27      }
28    }
29  }
30 }
```

6、我们看上列中 字段类型是我自己定义的 那么 我们不定义类型 会是什么情况呢?

```
1 PUT /test3/_doc/1
2 {
3   "name": "狂神说",
4   "age": 13,
5   "birth": "1997-01-05"
6 }
7 # 输出
```



```

8  {
9    "_index" : "test3",
10   "_type" : "_doc",
11   "_id" : "1",
12   "_version" : 1,
13   "result" : "created",
14   "_shards" : {
15     "total" : 2,
16     "successful" : 1,
17     "failed" : 0
18   },
19   "_seq_no" : 0,
20   "_primary_term" : 1
21 }
22

```

查看一下test3索引:

```
1 GET test3
```

返回结果:

```

1  {
2    "test3" : {
3      "aliases" : { },
4      "mappings" : {
5        "properties" : {
6          "age" : {
7            "type" : "long"
8          },
9          "birth" : {
10           "type" : "date"
11         },
12         "name" : {
13           "type" : "text",
14           "fields" : {
15             "keyword" : {
16               "type" : "keyword",
17               "ignore_above" : 256
18             }
19           }
20         }
21       }
22     },
23     "settings" : {
24       "index" : {
25         "creation_date" : "1585384497051",
26         "number_of_shards" : "1",
27         "number_of_replicas" : "1",
28         "uuid" : "xESBKF1XTpCAZOgMqBNuBA",
29         "version" : {
30           "created" : "7060199"
31         },
32         "provided_name" : "test3"
33       }
34     }
35   }
36 }

```

我们看上列没有给字段指定类型那么es就会默认给我配置字段类型！

对比关系型数据库：

PUT test1/type1/1：索引test1相当于关系型数据库的库，类型type1就相当于表，1代表数据中的主键id

这里需要补充的是，在elasticsearch5版本前，一个索引下可以创建多个类型，但是在elasticsearch5后，一个索引只能对应一个类型，而id相当于关系型数据库的主键id若不指定就会默认生成一个20位的uuid，属性相当关系型数据库的column(列)。

而结果中的result则是操作类型，现在是created，表示第一次创建。如果再次点击执行该命令那么result则会是updated，我们细心则会发现_version开始是1，现在你每点击一次就会增加一次。表示第几次更改。

7、我们在来学一条命令 (elasticsearch 中的索引的情况)：

```
1 GET _cat/indices?v
```

返回结果：查看我们所有索引的状态健康情况 分片，数据储存大小等等。

	health	status	index	uuid	pri	rep	docs.count	docs.deleted	store.size	pri.store.size
2	yellow	open	test2	71TUZ84wRTW5P81KeN4I4Q	1	1	0	0	283b	283b
3	yellow	open	test3	tpIP1z3wQz6Uqp2xtPrVvA	1	1	1	0	4.3kb	4.3kb
4	green	open	.kibana_task_manager_1	1mgzW7HzTJ-JXIts2QdWA	1	0	2	0	49.2kb	49.2kb
5	green	open	.apm-agent-configuration	UOPaBD02SEue8Eo4pqtphg	1	0	0	0	283b	283b
6	green	open	.kibana_1	AVFXLZTqSe-8v1S8mnUx5w	1	0	14	4	53.3kb	53.3kb
7	yellow	open	test1	yD_14X6uQyqe-6MBTXsM0w	1	1	1	0	3.7kb	3.7kb

8、那么怎么删除一条索引呢(库)呢？

```
1 DELETE /test1
```

返回：

```
1 {
2   "acknowledged" : true    # 表示删除成功！
3 }
```

增删改查命令

创建数据PUT

第一条数据：

```
1 PUT /kuangshen/user/1
2 {
3   "name": "狂神说",
4   "age": 18,
5   "desc": "一顿操作猛如虎，一看工资2500",
6   "tags": ["直男", "技术宅", "温暖"]
7 }
```

第二条数据：

```

1 PUT /kuangshen/user/2
2 {
3   "name": "张三",
4   "age": 3,
5   "desc": "法外狂徒",
6   "tags": ["渣男", "旅游", "交友"]
7 }

```

第三条数据：

```

1 PUT /kuangshen/user/3
2 {
3   "name": "李四",
4   "age": 30,
5   "desc": "mmp, 不知道怎么形容",
6   "tags": ["靓女", "旅游", "唱歌"]
7 }

```

查看下数据：

The screenshot shows the Elasticsearch Kibana interface. The 'Data View' tab is active, displaying a table of search results for the 'kuangshen' index. The table has columns: _index, _type, _id, _score, name, age, and desc. The results are as follows:

_index	_type	_id	_score	name	age	desc
kuangshen	user	1	1	狂神说	18	一顿操作猛如虎，一看工资2500
kuangshen	user	2	1	张三	3	法外狂徒
kuangshen	user	3	1	李四	1	mmp, 不知道怎么形容

注意△：当执行命令时，如果数据不存在，则新增该条数据，如果数据存在则修改该条数据。

咱们通过 GET 命令查询一下：

```

1 GET kuangshen/user/1

```

返回结果：

```

1 {
2   "_index" : "kuangshen",
3   "_type" : "user",
4   "_id" : "1",
5   "_version" : 1,
6   "_seq_no" : 0,
7   "_primary_term" : 1,
8   "found" : true,
9   "_source" : {
10    "name" : "狂神说",
11    "age" : 18,
12    "desc" : "一顿操作猛如虎，一看工资2500",
13    "tags" : [
14      "直男",
15      "技术宅",
16      "温暖"
17    ]
18  }
19 }

```

如果你想更新数据 可以覆盖这条数据：

```

1 PUT /kuangshen/user/1
2 {
3   "name": "狂神说Java",
4   "age": 18,
5   "desc": "一顿操作猛如虎，一看工资2.5",
6   "tags": ["直男", "技术宅", "温暖"]
7 }

```

返回结果:

```

1 {
2   "_index" : "kuangshen",
3   "_type" : "user",
4   "_id" : "1",
5   "_version" : 2,
6   "result" : "updated",
7   "_shards" : {
8     "total" : 2,
9     "successful" : 1,
10    "failed" : 0
11  },
12   "_seq_no" : 3,
13   "_primary_term" : 1
14 }

```

已经修改了 那么 **PUT** 可以更新数据但是。麻烦的是 原数据你还要重写一遍要 这不符合我们规矩。

更新数据 POST

我们使用 **POST** 命令，在 **id** 后面跟 **_update**，要修改的内容放到 **doc** 文档(属性)中即可。

```

1 POST /kuangshen/user/1/_update
2 {
3   "doc": {
4     "name": "狂神说Java",
5     "desc": "关注狂神公众号每日更新文章哦"
6   }
7 }

```

返回结果:

```

1 {
2   "_index" : "kuangshen",
3   "_type" : "user",
4   "_id" : "1",
5   "_version" : 3,
6   "result" : "updated",
7   "_shards" : {
8     "total" : 2,
9     "successful" : 1,
10    "failed" : 0
11  },
12   "_seq_no" : 4,
13   "_primary_term" : 1
14 }

```

条件查询_search?q=

简单的查询，我们上面已经不知不觉的使用熟悉了：

```
1 GET kuangshen/user/1
```

我们来学习下条件查询_search?q=

```
1 GET kuangshen/user/_search?q=name:狂神说
```

通过_serarch?q=name:狂神说 查询条件是name属性有狂神说的那些数据。

别忘了_search 和 from 属性中间的分隔符？。

返回结果：

```
1 {
2   "took" : 16,
3   "timed_out" : false,
4   "_shards" : {
5     "total" : 1,
6     "successful" : 1,
7     "skipped" : 0,
8     "failed" : 0
9   },
10  "hits" : {
11    "total" : {
12      "value" : 1,          # 一共1条数据
13      "relation" : "eq"
14    },
15    "max_score" : 1.4229509,
16    "hits" : [
17      {
18        "_index" : "kuangshen",
19        "_type" : "user",
20        "_id" : "1",
21        "_score" : 1.4229509,
22        "_source" : {
23          "name" : "狂神说Java",
24          "age" : 18,
25          "desc" : "关注狂神公众号每日更新文章哦",
26          "tags" : [
27            "直男",
28            "技术宅",
29            "温暖"
30          ]
31        }
32      }
33    ]
34  }
35 }
```

我们看一下结果 返回并不是 数据本身，是给我们了一个 **hits**，还有_score得分，就是根据算法算出和查询条件匹配度高得分就搞。

```
1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match":{
5       "name": "狂神"
6     }
7   }
8 }
```

上例，查询条件是一步步构建出来的，将查询条件添加到 match 中即可。返回结果还是一样的：

```
1 {
2   "took" : 0,
3   "timed_out" : false,
4   "_shards" : {
5     "total" : 1,
6     "successful" : 1,
7     "skipped" : 0,
8     "failed" : 0
9   },
10  "hits" : {
11    "total" : {
12      "value" : 1,
13      "relation" : "eq"
14    },
15    "max_score" : 1.6285465,
16    "hits" : [
17      {
18        "_index" : "kuangshen",
19        "_type" : "user",
20        "_id" : "1",
21        "_score" : 1.6285465,
22        "_source" : {
23          "name" : "狂神说Java",
24          "age" : 18,
25          "desc" : "关注狂神公众号每日更新文章哦",
26          "tags" : [
27            "直男",
28            "技术宅",
29            "温暖"
30          ]
31        }
32      }
33    ]
34  }
35 }
```

除此之外，我们还可以查询全部：

```
1 GET kuangshen/user/_search #这是一个查询但是没有条件
```

```
1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match_all": {}
5   }
6 }
```

match_all的值为空，表示没有查询条件，就像select * from table_name一样。

返回结果：全部查询出来了！

如果有个需求，我们仅是需要查看 **name** 和 **desc** 两个属性，其他的不要怎么办？

```
1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match_all": {}
5   },
6   "_source": ["name","desc"]
7 }
```

如上例所示，在查询中，通过 **_source** 来控制仅返回 **name** 和 **age** 属性。

```
1 {
2   "took" : 1,
3   "timed_out" : false,
4   "_shards" : {
5     "total" : 1,
6     "successful" : 1,
7     "skipped" : 0,
8     "failed" : 0
9   },
10  "hits" : {
11    "total" : {
12      "value" : 3,
13      "relation" : "eq"
14    },
15    "max_score" : 1.0,
16    "hits" : [
17      {
18        "_index" : "kuangshen",
19        "_type" : "user",
20        "_id" : "2",
21        "_score" : 1.0,
22        "_source" : {
23          "name" : "张三",
24          "desc" : "法外狂徒"
25        }
26      },
27      {
28        "_index" : "kuangshen",
29        "_type" : "user",
30        "_id" : "3",
31        "_score" : 1.0,
32        "_source" : {
33          "name" : "李四",
34          "desc" : "mmp, 不知道怎么形容"
35        }
36      }
37    ]
38  }
39 }
```

```

36     },
37     {
38         "_index" : "kuangshen",
39         "_type" : "user",
40         "_id" : "1",
41         "_score" : 1.0,
42         "_source" : {
43             "name" : "狂神说Java",
44             "desc" : "关注狂神公众号每日更新文章哦"
45         }
46     }
47 ]
48 }
49 }

```

一般的，我们推荐使用构建查询，以后在与程序交互时的查询等也是使用构建查询方式处理查询条件，因为该方式可以构建更加复杂的查询条件，也更加一目了然

排序查询

我们说到排序 有人就会想到：正序 或 倒序 那么我们先来倒序：

```

1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match_all": {}
5   },
6   "sort": [
7     {
8       "age": {
9         "order": "desc"
10      }
11    }
12  ]
13 }

```

上例，在条件查询的基础上，我们又通过 sort 来做排序，排序对象是 age，order 是 desc 降序。

```

1 {
2   "took" : 0,
3   "timed_out" : false,
4   "_shards" : {
5     "total" : 1,
6     "successful" : 1,
7     "skipped" : 0,
8     "failed" : 0
9   },
10  "hits" : {
11    "total" : {
12      "value" : 3,
13      "relation" : "eq"
14    },
15    "max_score" : null,
16    "hits" : [
17      {
18        "_index" : "kuangshen",
19        "_type" : "user",

```



```
20     "_id" : "3",
21     "_score" : null,
22     "_source" : {
23         "name" : "李四",
24         "age" : 30,
25         "desc" : "mmp, 不知道怎么形容",
26         "tags" : [
27             "靓女",
28             "旅游",
29             "唱歌"
30         ]
31     },
32     "sort" : [
33         30
34     ]
35 },
36 {
37     "_index" : "kuangshen",
38     "_type" : "user",
39     "_id" : "1",
40     "_score" : null,
41     "_source" : {
42         "name" : "狂神说Java",
43         "age" : 18,
44         "desc" : "关注狂神公众号每日更新文章哦",
45         "tags" : [
46             "直男",
47             "技术宅",
48             "温暖"
49         ]
50     },
51     "sort" : [
52         18
53     ]
54 },
55 {
56     "_index" : "kuangshen",
57     "_type" : "user",
58     "_id" : "2",
59     "_score" : null,
60     "_source" : {
61         "name" : "张三",
62         "age" : 3,
63         "desc" : "法外狂徒",
64         "tags" : [
65             "渣男",
66             "旅游",
67             "交友"
68         ]
69     },
70     "sort" : [
71         3
72     ]
73 }
74 ]
75 }
76 }
```

正序，就是 desc 换成了 asc

```
1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match_all": {}
5   },
6   "sort": [
7     {
8       "age": {
9         "order": "asc"
10      }
11    }
12  ]
13 }
```

注意:在排序的过程中，只能使用可排序的属性进行排序。那么可以排序的属性有哪些呢？

- 数字
- 日期
- ID

其他都不行！

分页查询

```
1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match_all": {}
5   },
6   "sort": [
7     {
8       "age": {
9         "order": "asc"
10      }
11    }
12  ],
13   "from": 0, # 从第n条开始
14   "size": 1 # 返回n条数据
15 }
```

返回结果：

```
1 {
2   "took" : 0,
3   "timed_out" : false,
4   "_shards" : {
5     "total" : 1,
6     "successful" : 1,
7     "skipped" : 0,
8     "failed" : 0
9   },
10  "hits" : {
11    "total" : {
12      "value" : 3,
13      "relation" : "eq"
```

```

14     },
15     "max_score" : null,
16     "hits" : [
17       {
18         "_index" : "kuangshen",
19         "_type" : "user",
20         "_id" : "2",
21         "_score" : null,
22         "_source" : {
23           "name" : "张三",
24           "age" : 3,
25           "desc" : "法外狂徒",
26           "tags" : [
27             "渣男",
28             "旅游",
29             "交友"
30           ]
31         },
32         "sort" : [
33           3
34         ]
35       }
36     ]
37   }
38 }

```

就返回了一条数据 是从第0条开始的返回一条数据 。可以再测试！

学到这里，我们也可以看到，我们的查询条件越来越多，开始仅是简单查询，慢慢增加条件查询，增加排序，对返回 结果进行限制。所以，我们可以说:对elasticsearch于 来说，所有的查询条件都是可插拔的，彼此之间用 分割。比如说，我们在查询中，仅对返回结果进行限制:

```

1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match_all": {}
5   },
6
7   "from": 0, # 从第n条开始
8   "size": 1 # 返回n条数据
9 }

```

布尔查询

先增加一个数据:

```

1 PUT /kuangshen/user/4
2 {
3   "name": "狂神说",
4   "age": 3,
5   "desc": "一顿操作猛如虎，一看工资2500",
6   "tags": ["直男", "技术宅", "温暖"]
7 }

```

must (and)

我要查询所有 name 属性为“狂神”的数据，并且年龄为18岁的！

```
1 GET kuangshen/user/_search
2 {
3   "query": {
4     "bool": {
5       "must": [
6         {
7           "match": {
8             "name": "狂神说"
9           }
10        },
11        {
12          "match": {
13            "age": 3
14          }
15        }
16      ]
17    }
18  }
19 }
```

我们通过在 bool 属性内使用 must 来作为查询条件！看结果，是不是有点像 **and** 的感觉，里面的条件需要都满足！

should (or)

那么我要查询name为狂神 或 age 为18 的呢？

```
1 GET kuangshen/user/_search
2 {
3   "query": {
4     "bool": {
5       "should": [
6         {
7           "match": {
8             "name": "狂神说"
9           }
10        },
11        {
12          "match": {
13            "age": 18
14          }
15        }
16      ]
17    }
18  }
19 }
```

返回结果：

```
1 {
2   "took" : 0,
3   "timed_out" : false,
4   "_shards" : {
5     "total" : 1,
6     "successful" : 1,
7     "skipped" : 0,
```

```

8      "failed" : 0
9    },
10   "hits" : {
11     "total" : {
12       "value" : 2,
13       "relation" : "eq"
14     },
15     "max_score" : 3.1522982,
16     "hits" : [
17       {
18         "_index" : "kuangshen",
19         "_type" : "user",
20         "_id" : "1",
21         "_score" : 3.1522982,
22         "_source" : {
23           "name" : "狂神说Java",
24           "age" : 18,
25           "desc" : "关注狂神公众号每日更新文章哦",
26           "tags" : [
27             "直男",
28             "技术宅",
29             "温暖"
30           ]
31         }
32       },
33       {
34         "_index" : "kuangshen",
35         "_type" : "user",
36         "_id" : "4",
37         "_score" : 2.4708953,
38         "_source" : {
39           "name" : "狂神说",
40           "age" : 3,
41           "desc" : "一顿操作猛如虎，一看工资2500",
42           "tags" : [
43             "直男",
44             "技术宅",
45             "温暖"
46           ]
47         }
48       }
49     ]
50   }
51 }
52

```

我们的返回结果 是不是 出现了一个 age : 3的。是不是有点像 **or** 呢

must_not (not)

我想要查询 年龄不是 18 的 数据

```

1 GET kuangshen/user/_search
2 {
3   "query": {
4     "bool": {
5       "must_not": [
6         {

```

```

7         "match": {
8             "age": 18
9         }
10    }
11  ]
12 }
13 }
14 }

```

Fitter

我要查询 name 为狂神的, age大于10的数据

```

1 GET kuangshen/user/_search
2 {
3   "query":{
4     "bool": {
5       "must": [
6         {
7           "match": {
8             "name": "狂神"
9           }
10        }
11      ],
12      "filter": {
13        "range": {
14          "age": {
15            "gt": 10
16          }
17        }
18      }
19    }
20  }
21 }

```

这里就用到了 filter 条件过滤查询, 过滤条件的范围用 range 表示, gt 表示大于, 大于多少呢?是10。其余操作如下:

- gt 表示大于
- gte 表示大于等于
- lt 表示小于
- lte 表示小于等于

要查询 name 是 狂神, age 在 25~30 之间的怎么查?

```

1 GET kuangshen/user/_search
2 {
3   "query":{
4     "bool": {
5       "must": [
6         {
7           "match": {
8             "name": "狂神"
9           }
10        }
11      ],
12      "filter": {
13        "range": {

```

```

14         "age": {
15             "gte": 25,
16             "lte": 30
17         }
18     }
19 }
20 }
21 }
22 }

```

短语检索

我要查询 tags为男的数据

```

1 GET kuangshen/user/_search
2 {
3     "query":{
4         "match": {
5             "tags": "男"
6         }
7     }
8 }

```

返回了所有标签中带 男 的记录!

既然按照标签检索, 那么, 能不能写多个标签呢?又该怎么写呢?

```

1 GET kuangshen/user/_search
2 {
3     "query":{
4         "match": {
5             "tags": "男 技术"
6         }
7     }
8 }

```

返回: 只要含有这个标签满足一个就给我返回这个数据了。

term查询精确查询

term 查询是直接通过倒排索引指定的 词条, 也就是精确查找。

term和match的区别:

- match是经过分析(analyzer)的, 也就是说, 文档是先被分析器处理了, 根据不同的分析器, 分析出的结果也会不同, 在会根据分词 结果进行匹配。
- term是不经过分词的, 直接去倒排索引查找精确的值。

注意 ⚠: 我们现在 用的es7版本 所以我们用 mappings properties 去给多个字段(fields)指定类型的时候,不能给我们的 索引制定类型:

```

1 PUT testdb
2 {
3     "mappings": {

```

```

4     "properties": {
5         "name":{
6             "type":"text"
7         },
8         "desc":{
9             "type": "keyword"
10        }
11    }
12 }
13 }
14 // 插入数据
15 PUT testdb/_doc/1
16 {
17     "name": "狂神说Java name",
18     "desc": "狂神说Java desc"
19 }
20 PUT testdb/_doc/2
21 {
22     "name": "狂神说Java name",
23     "desc": "狂神说Java desc2"
24 }

```

上述中testdb索引中,字段name在被查询时会被分析器进行分析后匹配查询。而属于keyword类型不会被分析器处理。

我们来验证一下:

```

1 GET _analyze
2 {
3     "analyzer": "keyword",
4     "text": "狂神说Java name"
5 }

```

结果:

```

1 {
2     "tokens" : [
3         {
4             "token" : "狂神说Java name",
5             "start_offset" : 0,
6             "end_offset" : 12,
7             "type" : "word",
8             "position" : 0
9         }
10    ]
11 }

```

是不是没有被分析啊。就是简单的一个字符串啊。再测试

```

1 GET _analyze
2 {
3     "analyzer": "standard",
4     "text": "狂神说Java name"
5 }

```

结果:

```

1 {

```



```

2   "tokens" : [
3     {
4       "token" : "狂",
5       "start_offset" : 0,
6       "end_offset" : 1,
7       "type" : "<IDEOGRAPHIC>",
8       "position" : 0
9     },
10    {
11      "token" : "神",
12      "start_offset" : 1,
13      "end_offset" : 2,
14      "type" : "<IDEOGRAPHIC>",
15      "position" : 1
16    },
17    {
18      "token" : "说",
19      "start_offset" : 2,
20      "end_offset" : 3,
21      "type" : "<IDEOGRAPHIC>",
22      "position" : 2
23    },
24    {
25      "token" : "java",
26      "start_offset" : 3,
27      "end_offset" : 7,
28      "type" : "<ALPHANUM>",
29      "position" : 3
30    },
31    {
32      "token" : "name",
33      "start_offset" : 8,
34      "end_offset" : 12,
35      "type" : "<ALPHANUM>",
36      "position" : 4
37    }
38  ]
39 }

```

那么我们看一下 们字符串是不是被分析了啊。

总结：keyword 字段类型不会被分析器分析！

现在我们来查询一下：

```

1  GET testdb/_search    // text 会被分析器分析 查询
2  {
3    "query": {
4      "term": {
5        "name": "狂"
6      }
7    }
8  }
9
10 GET testdb/_search    // keyword 不会被分析所以直接查询
11 {
12   "query": {
13     "match": {

```

```
14     "desc": "狂神说Java desc"
15   }
16 }
17 }
```

查找多个精确值(terms)

官网地址: https://www.elastic.co/guide/cn/elasticsearch/guide/current/finding_multiple_exact_values.html

```
1  PUT testdb/_doc/3
2  {
3    "t1": "22",
4    "t2": "2020-4-16"
5  }
6  PUT testdb/_doc/4
7  {
8    "t1": "33",
9    "t2": "2020-4-17"
10 }
11
12 # 查询 精确查找多个值
13 GET testdb/_search
14 {
15   "query": {
16     "bool": {
17       "should": [
18         {
19           "term": {
20             "t1": "22"
21           }
22         },
23         {
24           "term": {
25             "t1": "33"
26           }
27         }
28       ]
29     }
30   }
31 }
```

除了bool查询之外:

```
1  GET testdb/_doc/_search
2  {
3    "query": {
4      "terms": {
5        "t1": ["22", "33"]
6      }
7    }
8  }
```

```

1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match": {
5       "name": "狂神"
6     }
7   },
8   "highlight" :{
9     "fields": {
10      "name":{}}
11   }
12 }
13 }

```

返回结果:

```

1  #! Deprecation: [types removal] specifying types in search requests is deprecated.
2  {
3    "took" : 62,
4    "timed_out" : false,
5    "_shards" : {
6      "total" : 1,
7      "successful" : 1,
8      "skipped" : 0,
9      "failed" : 0
10   },
11   "hits" : {
12     "total" : {
13       "value" : 2,
14       "relation" : "eq"
15     },
16     "max_score" : 1.6472635,
17     "hits" : [
18       {
19         "_index" : "kuangshen",
20         "_type" : "user",
21         "_id" : "4",
22         "_score" : 1.6472635,
23         "_source" : {
24           "name" : "狂神说",
25           "age" : 3,
26           "desc" : "一顿操作猛如虎，一看工资2500",
27           "tags" : [
28             "直男",
29             "技术宅",
30             "温暖"
31           ]
32         },
33         "highlight" : {
34           "name" : [
35             "<em>狂</em><em>神</em>说"
36           ]
37         }
38       },
39       {
40         "_index" : "kuangshen",

```

```

41     "_type" : "user",
42     "_id" : "1",
43     "_score" : 1.4348655,
44     "_source" : {
45         "name" : "狂神说Java",
46         "age" : 18,
47         "desc" : "关注狂神公众号每日更新文章哦",
48         "tags" : [
49             "直男",
50             "技术宅",
51             "温暖"
52         ]
53     },
54     "highlight" : {
55         "name" : [
56             "<em>狂</em><em>神</em>说Java"
57         ]
58     }
59 }
60 ]
61 }
62 }

```

我们可以看到 已 `<em>狂神</em>` 经帮我们加上了一个 `<em>` 标签

这是es帮我们加的标签。那我也可以自己自定义样式

```

1 GET kuangshen/user/_search
2 {
3   "query":{
4     "match": {
5       "name": "狂神"
6     }
7   },
8   "highlight" :{
9     "pre_tags": "<b class='key' style='color:red'>",
10    "post_tags": "</b>",
11    "fields": {
12      "name":{}
13    }
14  }
15 }

```

结果:

```

1  #! Deprecation: [types removal] specifying types in search requests is deprecated.
2  {
3    "took" : 1,
4    "timed_out" : false,
5    "_shards" : {
6      "total" : 1,
7      "successful" : 1,
8      "skipped" : 0,
9      "failed" : 0
10   },
11   "hits" : {
12     "total" : {

```

```

13     "value" : 2,
14     "relation" : "eq"
15 },
16 "max_score" : 1.6472635,
17 "hits" : [
18     {
19         "_index" : "kuangshen",
20         "_type" : "user",
21         "_id" : "4",
22         "_score" : 1.6472635,
23         "_source" : {
24             "name" : "狂神说",
25             "age" : 3,
26             "desc" : "一顿操作猛如虎，一看工资2500",
27             "tags" : [
28                 "直男",
29                 "技术宅",
30                 "温暖"
31             ]
32         },
33         "highlight" : {
34             "name" : [
35                 "<b class='key' style='color:red'>狂</b><b class='key'
36                 style='color:red'>神</b>说"
37             ]
38         },
39     {
40         "_index" : "kuangshen",
41         "_type" : "user",
42         "_id" : "1",
43         "_score" : 1.4348655,
44         "_source" : {
45             "name" : "狂神说Java",
46             "age" : 18,
47             "desc" : "关注狂神公众号每日更新文章哦",
48             "tags" : [
49                 "直男",
50                 "技术宅",
51                 "温暖"
52             ]
53         },
54         "highlight" : {
55             "name" : [
56                 "<b class='key' style='color:red'>狂</b><b class='key'
57                 style='color:red'>神</b>说Java"
58             ]
59         }
60     ]
61 }
62 }

```

需要注意的是:自定义标签中属性或样式中的逗号一律用英文状态的单引号表示，应该与外部 es 语法的 双引号区分开。

注意 elasticsearch 在第一个版本的开始 每个文档都储存在一个索引中，并分配一个 映射类型，映射类型用于表示被索引的文档或者实体的类型，这样带来了一些问题，导致后来在 elasticsearch6.0.0 版本中一个文档只能包含一个映射类型，而在 7.0.0 中，映射类型则将被弃用，到了 8.0.0 中则将完全被删除。

只要记得，一个索引下面只能创建一个类型就行了，其中各字段都具有唯一性，如果在创建映射的时候，如果没有指定文档类型，那么该索引的默认索引类型是 `_doc`，不指定文档id则会内部帮我们生成一个id字符串。

API创建索引及文档

找文档

网上的es教程大都十分老旧，而且es的版本众多，个别版本的差异还较大，另外es本身提供多种api，导致许多文章各种乱七八糟实例！所以后面直接放弃，从官网寻找方案，这里我使用elasticsearch最新的7.6.1版本来讲解：

- 1、进入es的官网指导文档 <https://www.elastic.co/guide/index.html>
- 2、找到 Elasticsearch Clients（这个就是客户端api文档）

Elasticsearch: Store, Search, and Analyze

- [Elasticsearch Reference \[7.6\] — other versions](#)
- [Elasticsearch Resiliency Status](#)
- [Painless Scripting Language \[7.6\] — other versions](#)
- [Plugins and Integrations \[7.6\] — other versions](#)
- [Elasticsearch Clients](#)
- [Elasticsearch for Apache Hadoop and Spark \[7.6\] — other versions](#)
- [Curator Index Management \[5.8\] — other versions](#)

- 3、我们使用java rest风格api，大家可以更加自己的版本选择特定的other versions。

Elasticsearch Clients

- [Java REST Client \[7.6\] — other versions](#)
- [Java API \[7.6\] — other versions](#)
- [JavaScript API \[7.x\] — other versions](#)
- [Ruby API \[7.x\] — other versions](#)
- [Go API](#)
- [.NET API \[7.x\] — other versions](#)
- [PHP API \[7.x\] — other versions](#)
- [Perl API](#)
- [Python API](#)
- [eland Client](#)
- [Rust API](#)
- [Community Contributed Clients](#)

最受欢迎的

-  [开始使用Elasticsearch](#)
-  [开启Kibana之旅](#)
-  [ELK之日志与指标分析：视频](#)

- 4、rest又分为high level和low level，我们直接选择high level下面的 Getting started

Java REST Client

最受欢迎的

-  开始使用Elasticsearch
-  开启Kibana之旅
-  ELK之日志与指标分析：视频

+ Java REST Client: 7.6 (current)

Overview

+ Java Low Level REST Client

- Java High Level REST Client

+ Getting started

+ Document APIs

+ Search APIs

+ Miscellaneous APIs

+ Index APIs

+ Cluster APIs

+ Ingest APIs

+ Snapshot APIs

+ Tasks APIs

+ Script APIs

+ Licensing APIs

+ Machine Learning APIs

5、向下阅读找到Maven依赖和基本配置！

< > X Home

https://www.elastic.co/guide/en/elasticsearch/client/java-rest/current/java-rest-high-getting-started-maven.html

收藏 小破站 官网 工具站 文章 MyWeb 百度 翻译

Maven configuration

Here is how you can configure the dependency using maven as a dependency manager. Add the following to your `pom.xml` file:

```
<dependency>
  <groupId>org.elasticsearch.client</groupId>
  <artifactId>elasticsearch-rest-high-level-client</artifactId>
  <version>7.6.1</version>
</dependency>
```

Gradle configuration

Here is how you can configure the dependency using gradle as a dependency manager. Add the following to your `build.gradle` file:

```
dependencies {
  compile 'org.elasticsearch.client:elasticsearch-rest-high-level-client:7.6.1'
}
```

+ Java REST Client: 7.6 (current)

Overview

+ Java Low Level REST Client

- Java High Level REST Client

- Getting started

Compatibility

Javadoc

Maven Repository

Dependencies

Initialization

RequestOptions

Asynchronous usage

+ Document APIs

+ Search APIs

+ Miscellaneous APIs

Java REST Client 说明

Java REST Client 有两种风格：

Java Low Level REST Client：用于Elasticsearch的官方低级客户端。它允许通过http与Elasticsearch集群通信。将请求编排和响应反编排留给用户自己处理。它兼容所有的Elasticsearch版本。（PS：学过WebService的话，对编排与反编排这个概念应该不陌生。可以理解为对请求参数的封装，以及对响应结果的解析）

Java High Level REST Client：用于Elasticsearch的官方高级客户端。它是基于低级客户端的，它提供很多API，并负责请求的编排与响应的反编排。（PS：就好比是，一个是传自己拼接好的字符串，并且自己解析返回的结果；而另一个是传对象，返回的结果也已经封装好了，直接是对象，更加规范了参数的名称以及格式，更加面向对象一点）

（PS：所谓低级与高级，我觉得一个很形象的比喻是，面向过程编程与面向对象编程）

网上很多教程比较老旧，都是使用TransportClient操作的，在Elasticsearch 7.0中不建议使用TransportClient，并且在8.0中会完全删除TransportClient。因此，官方更建议我们用Java High Level REST Client，它执行HTTP请求，而不是序列号的Java请求。既然如此，这里我们就直接用高级了。

配置基本项目依赖

1、新建一个springboot (2.2.5版) 项目 `kuang-elasticsearch`，导入web依赖即可！

2、配置es的依赖！

```
1 <properties>
2   <java.version>1.8</java.version>
3   <!-- 这里SpringBoot默认配置的版本不匹配，我们需要自己配置版本！ -->
4   <elasticsearch.version>7.6.1</elasticsearch.version>
5 </properties>
6
7 <dependency>
8   <groupId>org.springframework.boot</groupId>
9   <artifactId>spring-boot-starter-data-elasticsearch</artifactId>
10 </dependency>
```

3、继续阅读文档到Initialization，我们看到需要构建RestHighLevelClient对象；

```
1 RestHighLevelClient client = new RestHighLevelClient(
2     RestClient.builder(
3         new HttpHost("localhost", 9200, "http"),
4         new HttpHost("localhost", 9201, "http")); // 构建客户端对象
5
6 // 操作....
7 // 高级客户端内部会创建低级客户端用于基于提供的builder执行请求。低级客户端维护一个连接池，
8 // 并启动一些线程，因此当你用完以后应该关闭高级客户端，并且在内部它将会关闭低级客户端，以释放这些资源。关闭客户端可以使用close()方法：
9 client.close(); // 关闭
```

4、我们编写一个配置类，提供这个bean来进行操作

```
1 package com.kuang.config;
2
3 import org.apache.http.HttpHost;
4 import org.elasticsearch.client.RestClient;
5 import org.elasticsearch.client.RestHighLevelClient;
6 import org.springframework.context.annotation.Bean;
7 import org.springframework.context.annotation.Configuration;
8
9 @Configuration
10 public class ElasticsearchClientConfig {
11
12     @Bean
13     public RestHighLevelClient restHighLevelClient() {
14         RestHighLevelClient client = new RestHighLevelClient(
15             RestClient.builder(
16                 new HttpHost("127.0.0.1", 9200, "http"));
17         return client;
18     }
19 }
```


5、常用方法工具类封装

```

1  package com.kuang.utils;
2
3  import com.alibaba.fastjson.JSON;
4  import org.elasticsearch.action.admin.indices.delete.DeleteIndexRequest;
5  import org.elasticsearch.action.bulk.BulkRequest;
6  import org.elasticsearch.action.bulk.BulkResponse;
7  import org.elasticsearch.action.delete.DeleteRequest;
8  import org.elasticsearch.action.delete.DeleteResponse;
9  import org.elasticsearch.action.get.GetRequest;
10 import org.elasticsearch.action.get.GetResponse;
11 import org.elasticsearch.action.index.IndexRequest;
12 import org.elasticsearch.action.index.IndexResponse;
13 import org.elasticsearch.action.search.SearchRequest;
14 import org.elasticsearch.action.search.SearchResponse;
15 import org.elasticsearch.action.support.master.AcknowledgedResponse;
16 import org.elasticsearch.action.update.UpdateRequest;
17 import org.elasticsearch.action.update.UpdateResponse;
18 import org.elasticsearch.client.RequestOptions;
19 import org.elasticsearch.client.RestHighLevelClient;
20 import org.elasticsearch.client.indices.CreateIndexRequest;
21 import org.elasticsearch.client.indices.CreateIndexResponse;
22 import org.elasticsearch.client.indices.GetIndexRequest;
23 import org.elasticsearch.common.unit.TimeValue;
24 import org.elasticsearch.common.xcontent.XContentType;
25 import org.elasticsearch.index.query.QueryBuilders;
26 import org.elasticsearch.rest.RestStatus;
27 import org.elasticsearch.search.builder.SearchSourceBuilder;
28 import org.elasticsearch.search.fetch.subphase.FetchSourceContext;
29 import org.springframework.beans.factory.annotation.Autowired;
30 import org.springframework.beans.factory.annotation.Qualifier;
31 import org.springframework.stereotype.Component;
32
33 import java.io.IOException;
34 import java.util.List;
35 import java.util.concurrent.TimeUnit;
36
37 @Component
38 public class EsUtils<T> {
39
40     @Autowired
41     @Qualifier("restHighLevelClient")
42     private RestHighLevelClient client;
43     /**
44      * 判断索引是否存在
45      * @param index
46      * @return
47      * @throws IOException
48      */
49     public boolean existsIndex(String index) throws IOException {
50         GetIndexRequest request = new GetIndexRequest(index);
51         boolean exists = client.indices().exists(request,
52             RequestOptions.DEFAULT);
53         return exists;
54     }
55 }

```

```

54     /**
55      * 创建索引
56      * @param index
57      * @throws IOException
58      */
59     public boolean createIndex(String index) throws IOException {
60         CreateIndexRequest request = new CreateIndexRequest(index);
61         CreateIndexResponse createIndexResponse
=client.indices().create(request, RequestOptions.DEFAULT);
62         return createIndexResponse.isAcknowledged();
63     }
64
65     /**
66      * 删除索引
67      * @param index
68      * @return
69      * @throws IOException
70      */
71     public boolean deleteIndex(String index) throws IOException {
72         DeleteIndexRequest deleteIndexRequest = new
DeleteIndexRequest(index);
73         AcknowledgedResponse response =
client.indices().delete(deleteIndexRequest, RequestOptions.DEFAULT);
74         return response.isAcknowledged();
75     }
76
77
78     /**
79      * 判断某索引下文档id是否存在
80      * @param index
81      * @param id
82      * @return
83      * @throws IOException
84      */
85     public boolean docExists(String index, String id) throws IOException {
86         GetRequest getRequest = new GetRequest(index,id);
87         //只判断索引是否存在不需要获取_source
88         getRequest.fetchSourceContext(new FetchSourceContext(false));
89         getRequest.storedFields("_none_");
90         boolean exists = client.exists(getRequest, RequestOptions.DEFAULT);
91         return exists;
92     }
93
94     /**
95      * 添加文档记录
96      * @param index
97      * @param id
98      * @param t 要添加的数据实体类
99      * @return
100     * @throws IOException
101     */
102     public boolean addDoc(String index,String id,T t) throws IOException {
103
104         IndexRequest request = new IndexRequest(index);
105         request.id(id);
106         //timeout
107         request.timeout(TimeValue.timeValueSeconds(1));
108         request.timeout("1s");

```

```

109         request.source(JSON.toJSONString(t), XContentType.JSON);
110         IndexResponse indexResponse = client.index(request,
RequestOptions.DEFAULT);
111         RestStatus status = indexResponse.status();
112         return status==RestStatus.OK || status== RestStatus.CREATED;
113     }
114
115     /**
116      * 根据id来获取记录
117      * @param index
118      * @param id
119      * @return
120      * @throws IOException
121      */
122     public GetResponse getDoc(String index, String id) throws IOException {
123         GetRequest request = new GetRequest(index,id);
124         GetResponse getResponse = client.get(request,
RequestOptions.DEFAULT);
125         return getResponse;
126     }
127
128     /**
129      * 批量添加文档记录
130      * 没有设置id ES会自动生成一个，如果要设置 IndexRequest的对象.id()即可
131      * @param index
132      * @param list
133      * @return
134      * @throws IOException
135      */
136     public boolean bulkAdd(String index, List<T> list) throws IOException {
137         BulkRequest bulkRequest = new BulkRequest();
138         //timeout
139         bulkRequest.timeout(TimeValue.timeValueMinutes(2));
140         bulkRequest.timeout("2m");
141
142         for (int i =0;i<list.size();i++){
143             bulkRequest.add(new IndexRequest(index)
144                 .source(JSON.toJSONString(list.get(i))));
145         }
146         BulkResponse bulkResponse = client.bulk(bulkRequest,
RequestOptions.DEFAULT);
147
148         return !bulkResponse.hasFailures();
149     }
150 }
151
152     /**
153      * 批量删除和更新就不写了可根据上面几个方法来写
154      */
155
156     /**
157      * 更新文档记录
158      * @param index
159      * @param id
160      * @param t
161      * @return
162      * @throws IOException
163      */

```

```

164 public boolean updateDoc(String index,String id,T t) throws IOException
165 {
166     UpdateRequest request = new UpdateRequest(index,id);
167     request.doc(JSON.toJSONString(t));
168     request.timeout(TimeValue.timeValueSeconds(1));
169     request.timeout("1s");
170     UpdateResponse updateResponse = client.update(
171         request, RequestOptions.DEFAULT);
172     return updateResponse.status()==RestStatus.OK;
173 }
174
175 /**
176  * 删除文档记录
177  * @param index
178  * @param id
179  * @return
180  * @throws IOException
181  */
182 public boolean deleteDoc(String index,String id) throws IOException {
183     DeleteRequest request = new DeleteRequest(index,id);
184     //timeout
185     request.timeout(TimeValue.timeValueSeconds(1));
186     request.timeout("1s");
187     DeleteResponse deleteResponse = client.delete(
188         request, RequestOptions.DEFAULT);
189
190     return deleteResponse.status()== RestStatus.OK;
191 }
192
193 /**
194  * 根据某字段来搜索
195  * @param index
196  * @param field
197  * @param key 要收搜的关键字
198  * @throws IOException
199  */
200 public void search(String index,String field ,String key,Integer
from,Integer size) throws IOException {
201     SearchRequest searchRequest = new SearchRequest(index);
202     SearchSourceBuilder sourceBuilder = new SearchSourceBuilder();
203     sourceBuilder.query(QueryBuilders.termQuery(field, key));
204     //控制搜索
205     sourceBuilder.from(from);
206     sourceBuilder.size(size);
207     //最大搜索时间。
208     sourceBuilder.timeout(new TimeValue(60, TimeUnit.SECONDS));
209     searchRequest.source(sourceBuilder);
210     SearchResponse searchResponse = client.search(searchRequest,
RequestOptions.DEFAULT);
211     System.out.println(JSON.toJSONString(searchResponse.getHits()));
212 }
213 }
214
215 }

```

测试创建索引:

```

1  @Test
2  void testCreateIndex() throws IOException {
3      CreateIndexRequest request = new CreateIndexRequest("kuang_index");
4      CreateIndexResponse createIndexResponse
5
6      =restHighLevelClient.indices().create(request,RequestOptions.DEFAULT);
7      System.out.println(createIndexResponse);
8  }

```

测试获取索引:

```

1  @Test
2  void testExistsIndex() throws IOException {
3      GetIndexRequest request = new GetIndexRequest("kuang_index");
4      boolean exists = restHighLevelClient.indices().exists(request,
5      RequestOptions.DEFAULT);
6      System.out.println(exists);
7  }

```

测试删除索引:

```

1  @Test
2  void testDeleteIndexRequest() throws IOException {
3      DeleteIndexRequest deleteIndexRequest = new
4      DeleteIndexRequest("kuang_index");
5      AcknowledgedResponse response
6      = restHighLevelClient.indices().delete(deleteIndexRequest,
7      RequestOptions.DEFAULT);
8      System.out.println(response.isAcknowledged());
9  }

```

测试添加文档记录:

创建一个实体类User

```

1  @Data
2  @AllArgsConstructor
3  @NoArgsConstructor
4  @Component
5  public class User {
6      private String name;
7      private int age;
8  }

```

测试添加文档记录

```

1  @Test
2  void testAddDocument() throws IOException {
3      // 创建对象
4      User user = new User("狂神说", 3);
5      // 创建请求
6      IndexRequest request = new IndexRequest("kuang_index");
7      // 规则
8      request.id("1");

```

```

9      request.timeout(TimeValue.timeValueSeconds(1));
10     request.timeout("1s");
11     request.source(JSON.toJSONString(user), XContentType.JSON);
12     // 发送请求
13     IndexResponse indexResponse = restHighLevelClient.index(request,
RequestOptions.DEFAULT);
14     System.out.println(indexResponse.toString());
15     RestStatus Status = indexResponse.status();
16     System.out.println(Status == RestStatus.OK || Status ==
RestStatus.CREATED);
17 }

```

测试：判断某索引下文档id是否存在

```

1 // 判断此id是否存在这个索引库中
2 @Test
3 void testIsExists() throws IOException {
4     GetRequest getRequest = new GetRequest("kuang_index","1");
5     // 不获取_source上下文 storedFields
6     getRequest.fetchSourceContext(new FetchSourceContext(false));
7     getRequest.storedFields("_none_");
8     // 判断此id是否存在!
9     boolean exists = restHighLevelClient.exists(getRequest,
RequestOptions.DEFAULT);
10    System.out.println(exists);
11 }

```

测试：根据id获取记录

```

1 // 获得文档记录
2 @Test
3 void testGetDocument() throws IOException {
4     GetRequest getRequest = new GetRequest("kuang_index","3");
5     GetResponse getResponse = restHighLevelClient.get(getRequest,
RequestOptions.DEFAULT);
6     System.out.println(getResponse.getSourceAsString()); // 打印文档内容
7     System.out.println(getResponse);
8 }

```

测试：更新文档记录

```

1 // 更新文档记录
2 @Test
3 void testUpdateDocument() throws IOException {
4     UpdateRequest request = new UpdateRequest("kuang_index","1");
5     request.timeout(TimeValue.timeValueSeconds(1));
6     request.timeout("1s");
7
8     User user = new User("狂神说", 18);
9     request.doc(JSON.toJSONString(user), XContentType.JSON);
10
11    UpdateResponse updateResponse = restHighLevelClient.update(
12        request, RequestOptions.DEFAULT);
13
14    System.out.println(updateResponse.status() == RestStatus.OK);
15 }

```

测试：删除文档记录

```

1 // 删除文档测试
2 @Test
3 void testDelete() throws IOException {
4     DeleteRequest request = new DeleteRequest("kuang_index", "3");
5     //timeout
6     request.timeout(TimeValue.timeValueSeconds(1));
7     request.timeout("1s");
8
9     DeleteResponse deleteResponse = restHighLevelClient.delete(
10         request, RequestOptions.DEFAULT);
11     System.out.println(deleteResponse.status() == RestStatus.OK);
12 }

```

测试：批量添加文档

```

1 // 批量添加数据
2 @Test
3 void testBulkRequest() throws IOException {
4     BulkRequest bulkRequest = new BulkRequest();
5     //timeout
6     bulkRequest.timeout(TimeValue.timeValueMinutes(2));
7     bulkRequest.timeout("2m");
8
9     ArrayList<User> userList = new ArrayList<>();
10    userList.add(new User("kuangshen1", 3));
11    userList.add(new User("kuangshen2", 3));
12    userList.add(new User("kuangshen3", 3));
13    userList.add(new User("qinjiang1", 3));
14    userList.add(new User("qinjiang2", 3));
15    userList.add(new User("qinjiang3", 3));
16
17    for (int i = 0; i < userList.size(); i++) {
18        bulkRequest
19            .add(new IndexRequest("kuang_index")
20                .id("" + (i + 1)))
21
22        .source(JSON.toJSONString(userList.get(i)), XContentType.JSON));
23    }
24    // bulk
25    BulkResponse bulkResponse =
26    restHighLevelClient.bulk(bulkRequest, RequestOptions.DEFAULT);
27
28    System.out.println(!bulkResponse.hasFailures());
29 }

```

查询测试：

```

1 // 查询测试
2 /**
3     * 使用QueryBuilder
4     * termQuery("key", obj) 完全匹配
5     * termsQuery("key", obj1, obj2..) 一次匹配多个值
6     * matchQuery("key", obj) 单个匹配, field不支持通配符, 前缀具高级特性
7     * multiMatchQuery("text", "field1", "field2"..); 匹配多个字段, field有通
8     配符或行
9     * matchAllQuery(); 匹配所有文件
10    */

```

```

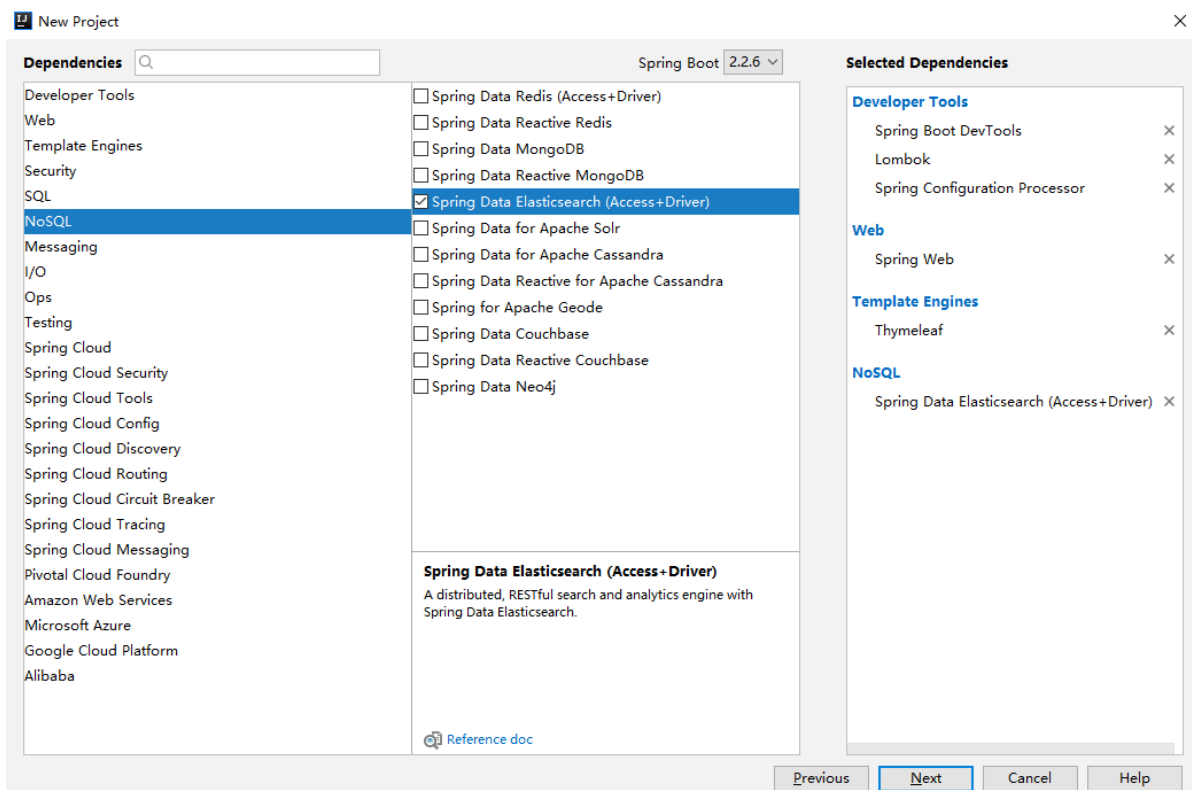
10 @Test
11 void testSearch() throws IOException {
12     SearchRequest searchRequest = new SearchRequest("kuang_index");
13     SearchSourceBuilder sourceBuilder = new SearchSourceBuilder();
14
15     // TermQueryBuilder termQueryBuilder = QueryBuilders.termQuery("name",
16     "qinjiang1");
17     MatchAllQueryBuilder matchAllQueryBuilder =
18     QueryBuilders.matchAllQuery();
19     sourceBuilder.query(matchAllQueryBuilder);
20
21     sourceBuilder.timeout(new TimeValue(60, TimeUnit.SECONDS));
22     searchRequest.source(sourceBuilder);
23
24     SearchResponse response = restHighLevelClient.search(searchRequest,
25     RequestOptions.DEFAULT);
26     System.out.println(JSON.toJSONString(response.getHits()));
27     System.out.println("=====SearchHit=====");
28     for (SearchHit documentFields : response.getHits().getHits()) {
29         System.out.println(documentFields.getSourceAsMap());
30     }
31 }

```

实战测试

初始化项目

- 1、启动es服务和客户端
- 2、使用springboot快速构建服务



- 3、修改版本依赖!


```

1 <properties>
2   <java.version>1.8</java.version>
3   <!-- 这里SpringBoot默认配置的版本不匹配，我们需要自己配置版本! -->
4   <elasticsearch.version>7.6.1</elasticsearch.version>
5 </properties>

```

4、配置 application.properties 文件

```

1 server.port=9090
2 # 关闭thymeleaf缓存
3 spring.thymeleaf.cache=false

```

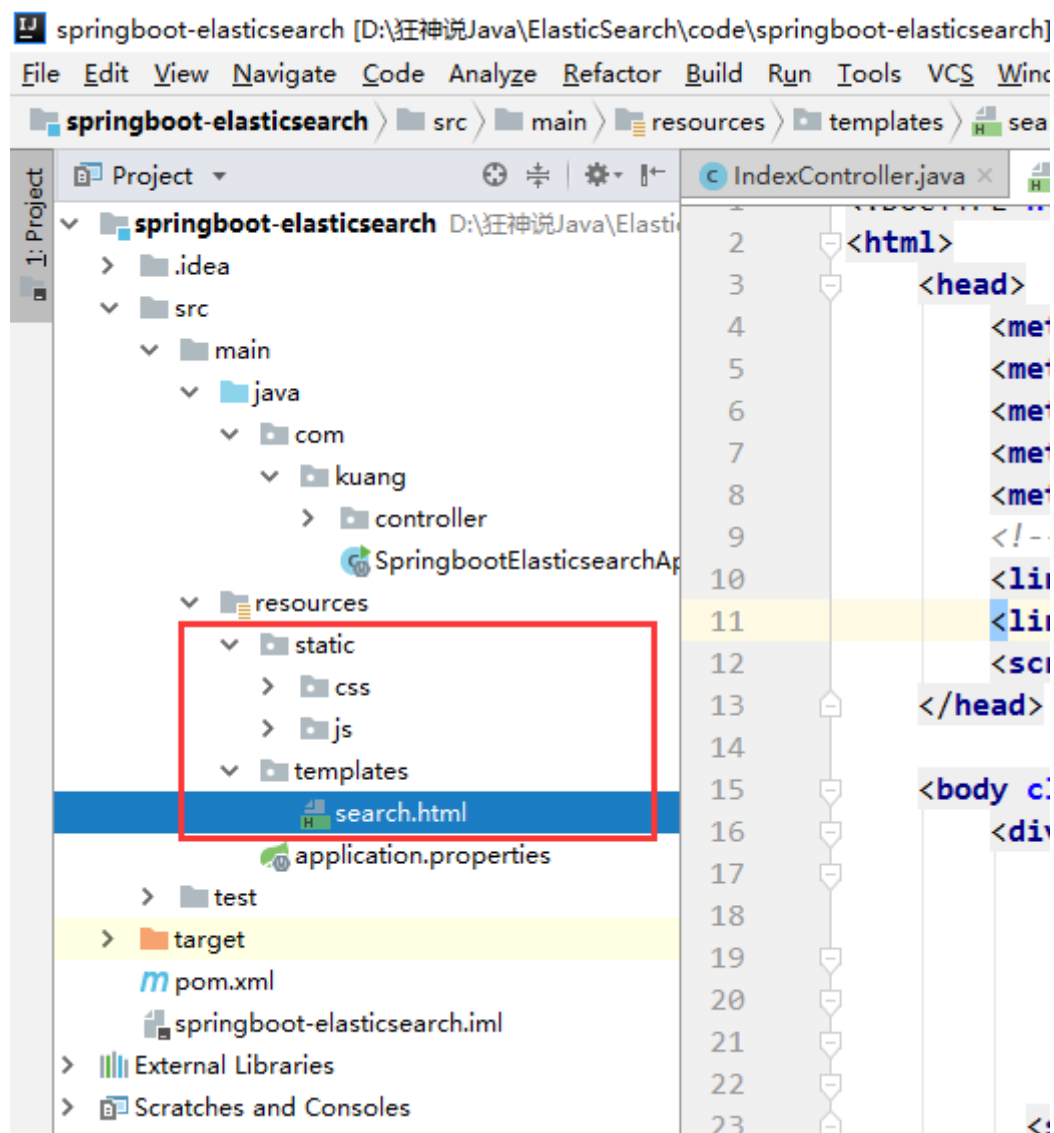
5、导入前端的素材！修改为Thymeleaf支持的格式！

```

1 <html xmlns:th="http://www.thymeleaf.org">

```

6、编写IndexController进行跳转测试！



jsoup讲解

1、导入jsoup的依赖

```

1 <!-- https://mvnrepository.com/artifact/org.jsoup/jsoup -->
2 <dependency>
3     <groupId>org.jsoup</groupId>
4     <artifactId>jsoup</artifactId>
5     <version>1.13.1</version>
6 </dependency>

```

2、编写一个工具类 HtmlParseUtil

```

1 public class HtmlParseUtil {
2     public static void main(String[] args) throws IOException {
3         // jsoup不能抓取ajax的请求，除非自己模拟浏览器进行请求！
4
5         // 1、https://search.jd.com/Search?keyword=java
6         String url = "https://search.jd.com/Search?keyword=java";
7         // 2、解析网页（需要联网）
8         Document document = Jsoup.parse(new URL(url), 30000);
9
10        // 3、抓取搜索到的数据！
11        // Document 就是我们JS的Document对象，你可以看到很多JS语法
12        Element element = document.getElementById("J_goodsList");
13
14        // 4、找到所有的li元素
15        Elements elements = element.getElementsByTag("li");
16
17        // 获取京东的商品信息
18        for (Element el : elements) {
19            // 这种网站，一般为了保证效率，一般会延时加载图片
20            // String img = el.getElementsByTag("img").eq(0).attr("src");
21            String img = el.getElementsByTag("img").eq(0).attr("source-data-lazy-img");
22            String price = el.getElementsByClass("p-price").eq(0).text();
23            String title = el.getElementsByClass("p-name").eq(0).text();
24
25            System.out.println(img);
26            System.out.println(price);
27            System.out.println(title);
28            System.out.println("=====");
29        }
30    }
31 }

```

3、封装一个实体类保存爬取下来的数据

```

1 @Data
2 @NoArgsConstructor
3 @AllArgsConstructor
4 public class Content {
5     private String title; // 商品名称
6     private String price; // 商品价格
7     private String img;    // 商品封面
8     // 大家可以自行扩展使用
9 }

```

4、封装为工具使用！

```

1 /**

```

```

2      * @author 狂神说Java 公众号：狂神说
3      * @param keywords 要搜索的关键词！
4      * @return 抓取的商品集合
5      */
6  public List<Content> parseJD(String keywords) throws Exception {
7      String url = "https://search.jd.com/Search?keyword="+keywords;
8      Document document = Jsoup.parse(new URL(url), 30000);
9      Element element = document.getElementById("J_goodsList");
10     Elements elements = element.getElementsByTag("li");
11
12     ArrayList<Content> goodsList = new ArrayList<>();
13
14     // 获取京东的商品信息
15     for (Element el : elements) {
16         String img = el.getElementsByTag("img").eq(0).attr("source-data-lazy-img");
17         String price = el.getElementsByClass("p-price").eq(0).text();
18         String title = el.getElementsByClass("p-name").eq(0).text();
19         // 封装获取的数据
20         Content content = new Content();
21         content.setImg(img);
22         content.setPrice(price);
23         content.setTitle(title);
24         goodsList.add(content);
25     }
26     return goodsList;
27 }

```

5、测试工具类的使用！

```

1  public static void main(String[] args) throws Exception {
2      new HtmlParseUtil().parseJD("vue").forEach(System.out::println);
3  }

```

搞定收工！简单爬虫编写完毕！我们这里的数据就使用爬取的即可，平时开发es的数据可能来自多个地方，你们可以从数据库查询获取也是一样的，后面我们来测试下效果！



业务编写

1、导入ElasticsearchClientConfig 配置类

```

1  @Configuration
2  public class ElasticsearchClientConfig {
3
4      @Bean
5      public RestHighLevelClient restHighLevelClient() {
6          RestHighLevelClient client = new RestHighLevelClient(
7              RestClient.builder(
8                  new HttpHost("127.0.0.1", 9200, "http"));
9          return client;
10     }
11
12 }

```

2、编写业务!

```

1  @Service
2  public class ContentService {
3
4      @Autowired
5      private RestHighLevelClient restHighLevelClient;
6
7      // 1、解析数据存入es
8      public Boolean parseContent(String keywords) throws Exception {
9          // 解析查询出来的数据
10         List<Content> contents = new HtmlParseUtil().parseJD(keywords);
11         // 封装数据到索引库中!
12         BulkRequest bulkRequest = new BulkRequest();
13         bulkRequest.timeout(TimeValue.timeValueMinutes(2));
14         bulkRequest.timeout("2m");
15         for(int i =0;i<contents.size();i++){
16             bulkRequest
17                 .add(new IndexRequest("jd_goods")
18
19 .source(JSON.toJSONString(contents.get(i)),XContentType.JSON));
20         }
21
22         BulkResponse bulkResponse =
23         restHighLevelClient.bulk(bulkRequest,RequestOptions.DEFAULT);
24         return !bulkResponse.hasFailures();
25     }
26
27     // 2、实现搜索功能，带分页处理
28     public List<Map<String, Object>> searchContentPage(String keyword, int
29     pageNo, int pageSize) throws IOException {
30
31         // 基本的参数判断!
32         if(pageNo <= 1){
33             pageNo = 1;
34         }
35         // 基本的条件搜索
36         SearchRequest searchRequest = new SearchRequest("jd_goods");
37         SearchSourceBuilder sourceBuilder = new SearchSourceBuilder();
38         // 分页
39         sourceBuilder.from(pageNo);
40         sourceBuilder.size(pageSize);
41         // 精准匹配 QueryBuilders 根据自己要求配置查询条件即可!

```

```

39     TermQueryBuilder termQueryBuilder = QueryBuilders.termQuery("title",
keyword);
40     sourceBuilder.query(termQueryBuilder);
41     sourceBuilder.timeout(new TimeValue(60, TimeUnit.SECONDS));
42
43     // 搜索
44     searchRequest.source(sourceBuilder);
45     SearchResponse response = restHighLevelClient.search(searchRequest,
RequestOptions.DEFAULT);
46     // 解析结果!
47     List<Map<String, Object>> list = new ArrayList<>();
48     for (SearchHit documentFields : response.getHits().getHits()) {
49         list.add(documentFields.getSourceAsMap());
50     }
51     return list;
52 }
53
54 }

```

3、controller

```

1  @RestController
2  public class ContentController {
3
4      @Autowired
5      private ContentService contentService;
6
7      @GetMapping("/parse/{keyword}")
8      public Boolean parse(@PathVariable("keyword") String keyword) throws
Exception {
9          return contentService.parseContent(keyword);
10     }
11
12     //http://localhost:9090/search/java/1/10
13     @GetMapping("/search/{keyword}/{pageNo}/{pageSize}")
14     public List<Map<String, Object>> search(@PathVariable("keyword") String
keyword,
15                                           @PathVariable("pageNo") int pageNo,
16                                           @PathVariable("pageSize") int pageSize) throws
Exception {
17
18         return contentService.searchContentPage(keyword, pageNo, pageSize);
19     }
20
21 }

```

前端逻辑

1、定义导入vue和axios的依赖!

```

1  <script th:src="@{/js/axios.js}"></script>
2  <script th:src="@{/js/vue.min.js}"></script>

```

2、初始化Vue对象，给外层div绑定app对象!

```

1 <script>
2
3   new Vue({
4     el: '#app',
5     data: {
6       keyword: '', // 搜索关键字
7       results: [] // 搜索的结果
8     }
9   })
10
11 </script>

```

3、绑定搜索框及相关事件!

```

<!-- 搜索 -->
<div id="mallSearch" class="mall-search">
  <form name="searchTop" class="mallSearch-form clearfix">
    <fieldset>
      <legend>天猫搜索</legend>
      <div class="mallSearch-input clearfix">
        <div class="s-combobox" id="s-combobox-689">
          <div class="s-combobox-input-wrap">
            <input v-model="keyword" type="text" autocomplete="off" value="dd" id="mq"
              class="s-combobox-input" aria-haspopup="true">
          </div>
        </div>
        <button type="submit" id="searchbtn" @click.prevent="searchKey">搜索</button>
      </div>
    </fieldset>
  </form>
  <ul class="relKeyTop">
    <li><a>狂神说Java</a></li>
    <li><a>狂神说前端</a></li>
    <li><a>狂神说Linux</a></li>
    <li><a>狂神说大数据</a></li>
    <li><a>狂神聊理财</a></li>
  </ul>
</div>

```

双向绑定

绑定事件

4、编写方法，获取后端传递的数据!

```

1 <script>
2
3   new Vue({
4     el: '#app',
5     data: {
6       keyword: '',
7       results: []
8     },
9     methods: {
10      searchKey(){
11        var keyword = this.keyword;
12        console.log(keyword);
13        axios.get('search/'+keyword+'/1/10').then(response=>{
14          console.log(response);
15          this.results = response.data;
16        });
17      }
18    }
19  })
20
21 </script>

```

5、渲染解析回来的数据!



搜索高亮

1、编写业务类，处理高亮字段

```
1 // 3、实现搜索功能，带高亮
2 public List<Map<String, Object>> searchContentHighlighter(String keyword,
3 int pageNo, int pageSize) throws IOException {
4     // 基本的参数判断!
5     if(pageNo <= 1){
6         pageNo = 1;
7     }
8     // 基本的条件搜索
9     SearchRequest searchRequest = new SearchRequest("jd_goods");
10    SearchSourceBuilder sourceBuilder = new SearchSourceBuilder();
11    // 分页
12    sourceBuilder.from(pageNo);
13    sourceBuilder.size(pageSize);
14    // 精准匹配 QueryBuilders 根据自己要求配置查询条件即可!
15    TermQueryBuilder termQueryBuilder = QueryBuilders.termQuery("title",
16 keyword);
17    sourceBuilder.query(termQueryBuilder);
18    // 高亮构建!
19    HighlightBuilder highlightBuilder = new HighlightBuilder(); //生成高亮查询
20    highlightBuilder.field("title"); //高亮查询字段
21    highlightBuilder.requireFieldMatch(false); //如果要多个字段高亮,这项要为
22    false
23    highlightBuilder.preTags("<span style=\"color:red\">"); //高亮设置
24    highlightBuilder.postTags("</span>");
25    sourceBuilder.highlighter(highlightBuilder);
26    sourceBuilder.timeout(new TimeValue(60, TimeUnit.SECONDS));
27
28    // 搜索
29    searchRequest.source(sourceBuilder);
30    SearchResponse response = restHighLevelClient.search(searchRequest,
    RequestOptions.DEFAULT);
```

```

31 // 解析结果!
32 List<Map<String, Object>> list = new ArrayList<>();
33 for (SearchHit hit : response.getHits()) {
34     //获取高亮字段
35     Map<String, HighlightField> highlightFields =
hit.getHighlightFields();
36     HighlightField titleField = highlightFields.get("title");
37
38     Map<String, Object> source = hit.getSourceAsMap();
39     //千万记得要记得判断是不是为空,不然你匹配的第一个结果没有高亮内容,那么就会报空指
    针异常,这个错误一开始真的搞了很久
40     if(titleField!=null){
41         Text[] fragments = titleField.fragments();
42         String name = "";
43         for (Text text : fragments) {
44             name += text;
45         }
46         source.put("title", name);    //高亮字段替换掉原本的内容
47     }
48     list.add(source);
49 }
50 return list;
51 }

```

2、controller层调用新的高亮业务!

```

1 //http://localhost:9090/search/java/1/10
2 @GetMapping("/search/{keyword}/{pageNo}/{pageSize}")
3 public List<Map<String, Object>> search(@PathVariable("keyword") String
keyword,
4                                         @PathVariable("pageNo") int pageNo,
5                                         @PathVariable("pageSize") int
pageSize) throws Exception {
6
7     return contentService.searchContentHighlighter(keyword,pageNo,pageSize);
8 }

```

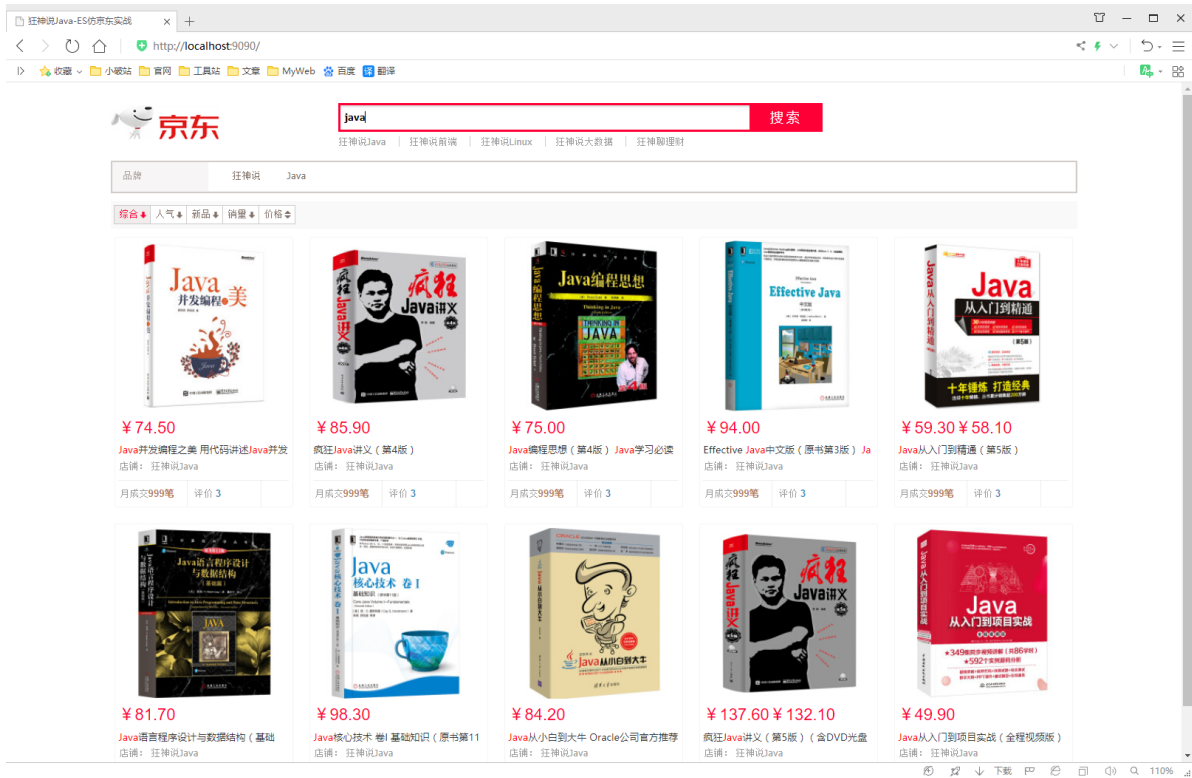
3、前端vue指令解析html!

```

1 <!--标题-->
2 <p class="productTitle">
3     <a v-html="result.title"> </a>
4 </p>

```

4、最终效果!



ES集群

为什么要实现集群

ES基本概念名词

Cluster

代表一个集群，集群中有多个节点，其中有一个为主节点，这个主节点是可以通过选举产生的，主从节点是对于集群内部来说的。es的一个概念就是去中心化，字面上理解就是无中心节点，这是对于集群外部来说的，因为从外部来看es集群，在逻辑上是个整体，你与任何一个节点的通信和与整个es集群通信是等价的。

Shards

代表索引分片，es可以把一个完整的索引分成多个分片，这样的好处是可以把一个大的索引拆分成多个，分布到不同的节点上。构成分布式搜索。分片的数量只能在索引创建前指定，并且索引创建后不能更改。

Replicas

代表索引副本，es可以设置多个索引的副本，副本的作用一是提高系统的容错性，当某个节点某个分片损坏或丢失时可以从副本中恢复。二是提高es的查询效率，es会自动对搜索请求进行负载均衡。

Recovery

代表数据恢复或叫数据重新分布，es在有节点加入或退出时会根据机器的负载对索引分片进行重新分配，挂掉的节点重新启动时也会进行数据恢复。

ES为什么要实现集群

在单台ES服务器节点上，随着业务量的发展索引文件慢慢增多，会影响到效率和内存存储问题等。

我们可以采用ES集群，将单个索引的分片到多个不同分布式物理机器上存储，从而可以实现高可用、容错性等。

ES集群中索引可能由多个分片构成，并且每个分片可以拥有多个副本。通过将一个单独的索引分为多个分片，我们可以处理不能在一个单一的服务器上面运行的大型索引，简单的说就是索引的大小过大，导致效率问题。不能运行的原因可能是内存也可能是存储。由于每个分片可以有多个副本，通过将副本分配到多个服务器，可以提高查询的负载能力。

ES是如何解决高并发

ES是一个分布式全文检索框架，隐藏了复杂的处理机制，内部使用分片机制、集群发现、分片负载均衡请求路由。

Shards 分片:代表索引分片，es可以把一个完整的索引分成多个分片，这样的好处是可以把一个大的索引拆分成多个，分布到不同的节点上。构成分布式搜索。分片的数量只能在索引创建前指定，并且索引创建后不能更改。

Replicas分片:代表索引副本，es可以设置多个索引的副本，副本的作用一是提高系统的容错性，当某个节点某个分片损坏或丢失时可以从副本中恢复。二是提高es的查询效率，es会自动对搜索请求进行负载均衡。

集群规划

搭建一个集群我们需要考虑如下几个问题：

- 1、我们需要多大规模的集群？
- 2、集群中的节点角色如何分配？
- 3、如何避免脑裂问题？
- 4、索引应该设置多少个分片？
- 5、分片应该设置几个副本？

下面我们就来分析和回答这几个问题：

我们需要多大规模的集群？

需要从以下两个方面考虑：

- 1、当前的数据量有多大？数据增长情况如何？
- 2、你的机器配置如何？cpu、多大内存、多大硬盘容量？

推算的依据：

ES JVM heap 最大可以设置32G 。

30G heap 大概能处理的数据量 10 T。如果内存很大如128G，可在一台机器上运行多个ES节点实例。

备注：集群规划满足当前数据规模+适量增长规模即可，后续可按需扩展。

两类应用场景：

A、用于构建业务搜索功能模块，且多是垂直领域的搜索。数据量级几千万到数十亿级别。一般2-4台机器的规模。

B、用于大规模数据的实时OLAP（联机处理分析），经典的如ELK Stack，数据规模可能达到千亿或更多。几十到上百节点的规模。

集群中的节点角色如何分配？

节点角色：

Master

node.master: true 节点可以作为主节点

DataNode

node.data: true 默认是数据节点。

Coordinate node 协调节点

如果仅担任协调节点，将上两个配置设为false。

说明：

一个节点可以充当一个或多个角色，默认三个角色都有

协调节点：一个节点只作为接收请求、转发请求到其他节点、汇总各个节点返回数据等功能的节点。就叫协调节点

如何分配：

A、小规模集群，不需严格区分。

B、中大规模集群（十个以上节点），应考虑单独的角色充当。特别并发查询量大，查询的合并量大，可以增加独立的协调节点。角色分开的好处是分工分开，不互影响。如不会因协调角色负载过高而影响数据节点的能力。

如何避免脑裂问题？

脑裂问题：

一个集群中只有一个A主节点，A主节点因为需要处理的东西太多或者网络过于繁忙，从而导致其他从节点ping不通A主节点，这样其他从节点就会认为A主节点不可用了，就会重新选出一个新的主节点B。过了一会A主节点恢复正常了，这样就出现了两个主节点，导致一部分数据来源于A主节点，另外一部分数据来源于B主节点，出现数据不一致问题，这就是**脑裂**。

尽量避免脑裂，需要添加最小数量的主节点配置：

discovery.zen.minimum_master_nodes: (有master资格节点数/2) + 1

这个参数控制的是，选举主节点时需要看到最少多少个具有master资格的活节点，才能进行选举。官方的推荐值是(N/2)+1，其中N是具有master资格的节点的数量。

常用做法（中大规模集群）：

1、Master 和 dataNode 角色分开，配置奇数个master，如3

2、单播发现机制，配置master资格节点：

discovery.zen.ping.multicast.enabled: false —— 关闭多播发现机制，默认是关闭的

discovery.zen.ping.unicast.hosts: ["master1", "master2", "master3"] —— 配置单播发现的主节点ip地址，其他从节点要加入进来，就得去询问单播发现机制里面配置的主节点我要加入到集群里面了，主节点同意以后才能加入，然后主节点再通知集群中的其他节点有新节点加入

3、配置选举发现数，及延长ping master的等待时长

discovery.zen.ping_timeout: 30（默认值是3秒）——其他节点ping主节点多久时间没有响应就认为主节点不可用了

discovery.zen.minimum_master_nodes: 2 —— 选举主节点时需要看到最少多少个具有master资格的活节点，才能进行选举

索引应该设置多少个分片？

说明：分片数指定后不可变，除非重索引。

思考：

分片对应的存储实体是什么？存储的实体是索引

分片是不是越多越好？不是

分片多有什么影响？分片多浪费存储空间、占用资源、影响性能

分片过多的影响：

每个分片本质上就是一个Lucene索引，因此会消耗相应的文件句柄，内存和CPU资源。

每个搜索请求会调度到索引的每个分片中。如果分片分散在不同的节点倒是问题不太。但当分片开始竞争相同的硬件资源时，性能便会逐步下降。

ES使用词频统计来计算相关性。当然这些统计也会分配到各个分片上。如果在大量分片上只维护了很少的数据，则将导致最终的文档相关性较差。

分片设置的可参考原则：

ElasticSearch推荐的最大JVM堆空间是30~32G，所以把你的分片最大容量限制为30GB，然后再对分片数量做合理估算。例如，你认为你的数据能达到200GB，推荐你最多分配7到8个分片。

在开始阶段，一个好的方案是根据你的节点数量按照1.5~3倍的原则来创建分片。例如，如果你有3个节点，则推荐你创建的分片数最多不超过9(3x3)个。当性能下降时，增加节点，ES会平衡分片的放置。

对于基于日期的索引需求，并且对索引数据的搜索场景非常少。也许这些索引量将达到成百上千，但每个索引的数据量只有1GB甚至更小。对于这种类似场景，建议只需要为索引分配1个分片。如日志管理就是一个日期的索引需求，日期索引会很多，但每个索引存放的日志数据量就很少。

分片应该设置几个副本？

说明：副本数是可以随时调整的！

思考：

副本的用途是什么？备份数据保证高可用数据不丢失，高并发的时候参与数据查询

针对它的用途，我们该如何设置它的副本数？一般一个分片有1-2个副本即可保证高可用

集群规模没变的情况下副本过多会有什么影响？副本多浪费存储空间、占用资源、影响性能

副本设置基本原则：

为保证高可用，副本数设置为2即可。要求集群至少要有3个节点，来分开存放主分片、副本。
如发现并发量大时，查询性能会下降，可增加副本数，来提升并发查询能力。

注意：新增副本时主节点会自动协调，然后拷贝数据到新增的副本节点

集群核心原理分析

1、每个索引会被分成多个分片shards进行存储，默认创建索引是分配5个分片进行存储。每个分片都会分布式部署在多个不同的节点上进行部署，该分片成为primary shards。

注意：索引的主分片primary shards定义好后，后面不能做修改。

2、为了实现高可用数据的高可用，主分片可以有对应的备分片replicas shards，replic shards分片承载了负责容错、以及请求的负载均衡。

注意：每一个主分片为了实现高可用，都会有自己对应的备分片，主分片对应的备分片不能存放同一台服务器上。主分片primary shards可以和其他replicas shards存放在同一个node节点上。

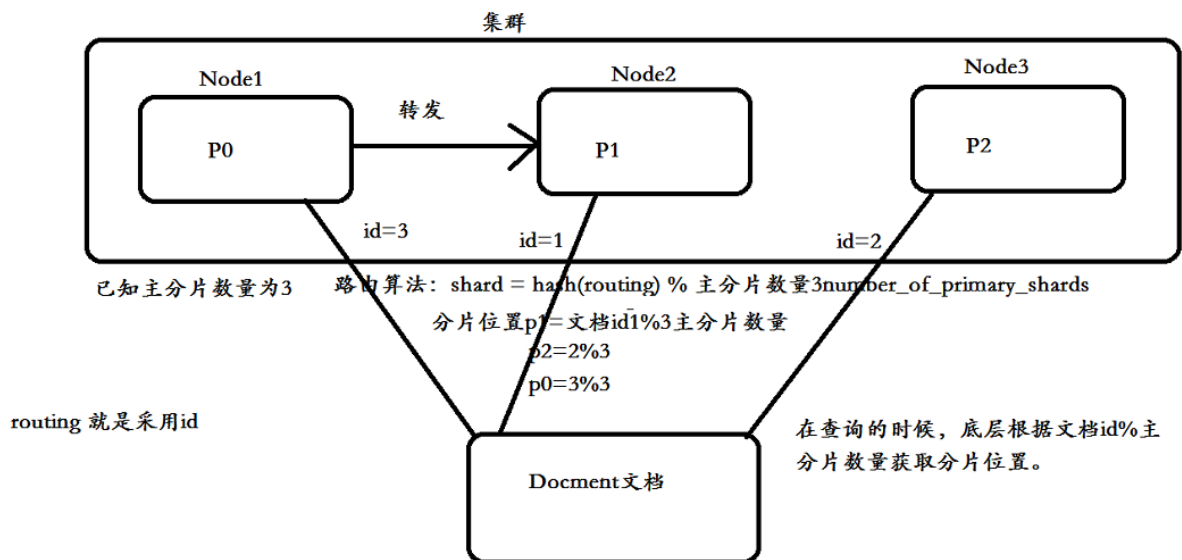
3、document routing（数据路由）

当客户端发起创建document的时候，es需要确定这个document放在该index哪个shard上。这个过程就是数据路由。

路由算法：shard = hash(routing) % number_of_primary_shards

如果number_of_primary_shards在查询的时候取余发生的变化，无法获取到该数据

注意：索引的主分片数量定义好后，不能被修改



1. ES核心存放的是那些？ 索引

2. 如果ES实现了集群的话，会将单台服务器节点的索引文件使用分片技术，分布式存放在多个不同的物理机器上。



3. 什么是分片技术？ 将数据拆分成多台节点进行存放。

4. 在ES分片技术中，分为主（primary）分片、副本Replicas分片

1. ES为了高可用，每个主的分片都会对应有一个备分的分片。

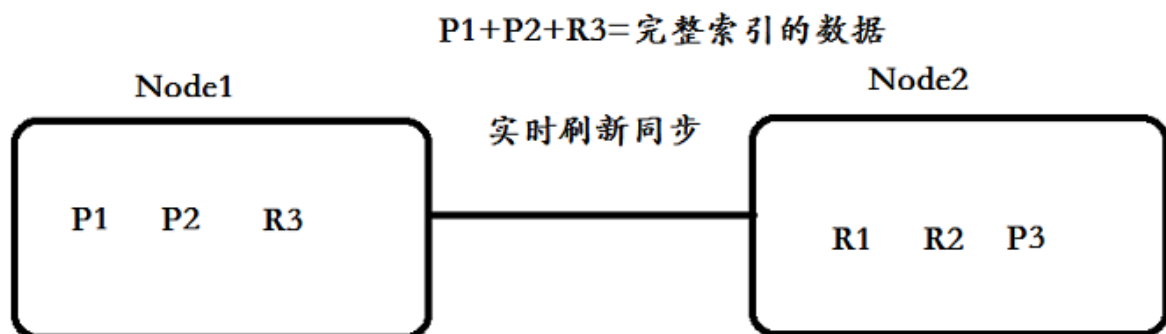
注意的是对应的主的分片和备的分片不能够在同一台节点上进行存放。

P 表示主分片

R 表示备分片

2. **单台ES服务器中是没有备用分片的**

$\text{number_of_shards} = 3 \quad 3 + 3 = 6$

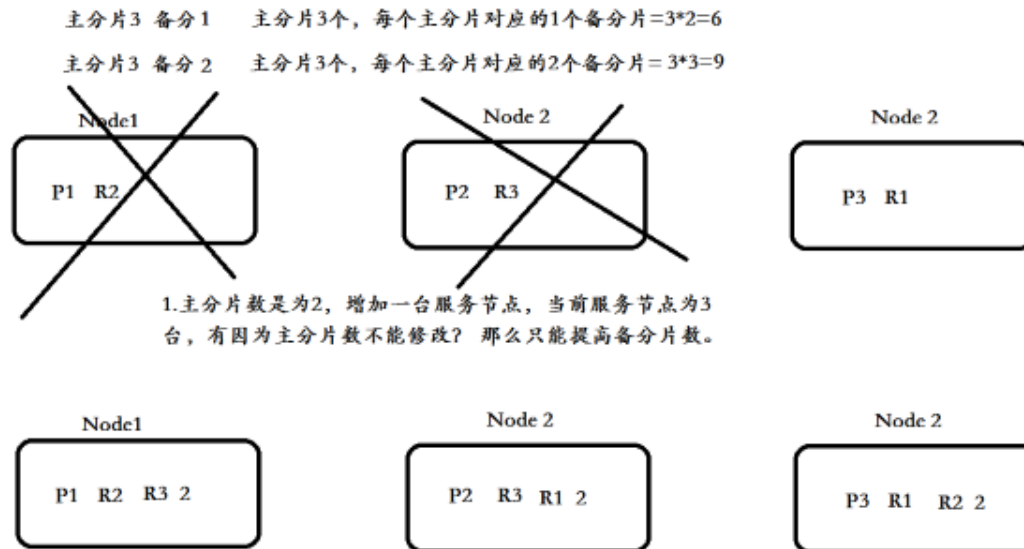


注意：在节点服务器上 既可存放主分片也可以存放备用分片

高可用视图分析

下图所示：上面的图，如果节点1与节点2宕机了，es集群数据就不完整了。

下图，如果节点1与节点2宕机了，es集群数据还是完整的



集群搭建

- 1、复制三个es的文件
- 2、进入elasticsearch的config目录，修改elasticsearch.yml的配置

```
1 # ===== Elasticsearch Configuration =====
2 # 配置es的集群名称，es会自动发现在同一网段下的es,如果在同一网段下有多个集群,就可以用这个
  属性来区分不同的集群。
3 cluster.name: elasticsearch
4 # 节点名称（要修改）
5 node.name: node-001
6 # 指定该节点是否有资格被选举成为node
7 node.master: true
8 # 指定该节点是否存储索引数据,默认为true。
9 node.data: true
10 # 设置绑定的ip地址还有其它节点和该节点交互的ip地址,本机ip
11 network.host: 127.0.0.1
12 # 指定http端口,你使用head,kopf等相关插件使用的端口 （要修改）
13 http.port: 9200
14 # 设置节点间交互的tcp端口,默认是9300。（要修改）
15 transport.tcp.port: 9300
16 #设置集群中master节点的初始列表,可以通过这些节点来自动发现新加入集群的节点。
17 #因为下两台elasticsearch的port端口会设置成9301 和 9302 所以写入两台#elasticsearch地
  址的完整路径
18 discovery.zen.ping.unicast.hosts:
19 ["127.0.0.1:9300","127.0.0.1:9301","127.0.0.1:9302"]
20 #如果要使用head,那么需要解决跨域问题,使head插件可以访问es
21 http.cors.enabled: true
  http.cors.allow-origin: ""
```

配置解析

IP访问限制、默认端口修改9200

这里有两个需要提醒下，第一个就是IP访问限制，第二个就是es实例的默认端口号9200。IP访问限制可以限定具体的IP访问服务器，这有一定的安全过滤作用。

```
1 # Set the bind address to a specific IP (IPv4 or IPv6):
2 #
3 network.host: 192.168.152.128
```

如果设置成0.0.0.0则是不限制任何IP访问。一般在生产的服务器可能会限定几台IP，通常用于管理使用。

默认的端口9200在一般情况下也有点风险，可以将默认的端口修改成另外一个，这还有一个原因就是怕开发人员误操作，连接上集群。当然，如果你的公司网络隔离做的很好也无所谓。

```
1 #
2 # Set a custom port for HTTP:
3 #
4 http.port: 9200
5 transport.tcp.port: 9300
```

这里的9300是集群内部通讯使用的端口，这个也可以修改掉。因为连接集群的方式有两种，通过扮演集群node也是可以进入集群的，所以还是安全起见，修改掉默认的端口。

说明：记得修改安装了ES的3台虚拟机（三个节点）的相同配置，要不然节点之间无法建立连接工作，也会报错。

集群发现IP列表、node、cluster名称

紧接着修改集群节点IP地址，这样可以让集群在规定的几个节点之间工作。elasticsearch，默认是使用自动发现IP机制。就是在当前网段内，只要能被自动感知到的IP就能自动加入到集群中。这有好处也有坏处。好处就是自动化了，当你的es集群需要云化的时候就会非常方便。但是也会带来一些不稳定的情况，如，master的选举问题、数据复制问题。

导致master选举的因素之一就是集群有节点进入。当数据复制发生的时候也会影响集群，因为要做数据平衡复制和冗余。这里面可以独立master集群，剔除master集群的数据节点能力。

固定列表的IP发现有两种配置方式，一种是互相依赖发现，一种是全量发现。各有优势吧，我是使用的依赖发现来做的。这有个很重要的参考标准，就是你的集群扩展速度有多快。因为这有个问题就是，当全量发现的时候，如果是初始化集群会有很大的问题，就是master全局会很长，然后节点之间的启动速度各不一样。所以我采用了靠谱点的依赖发现。

你需要在192.168.152.128的elasticsearch中配置成：

```
1 # ----- Discovery -----
2 #
3 # Pass an initial list of hosts to perform discovery when new node is
4 # started:
5 # The default list of hosts is ["127.0.0.1", "[::1]"]
6 #
7 discovery.zen.ping.unicast.hosts: [
8     "192.168.152.129:9300","192.168.152.130:9300" ]
```

让他去发现129，130的机器，以此类推，完成剩下的129和130机器的配置。

然后你需要配置下集群名称，就是你当前节点所在集群的名称，这有助于你规划你的集群。集群中的所有节点的集群名称必须一样，只有集群名称一样才能组成一个逻辑集群。


```
1 # ----- cluster -----  
  ---  
2 #  
3 # Use a descriptive name for your cluster:  
4 #  
5 cluster.name: mycluster
```

配置你当前节点的名称

```
1 # ----- Node -----  
  ---  
2 #  
3 # Use a descriptive name for the node:  
4 #  
5 node.name: node-1
```

以此类推，完成另外两个节点的配置。cluster.name的名称必须保持一致。然后分别设置node.name。

说明：

这里搭建的是一个简单的集群，没有做集群节点角色的区分，所以3个节点默认的角色有主节点、数据节点、协调节点

选举ES主节点的逻辑：

选举的大概逻辑，它会根据分片的数据的前后新鲜程度来作为选举的一个重要逻辑。（日志、数据、时间都会作为集群master全局的重要指标）

因为考虑到数据一致性问题，当然是用最新的数据节点作为master，然后进行新数据的复制和刷新其他node。