

# 聚合搜索平台项目笔记

## 项目介绍

一个企业级聚合搜索平台（简化版的搜索中台）。

项目的意义：

用户角度：允许用户在同一个页面集中搜索出不同来源、不同类型的内容，提升用户的检索效率和搜索体验。

企业角度：当企业有多个项目的数据需要被搜索时，无需针对每个项目单独开发搜索功能，可以直接将数据接入搜索中台，提升开发效率。

## 特点

1. 比较新颖，且项目周期不会很长（计划一个月），争取带大家快速搞定，成为简历上的亮眼项目
2. 选用 Vue 3 + Spring Boot 2.7 新版本实现，包含完整的前后端，并且依然是从 0 到 1 带大家开发
3. 前端讲得非常细；后端也包含较多的编程技巧，会带大家学习实践数据抓取 => 数据同步 => 搜索引擎等一系列业务开发
4. 项目中会分享很多架构设计知识，带大家提升一个思考的 level

## 项目资料

完整源码见星球代码库：<https://t.zsxq.com/08DElCxT7> <<https://t.zsxq.com/08DElCxT7>> （前端：yuso-frontend，后端：yuso-backend）

全部直播回放 + 大纲：[https://www.code-](https://www.code-nav.cn/course/1790979621621641217/section/1790981240778174466?type=)

[nav.cn/course/1790979621621641217/section/1790981240778174466?type=](https://www.code-nav.cn/course/1790979621621641217/section/1790981240778174466?type=)  
<[\[nav.cn/course/1790979621621641217/section/1790981240778174466?type=>\]\(https://www.code-nav.cn/course/1790979621621641217/section/1790981240778174466?type=\)](https://www.code-</a></p></div><div data-bbox=)

# 第一期

单集直播回放: <https://t.zsxq.com/0cVj0mO5v> <<https://t.zsxq.com/0cVj0mO5v>>

主要内容:

1. 项目介绍: 背景、技术栈、业务流程、项目计划
2. 前端项目初始化
3. 后端项目初始化
4. 万用项目模板介绍
5. 前端聚合搜索页面开发
6. 前后端联调

## 技术栈介绍

全栈项目

### 前端

- Vue
- Ant Design Vue
- Lodash

### 后端

- Spring Boot
- MySQL
- Elasticsearch (Elastic Stack) 搜索引擎
- 数据抓取
  - 离线
  - 实时
- 数据同步 (4 种同步方式)
  - 定时
  - 双写

◦ Logstash

◦ Canal

• JMeter 压力测试

• Guava Retrying 保证 API 的稳定性?

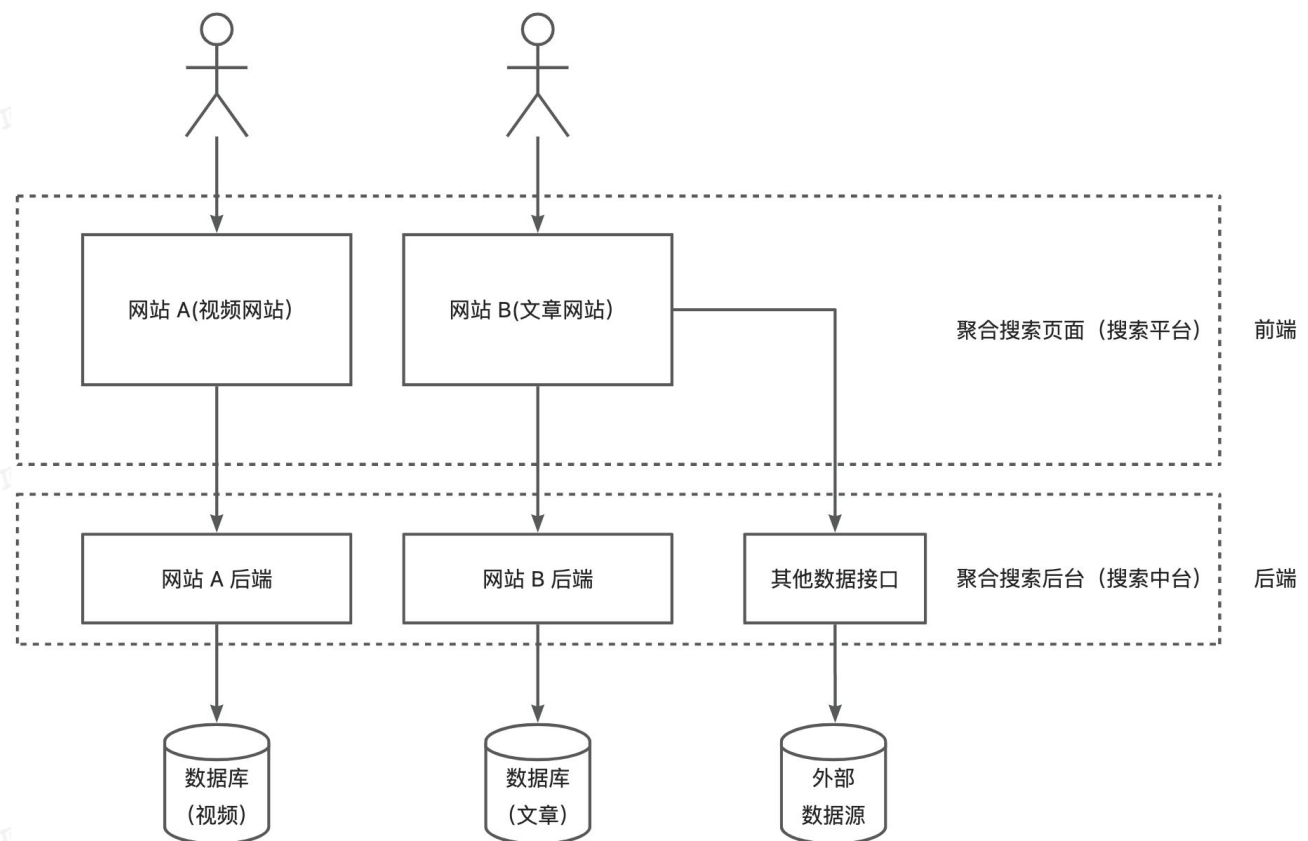
## 业务流程

1. 先得到各种不同分类的数据

2. 提供一个搜索页面（单一搜索 + 聚合搜索），支持搜索

还可以做一些优化，比如关键词高亮、搜索建议、防抖节流：

项目架构图：



## 前端初始化

步骤：

1. 参考 Ant Design 组件库的官方文档来搭建初始化项目

(<https://2x.antdv.com/docs/vue/getting-started-cn>

<<https://2x.antdv.com/docs/vue/getting-started-cn>> )，并整合组件库。

2. 删减不需要的页面和路由

## 后端初始化

步骤：

1. 使用星球 springboot-init 万用项目模板
2. 直接使用 swagger 文档在线测试接口

## 前端开发

组件库使用方式：从上到下依次在组件库文档中找到对应组件，复制粘贴 + 修改，完成基本页面开发。

## 记录搜索状态

目标：用 url 记录页面搜索状态，当用户刷新页面时，能够从 url 还原之前的搜索状态

需要双向同步：url  $\leftrightarrow$  页面状态

核心小技巧：把同步状态改成单向，即只允许 url 来改变页面状态，不允许反向

分步骤来实现，思路更清晰：

1. 让用户在操作的时候，改变 url 地址（点击搜索框，搜索内容填充到 url 上？切换 tab 时，也要填充）
2. 当 url 改变的时候，去改变页面状态（监听 url 的改变）

## 前后端联调

使用 Axios 向后端发送请求。

步骤（请参考官方文档：<https://www.axios-http.cn/docs/intro> <<https://www.axios-http.cn/docs/intro>>）：

1. 前端整合 Axios
2. 自定义 Axios 实例
3. 发送请求

## 第二期

单集直播回放: <https://t.zsxq.com/0ck9qfV1v> <<https://t.zsxq.com/0ck9qfV1v>>

主要内容:

1. 多数据源获取 (包含几种爬虫方式的讲解和实践)
  - a. 文章 (内部)
  - b. 用户 (内部)
  - c. 图片 (外部, 不是我们自己的项目、自己的用户生产的数据)
2. 前后端接口调通
3. 聚合搜索业务场景分析
4. 聚合搜索接口开发

## 获取不同类型的数据源

### 数据抓取流程

1. 分析数据源, 怎么获取?
2. 拿到数据后, 怎么处理?
3. 写入数据库等存储

### 数据抓取的几种方式

1. 直接请求数据接口 (最方便), 可使用 HttpClient、OKHttp、RestTemplate、Hutool (<https://hutool.cn/> <<https://hutool.cn/>> ) 等客户端发送请求
2. 等网页渲染出明文内容后, 从前端完整页面中解析出需要的内容
3. 有一些网站可能是动态请求的, 他不会一次性加载所有的数据, 而是要你点某个按钮、输入某个验证码才会显示出数据。可使用无头浏览器: selenium、node.js puppeteer

**注意, 爬虫技术不能滥用, 千万不要给别人的系统造成压力、不要侵犯他人权益!**

内部没有，可以从互联网上获取基础数据 => 爬虫

可使用该网站进行测试抓取：<https://www.code-nav.cn/learn/passage> <<https://www.code-nav.cn/learn/passage>> （注意！仅做测试，不要频繁请求！！！），获取到文章后要入库。

离线抓取：定时获取或者只获取一次

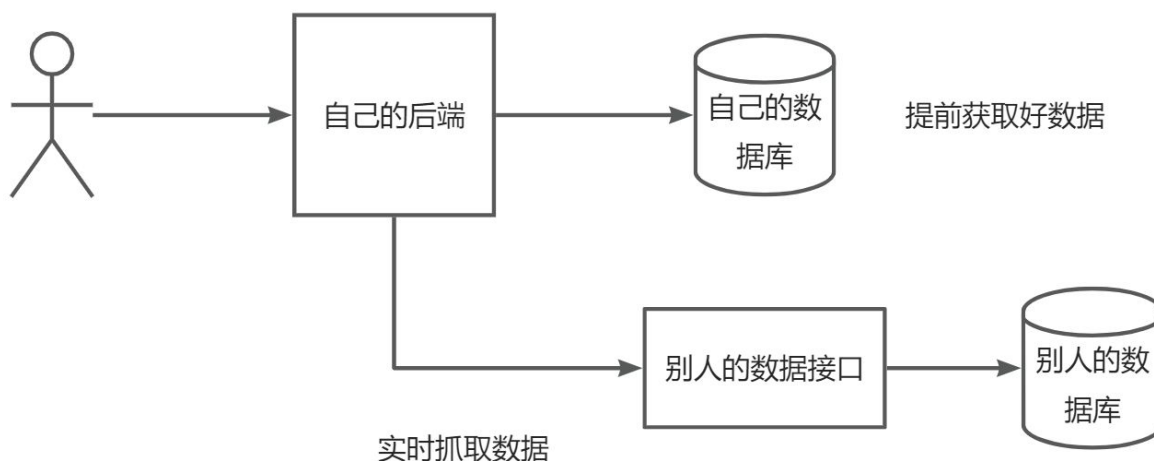
## 2、用户获取

每个网站的用户基本都是自己的，一般无需从站外获取。

## 3、图片获取

实时抓取：我们自己的数据库不存这些数据，用户要搜的时候，直接从别人的接口（网站 / 数据库）去搜。

流程如图：



jsoup 解析库：支持发送请求获取到 HTML 文档，然后从中解析出需要的字段。

## 现有业务场景分析

目前是在页面加载时，调用三个接口分别获取文章、图片、用户数据。

几种不同的业务场景：

- 1) 其实可以用用户点某个 tab 的时候，只调用这个 tab 的接口，比如：[https://www.code-nav.cn/search/all?current=2&pageSize=8&searchText=&sortField=\\_score&sortOrder=descend](https://www.code-nav.cn/search/all?current=2&pageSize=8&searchText=&sortField=_score&sortOrder=descend)

<[https://www.code-nav.cn/search/all?](https://www.code-nav.cn/search/all?current=2&pageSize=8&searchText=&sortField=_score&sortOrder=descend)

[current=2&pageSize=8&searchText=&sortField=\\_score&sortOrder=descend](https://www.code-nav.cn/search/all?current=2&pageSize=8&searchText=&sortField=_score&sortOrder=descend)> )

2) 如果是针对聚合内容的网页，其实可以一个请求搞定，比如：<https://tophub.today/>

<<https://tophub.today/>>

3) 有可能还要查询其他的信息，比如其他数据的总数，同时给用户反馈，比如 B 站搜索页

**你要根据实际情况去选择方式！**

目前设计存在的问题：

1. 请求数量比较多，可能会收到浏览器的限制
2. 请求不同接口的参数可能不一致，增加前后端沟通成本
3. 前端写调用多个接口的代码，重复代码

## 聚合接口

1) 请求数量比较多，可能会收到浏览器的限制 => 用一个接口请求完所有的数据（后端可以并发，几乎没有并发数量限制）

```
1 {  
2   user = userService.query  
3   post = postService.query  
4   picture = pictureService.query  
5   return user + post + picture  
6 }
```

2) 请求不同接口的参数可能不一致，增加前后端沟通成本 => 用一个接口把请求参数统一，前端每次传固定的参数，后端去对参数进行转换

```
1 {  
2   前端统一传 searchText  
3   后端把 searchText 转换为 userName => queryUser  
4 }
```

统一返回结果：比如都使用 Page 页面封装

3) 前端写调用多个接口的代码，重复代码 => 用一个接口，通过不同的参数去区分查询的数据源

```
1 {  
2   前端传 type 调用后端同一个接口，后端根据 type 调用不同的 service 查询  
3   比如: type = user, userService.query  
4 }
```

**注意，并发不一定更快！可能存在短板效应。要以实际测试结果为准！**

## 第三期

单集直播回放: <https://t.zsxq.com/0cKPukuOa> <<https://t.zsxq.com/0cKPukuOa>>

主要内容:

1. 搜索接口优化 (运用 3 种设计模式)
2. 前端调整搜索接口调用
3. 从 0 开始学习 Elastic Stack
  - a. Elasticsearch 安装
  - b. 入门实践
  - c. Kibana

## 聚合接口优化

思考: 怎么样能让前端又能一次搜出所有数据、又能够分别获取某一类数据 (比如分页场景)

解决方案:

新增 type 字段: 前端传 type 调用后端同一个接口, 后端根据 type 调用不同的 service 查询

比如前端传递 type = user, 后端执行 userService.query

逻辑;

1. 如果 type 为空, 那么搜索出所有的数据
2. 如果 type 不为空
  - a. 如果 type 合法, 那么查出对应数据
  - b. 否则报错

问题: type 增多后, 要把查询逻辑堆积在 controller 代码里么?

思考: 怎么能让搜索系统 **更轻松地** 接入更多的数据源?

## 门面模式

介绍: 帮助我们用户 (客户端) 去更轻松地实现功能, 不需要关心门面背后的细节。

聚合搜索类业务基本都是门面模式: 即前端不用关心后端从哪里、怎么去取不同来源、怎么去聚合不同来源的数据, 更方便地获取到内容。

当调用你系统（接口）的客户端觉得麻烦的时候，你就应该思考，是不是可以抽象一个门面了。

## 适配器模式

1) 定制统一的数据源接入规范（标准）：

- 什么数据源允许接入？
- 你的数据源接入时要满足什么要求？
- 需要接入方注意什么事情？

本系统要求：任何接入我们系统的数据，它必须要能够根据关键词搜索、并且支持分页搜索。

通过声明接口的方式来定义规范。

2) 假如说我们的数据源已经支持了搜索，但是原有的方法参数和我们的规范不一致，怎么办？

使用适配器模式：通过转换，让两个系统能够完成对接。

## 注册器模式（本质也是单例）

提前通过一个 map 或者其他类型存储好后面需要调用的对象。

效果：替代了 if... else..., 代码量大幅度减少，可维护可扩展。

## 搜索优化

现有问题：搜索不够灵活。

比如搜“鱼皮rapper”无法搜到“鱼皮是 rapper”，因为 MySQL 数据库的 like 是包含查询。

需要分词搜索

## Elastic Stack（一套技术栈）

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

包含了数据的整合 => 提取 => 存储 => 使用，一整套！

各组件介绍：

- beats 套件：从各种不同类型的文件 / 应用中采集数据。比如：a,b,c,d,e,aa,bb,cc
- Logstash：从多个采集器或数据源来抽取 / 转换数据，向 es 输送。比如：a,bb,cc
- elasticsearch：存储、查询数据

- kibana: 可视化 es 的数据

## 安装 ES

elasticsearch: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/setup.html>

<<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/setup.html>>

<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/zip-windows.html>

<<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/zip-windows.html>>

只要是一套技术，所有版本必须一致!!! 此处都用 7.17 版本。

## 安装 Kibana

kibana: <https://www.elastic.co/guide/en/kibana/7.17/introduction.html>

<<https://www.elastic.co/guide/en/kibana/7.17/introduction.html>>

<https://www.elastic.co/guide/en/kibana/7.17/install.html>

<<https://www.elastic.co/guide/en/kibana/7.17/install.html>>

## Elasticsearch 概念

你就把它当成 MySQL 一样的数据库去学习和理解。

入门学习:

- Index 索引 => MySQL 里的表 (table)
- 建表、增删改查 (查询需要花费的学习时间最多)
- 用客户端去调用 ElasticSearch (3 种)
- 语法: SQL、代码的方法 (4 种语法)

ES 相比于 MySQL, 能够自动帮我们做分词, 能够非常高效、灵活地查询内容。

## 索引 (倒排索引)

正向索引: 理解为书籍的目录, 可以快速帮你找到对应的内容 (怎么根据页码找到文章)

倒排索引: 怎么根据内容找到文章

文章 A: 你好, 我是 rapper

文章 B: 鱼皮你好, 我是 coder

切词:

你好, 我是, rapper

鱼皮, 你好, 我是, coder

构建倒排索引表:

| 词      | 内容 id   |
|--------|---------|
| 你好     | 文章 A, B |
| 我是     | 文章 A, B |
| rapper | 文章 A    |
| 鱼皮     | 文章 B    |
| coder  | 文章 B    |

用户搜: "鱼皮 rapper"

ES 先切词: 鱼皮, rapper

然后去倒排索引表找对应的文章

## ES 的几种调用方式

### 1) restful api 调用 (http 请求)

GET 请求: <http://localhost:9200/> <<http://localhost:9200/>>

curl 可以模拟发送请求: curl -X GET "localhost:9200/?pretty"

ES 的启动端口

1. 9200: 给外部用户 (给客户端调用) 的端口
2. 9300: 给 ES 集群内部通信的 (外部调用不了的)

### 2) kibana devtools

自由地对 ES 进行操作 (本质也是 restful api)

devtools 不建议生产环境使用

### 3) 客户端调用

java 客户端、go 客户端等。

参考文档: [https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/\\_getting\\_started.html](https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/_getting_started.html)  
<[https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/\\_getting\\_started.html](https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/_getting_started.html)>

## ES 的语法

### DSL

json 格式, 好理解; 和 http 请求最兼容, 应用最广, 也是鱼皮个人比较推荐的

#### 建表、插入数据

```
1  POST post/_doc/ro1R-oYBDwC2db8WGd1L
2  {
3    "title": "鱼皮2",
4    "desc": "鱼皮的描述2"
5  }
```

#### 查询

DSL 语法: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/query-dsl.html>  
<<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/query-dsl.html>> (忘了就查, 不用背)

```
1 GET post/_search
2 {
3   "query": {
4     "match_all": { }
5   },
6   "sort": [
7     {
8       "@timestamp": "desc"
9     }
10  ]
11 }
```

根据 id 查询:

```
1 GET post/_doc/ro1R-oYBDwC2db8WGd1L
```

## 修改

```
1 POST post/_doc/ro1R-oYBDwC2db8WGd1L
2 {
3   "title": "鱼皮2",
4   "desc": "鱼皮的描述2"
5 }
```

## 删除

删除普通索引:

```
1 DELETE index_name
```

删除数据流式索引:

```
1 DELETE _data_stream/logs-my_app-default
```

## EQl

专门查询 ECS 文档 (标准指标文档) 的数据的语法, 更加规范, 但只适用于特定场景 (比如事件流)

文档: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/eql.html>

<<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/eql.html>>

示例:

```
1  POST my_event/_doc
2  {
3    "title": "鱼333333皮",
4    "@timestamp": "2099-05-06T16:21:15.000Z",
5    "event": {
6      "original": "192.0.2.42 - - [06/May/2099:16:21:15 +0000] \"GET /images/bg
7    }
8  }
9
10
11  GET my_event/_eql/search
12  {
13    "query": ""
14    any where 1 == 1
15    ""
16  }
```

## SQL

文档: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/sql-getting-started.html> <<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/sql-getting-started.html>>

学习成本低, 但是可能需要插件支持、**性能较差**

示例:

```
1  POST /_sql?format=txt
2  {
3    "query": "SELECT * FROM post where title like '%鱼皮%'"
4  }
```

## Painless Scripting language

编程式取值, 更灵活, 但是学习成本高

## Mapping

文档: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/explicit-mapping.html>  
<<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/explicit-mapping.html>>

可以理解为数据库的表结构, 有哪些字段、字段类型。

ES 支持动态 mapping, 表结构可以动态改变, 而不像 MySQL 一样必须手动建表, 没有的字段就不能插入。

显示创建 mapping:

```
1  GET user/_mapping
2
3  PUT user
4  {
5    "mappings": {
6      "properties": {
7        "age": { "type": "integer" },
8        "email": { "type": "keyword" },
9        "name": { "type": "text" }
10     }
11   }
12 }
```

## 第四期

单集直播回放: <https://t.zsxq.com/0crCWmhHC> <<https://t.zsxq.com/0crCWmhHC>>

主要内容:

1. ES 搜索引擎实战 (2 种 Java 客户端操作方式)
2. 数据同步实战 (4 种同步方式)
3. Kibana 搭建看板
4. JMeter 接口性能测试
5. 其他扩展思路

## ElasticStack 概念

ES 索引 (Index) => 表

ES field (字段) => 列

倒排索引

调用语法 (DSL、EQL、SQL 等)

Mapping 表结构

- 自动生成 mapping
- 手动指定 mapping

## 分词器

指定了分词的规则。

内置分词器: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/analysis-analyzers.html> <<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/analysis-analyzers.html>>

空格分词器: whitespace, 结果 The、quick、brown、fox.

示例:

```
1 POST _analyze
2 {
3   "analyzer": "whitespace",
4   "text":     "The quick brown fox."
5 }
```

标准分词规则, 结果: is、this、déjà、vu

```
1 POST _analyze
2 {
3   "tokenizer": "standard",
4   "filter": [ "lowercase", "asciifolding" ],
5   "text":     "Is this déjà vu?"
6 }
```

关键词分词器: 就是不分词, 整句话当作专业术语

```
1 POST _analyze
2 {
3   "analyzer": "keyword",
4   "text":     "The quick brown fox."
5 }
```

## IK 分词器 (ES 插件)

中文友好: <https://github.com/medcl/elasticsearch-analysis-ik>

<<https://github.com/medcl/elasticsearch-analysis-ik>>

下载地址: <https://github.com/medcl/elasticsearch-analysis-ik/releases/tag/v7.17.7>

<<https://github.com/medcl/elasticsearch-analysis-ik/releases/tag/v7.17.7>> (注意版本一致)

思考: 怎么样让 ik 按自己的想法分词?

解决方案: 自定义词典 (自己尝试)

ik\_smart 和 ik\_max\_word 的区别? 举例: “小黑子”

ik\_smart 是智能分词, 尽量选择最像一个词的拆分方式, 比如 “小”、“黑子”

ik\_max\_word 尽可能地分词, 可以包括组合词, 比如 “小黑”、“黑子”

## 打分机制

比如有 3 条内容:

1. 鱼皮是狗
2. 鱼皮是小黑子
3. 我是小黑子

用户搜索:

1. 鱼皮, 第一条分数最高, 因为第一条匹配了关键词, 而且更短 (匹配比例更大)
2. 鱼皮小黑子 => 鱼皮、小、黑子, 排序结果: 2 > 3 > 1

参考文章: <https://liyupi.blog.csdn.net/article/details/119176943>

<<https://liyupi.blog.csdn.net/article/details/119176943>>

官方参考文章: <https://www.elastic.co/guide/en/elasticsearch/guide/master/controlling-relevance.html> <<https://www.elastic.co/guide/en/elasticsearch/guide/master/controlling-relevance.html>>

## ES 调用方式

3 种:

1. HTTP Restful 调用
2. kibana 操作 (dev tools)
3. 客户端操作 (Java)

## Java 操作 ES

3 种方式:

1) ES 官方的 Java API

<https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/introduction.html>  
<<https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/introduction.html>>

快速开始: <https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/connecting.html> <<https://www.elastic.co/guide/en/elasticsearch/client/java-api-client/7.17/connecting.html>>

2) ES 以前的官方 Java API, HighLevelRestClient (已废弃, 不建议用)

3) Spring Data Elasticsearch (推荐)

spring-data 系列: spring 提供的操作数据的框架

spring-data-redis: 操作 redis 的一套方法

spring-data-mongodb: 操作 mongodb 的一套方法

spring-data-elasticsearch: 操作 elasticsearch 的一套方法

官方文档: <https://docs.spring.io/spring-data/elasticsearch/docs/4.4.10/reference/html/>  
<<https://docs.spring.io/spring-data/elasticsearch/docs/4.4.10/reference/html/>>

自定义方法: 用户可以指定接口的方法名称, 框架帮你自动生成查询

## 用 ES 实现搜索接口

### 1、建表 (建立索引)

数据库表结构:

```

1  -- 帖子表
2  create table if not exists post
3  (
4      id          bigint auto_increment comment 'id' primary key,
5      title       varchar(512)                null comment '标题',
6      content     text                        null comment '内容',
7      tags        varchar(1024)              null comment '标签列表 (json)',
8      thumbNum    int      default 0          not null comment '点赞数',
9      favourNum   int      default 0          not null comment '收藏数',
10     userId      bigint                      not null comment '创建用户 id',
11     createTime  datetime default CURRENT_TIMESTAMP not null comment '创建时间',
12     updateTime  datetime default CURRENT_TIMESTAMP not null on update CURRENT_TIMESTAMP,
13     isDelete    tinyint  default 0          not null comment '是否删除',
14     index idx_userId (userId)
15 ) comment '帖子' collate = utf8mb4_unicode_ci;

```

## ES Mapping:

id (可以不放到字段设置里)

ES 中, 尽量存放需要用户筛选 (搜索) 的数据

aliases: 别名 (为了后续方便数据迁移)

字段类型是 text, 这个字段是可被分词的、可模糊查询的; 而如果是 keyword, 只能完全匹配、精确查询。

analyzer (存储时生效的分词器): 用 ik\_max\_word, 拆的更碎、索引更多, 更有可能被搜出来

search\_analyzer (查询时生效的分词器): 用 ik\_smart, 更偏向于用户想搜的分词

如果想要让 text 类型的分词字段也支持精确查询, 可以创建 keyword 类型的子字段:

```

1  "fields": {
2      "keyword": {
3          "type": "keyword",
4          "ignore_above": 256 // 超过字符数则忽略查询
5      }
6  }

```

建表结构:

```
1  POST post_v1
2  {
3    "aliases": {
4      "post": {}
5    },
6    "mappings": {
7      "properties": {
8        "title": {
9          "type": "text",
10         "analyzer": "ik_max_word",
11         "search_analyzer": "ik_smart",
12         "fields": {
13           "keyword": {
14             "type": "keyword",
15             "ignore_above": 256
16           }
17         }
18       },
19       "content": {
20         "type": "text",
21         "analyzer": "ik_max_word",
22         "search_analyzer": "ik_smart",
23         "fields": {
24           "keyword": {
25             "type": "keyword",
26             "ignore_above": 256
27           }
28         }
29       },
30       "tags": {
31         "type": "keyword"
32       },
33       "userId": {
34         "type": "keyword"
35       },
36       "createTime": {
37         "type": "date"
38       },
39       "updateTime": {
40         "type": "date"
41       },
42       "isDelete": {
43         "type": "keyword"
44       }
45     }
46   }
47 }
```

## 2、增删改查

第一种方式：ElasticsearchRepository<PostEsDTO, Long>，默认提供了简单的增删改查，多用于可预期的、相对没那么复杂的查询、自定义查询，返回结果相对简单直接。

接口代码：

```
1 public interface CrudRepository<T, ID> extends Repository<T, ID> {
2     <S extends T> S save(S entity);
3
4     <S extends T> Iterable<S> saveAll(Iterable<S> entities);
5
6     Optional<T> findById(ID id);
7
8     boolean existsById(ID id);
9
10    Iterable<T> findAll();
11
12    Iterable<T> findAllById(Iterable<ID> ids);
13
14    long count();
15
16    void deleteById(ID id);
17
18    void delete(T entity);
19
20    void deleteAllById(Iterable<? extends ID> ids);
21
22    void deleteAll(Iterable<? extends T> entities);
23
24    void deleteAll();
25 }
```

ES 中，\_开头的字段表示系统默认字段，比如 \_id，如果系统不指定，会自动生成。但是不会在 \_source 字段中补充 id 的值，所以建议大家手动指定。

支持根据方法名自动生成方法，比如：

```
1 List<PostEsDTO> findByTitle(String title);
```

第二种方式：Spring 默认给我们提供的操作 es 的客户端对象 ElasticsearchRestTemplate，也提供了增删改查，它的增删改查更灵活，适用于更复杂的操作，返回结果更完整，但需要自己解析。

**对于复杂的查询，建议用第二种方式。**

三个步骤：

1. 取参数

2. 把参数组合为 ES 支持的搜索条件

### 3. 从返回值中取结果

## 3、查询 DSL

参考文档：

- <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/query-filter-context.html> <<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/query-filter-context.html>>
- <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/query-dsl-bool-query.html> <<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/query-dsl-bool-query.html>>

示例代码：

```
1 GET post/_search
2 {
3   "query": {
4     "bool": { // 组合条件
5       "must": [ // 必须都满足
6         { "match": { "title": "鱼皮" }}, // match 模糊查询
7         { "match": { "content": "知识星球" }}
8       ],
9       "filter": [
10        { "term": { "status": "published" }}, // term 精确查询
11        { "range": { "publish_date": { "gte": "2015-01-01" }}} // range 范围查
12      ]
13    }
14  }
15 }
```

wildcard 模糊查询

regexp 正则匹配查询

查询结果中，score 代表匹配分数

建议先测试 DSL、再翻译成 Java

```
1  {
2  "query": {
3    "bool": {
4      "must_not": [
5        {
6          "match": {
7            "title": ""
8          }
9        },
10     ]
11    "should": [
12      {
13        "match": {
14          "title": ""
15        }
16      },
17      {
18        "match": {
19          "desc": ""
20        }
21      }
22    ],
23    "filter": [
24      {
25        "term": {
26          "isDelete": 0
27        }
28      },
29      {
30        "term": {
31          "id": 1
32        }
33      },
34      {
35        "term": {
36          "tags": "java"
37        }
38      },
39      {
40        "term": {
41          "tags": "框架"
42        }
43      }
44    ],
45    "minimum_should_match": 0
46  }
47 },
48 "from": 0, // 分页
49 "size": 5, // 分页
50 "_source": ["name", "_createTime", "desc", "reviewStatus", "priority", "tag
```

```
51     "sort": [ // 排序
52         {
53             "priority": {
54                 "order": "desc"
55             }
56         },
57         {
58             "_score": {
59                 "order": "desc"
60             }
61         },
62         {
63             "publishTime": {
64                 "order": "desc"
65             }
66         }
67     ]
68 }
```

翻译为 Java:

```

1      BoolQueryBuilder boolQueryBuilder = QueryBuilders.boolQuery();
2      // 过滤
3      boolQueryBuilder.filter(QueryBuilders.termQuery("isDelete", 0));
4      if (id != null) {
5          boolQueryBuilder.filter(QueryBuilders.termQuery("id", id));
6      }
7      if (notId != null) {
8          boolQueryBuilder.mustNot(QueryBuilders.termQuery("id", notId));
9      }
10     if (userId != null) {
11         boolQueryBuilder.filter(QueryBuilders.termQuery("userId", userId));
12     }
13     // 必须包含所有标签
14     if (CollectionUtils.isNotEmpty(tagList)) {
15         for (String tag : tagList) {
16             boolQueryBuilder.filter(QueryBuilders.termQuery("tags", tag));
17         }
18     }
19     // 包含任何一个标签即可
20     if (CollectionUtils.isNotEmpty(orTagList)) {
21         BoolQueryBuilder orTagBoolQueryBuilder = QueryBuilders.boolQuery(
22             for (String tag : orTagList) {
23                 orTagBoolQueryBuilder.should(QueryBuilders.termQuery("tags",
24                     tag));
25                 orTagBoolQueryBuilder.minimumShouldMatch(1);
26                 boolQueryBuilder.filter(orTagBoolQueryBuilder);
27             }
28         // 按关键词检索
29         if (StringUtils.isNotBlank(searchText)) {
30             boolQueryBuilder.should(QueryBuilders.matchQuery("title", searchText));
31             boolQueryBuilder.should(QueryBuilders.matchQuery("content", searchText));
32             boolQueryBuilder.minimumShouldMatch(1);
33         }
34         // 按标题检索
35         if (StringUtils.isNotBlank(title)) {
36             boolQueryBuilder.should(QueryBuilders.matchQuery("title", title));
37             boolQueryBuilder.minimumShouldMatch(1);
38         }
39         // 按内容检索
40         if (StringUtils.isNotBlank(content)) {
41             boolQueryBuilder.should(QueryBuilders.matchQuery("content", content));
42             boolQueryBuilder.minimumShouldMatch(1);
43         }
44         // 排序
45         SortBuilder<?> sortBuilder = SortBuilders.scoreSort();
46         if (StringUtils.isNotBlank(sortField)) {
47             sortBuilder = SortBuilders.fieldSort(sortField);
48             sortBuilder.order(CommonConstant.SORT_ORDER_ASC.equals(sortOrder)
49                 ? SortBuilders.ORDER_ASCENDING : SortBuilders.ORDER_DESCENDING);
50         }
51         // 分页

```

```

51     PageRequest pageRequest = PageRequest.of((int) current, (int) pageSiz
52     // 构造查询
53     NativeSearchQuery searchQuery = new NativeSearchQueryBuilder().withQu
54     .withPageable(pageRequest).withSorts(sortBuilder).build();
55     SearchHits<PostEsDTO> searchHits = elasticSearchRestTemplate.search(s

```

动静分离设计：先模糊筛选静态数据，查出数据后，再根据查到的内容 id 去数据库查找到 **动态数据**。

## 数据同步

一般情况下，如果做查询搜索功能，使用 ES 来模糊搜索，但是数据是存放在数据库 MySQL 里的，所以说我们需要把 MySQL 中的数据和 ES 进行同步，保证数据一致（以 MySQL 为主）。

MySQL => ES（单向）

首次安装完 ES，把 MySQL 数据全量同步到 ES 里，写一个单次脚本

4 种方式，全量同步（首次）+ 增量同步（新数据）：

1. 定时任务，比如 1 分钟 1 次，找到 MySQL 中过去几分钟内（至少是定时周期的 2 倍）发生改变的数据，然后更新到 ES。

优点：简单易懂、占用资源少、不用引入第三方中间件

缺点：有时间差

应用场景：数据短时间内不同步影响不大、或者数据几乎不发生修改

2. 双写：写数据的时候，必须也去写 ES；更新删除数据库同理。（事务：建议先保证 MySQL 写成功，如果 ES 写失败了，可以通过定时任务 + 日志 + 告警进行检测和修复（补偿））
3. 用 Logstash 数据同步管道（一般要配合 kafka 消息队列 + beats 采集器）：
4. Canal 监听 MySQL Binlog，实时同步

## Logstash

**传输** 和 **处理** 数据的管道

<https://www.elastic.co/guide/en/logstash/7.17/getting-started-with-logstash.html>

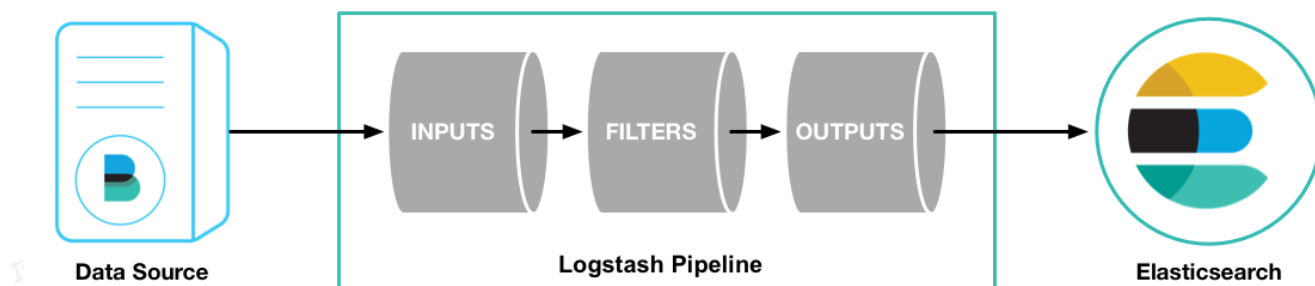
<<https://www.elastic.co/guide/en/logstash/7.17/getting-started-with-logstash.html>>

[https://artifacts.elastic.co/downloads/logstash/logstash-7.17.9-windows-x86\\_64.zip](https://artifacts.elastic.co/downloads/logstash/logstash-7.17.9-windows-x86_64.zip)

<[https://artifacts.elastic.co/downloads/logstash/logstash-7.17.9-windows-x86\\_64.zip](https://artifacts.elastic.co/downloads/logstash/logstash-7.17.9-windows-x86_64.zip)>

好处：用起来方便，插件多

缺点：成本更大、一般要配合其他组件使用（比如 kafka）



事件 Demo:

```
1 cd logstash-7.17.9
2 .\bin\logstash.bat -e "input { stdin { } } output { stdout { } }"
```

快速开始文档: <https://www.elastic.co/guide/en/logstash/7.17/running-logstash-windows.html> <<https://www.elastic.co/guide/en/logstash/7.17/running-logstash-windows.html>>

监听 udp 并输出:

```
1 # Sample Logstash configuration for receiving
2 # UDP syslog messages over port 514
3
4 input {
5   udp {
6     port => 514
7     type => "syslog"
8   }
9 }
10
11 output {
12   stdout { codec => rubydebug }
13 }
```

要把 MySQL 同步给 Elasticsearch。

问题 1: 找不到 mysql 的包

Error: unable to load mysql-connector-java-5.1.36-bin.jar from :jdbc\_driver\_library, file not readable (please check user and group permissions for the path)

Exception: LogStash::PluginLoadingError

解决: 修改 Logstash 任务配置中的 jdbc\_driver\_library 为驱动包的绝对路径 (驱动包可以从 maven 仓库中拷贝)

增量配置: 是不是可以只查最新更新的? 可以记录上次更新的数据时间, 只查出来 > 该更新时间的数据

小知识: 预编译 SQL 的优点?

1. 灵活

2. 模板好懂
3. 快 (有缓存)
4. 部分防注入

sql\_last\_value 是取上次查到的数据的最后一行的指定的字段, 如果要全量更新, 只要删除掉 E:\software\ElasticStack\logstash-7.17.9\data\plugins\inputs\jdbc\logstash\_jdbc\_last\_run 文件即可 (这个文件存储了上次同步到的数据)

```
1 input {
2   jdbc {
3     jdbc_driver_library => "E:\software\ElasticStack\logstash-7.17.9\config\m
4     jdbc_driver_class => "com.mysql.jdbc.Driver"
5     jdbc_connection_string => "jdbc:mysql://localhost:3306/my_db"
6     jdbc_user => "root"
7     jdbc_password => "123456"
8     statement => "SELECT * from post where updateTime > :sql_last_value"
9     tracking_column => "updatetime"
10    tracking_column_type => "timestamp"
11    use_column_value => true
12    parameters => { "favorite_artist" => "Beethoven" }
13    schedule => "*/5 * * * * *"
14    jdbc_default_timezone => "Asia/Shanghai"
15  }
16 }
17
18 output {
19   stdout { codec => rubydebug }
20 }
21
```

注意查询语句中要按 updateTime 排序, 保证最后一条是最大的:

```

1 input {
2   jdbc {
3     jdbc_driver_library => "E:\software\ElasticStack\logstash-7.17.9\config\m
4     jdbc_driver_class => "com.mysql.jdbc.Driver"
5     jdbc_connection_string => "jdbc:mysql://localhost:3306/my_db"
6     jdbc_user => "root"
7     jdbc_password => "123456"
8     statement => "SELECT * from post where updateTime > :sql_last_value and u
9     tracking_column => "updatetime"
10    tracking_column_type => "timestamp"
11    use_column_value => true
12    parameters => { "favorite_artist" => "Beethoven" }
13    schedule => "*/5 * * * * *"
14    jdbc_default_timezone => "Asia/Shanghai"
15  }
16 }
17
18 output {
19   stdout { codec => rubydebug }
20
21   elasticsearch {
22     hosts => "http://localhost:9200"
23     index => "post_v1"
24     document_id => "%{id}"
25   }
26 }

```

两个问题：

1. 字段全变成小写了
2. 多了一些我们不想同步的字段

可以编写过滤：

```

1 input {
2   jdbc {
3     jdbc_driver_library => "E:\software\ElasticStack\logstash-7.17.9\config\m
4     jdbc_driver_class => "com.mysql.jdbc.Driver"
5     jdbc_connection_string => "jdbc:mysql://localhost:3306/my_db"
6     jdbc_user => "root"
7     jdbc_password => "123456"
8     statement => "SELECT * from post where updateTime > :sql_last_value and u
9     tracking_column => "updatetime"
10    tracking_column_type => "timestamp"
11    use_column_value => true
12    parameters => { "favorite_artist" => "Beethoven" }
13    schedule => "*/5 * * * * *"
14    jdbc_default_timezone => "Asia/Shanghai"
15  }
16 }
17
18 filter {
19   mutate {
20     rename => {
21       "updatetime" => "updateTime"
22       "userid" => "userId"
23       "createtime" => "createTime"
24       "isdelete" => "isDelete"
25     }
26     remove_field => ["thumbnum", "favournum"]
27   }
28 }
29
30 output {
31   stdout { codec => rubydebug }
32
33   elasticsearch {
34     hosts => "127.0.0.1:9200"
35     index => "post_v1"
36     document_id => "%{id}"
37   }
38 }

```

## 订阅数据库流水的同步方式 Canal

<https://github.com/alibaba/canal/> <<https://github.com/alibaba/canal/>>

优点：实时同步，实时性非常强

原理：数据库每次修改时，会修改 binlog 文件，只要监听该文件的修改，就能第一时间得到消息并处理

canal: 帮你监听 binlog, 并解析 binlog 为你可以理解的内容。

它伪装成了 MySQL 的从节点, 获取主节点给的 binlog, 如图:

快速开始: <https://github.com/alibaba/canal/wiki/QuickStart>

<<https://github.com/alibaba/canal/wiki/QuickStart>>

windows 系统, 找到你本地的 mysql 安装目录, 在根目录下新建 my.ini 文件:

```
1 [mysqld]
2 log-bin=mysql-bin # 开启 binlog
3 binlog-format=ROW # 选择 ROW 模式
4 server_id=1 # 配置 MySQL replaction 需要定义, 不要和 canal 的 slaveId 重复
```

如果 java 找不到, 修改 startup.bat 脚本为你自己的 java home:

```
1
2 set JAVA_HOME=C:\Users\59278\.jdk\corretto-1.8.0_302
3 echo %JAVA_HOME%
4 set PATH=%JAVA_HOME%\bin;%PATH%
5 echo %PATH%
6
```

问题: mysql 无法链接, Caused by: java.io.IOException: caching\_sha2\_password Auth failed

解决方案:

<https://github.com/alibaba/canal/issues/3902>

<<https://github.com/alibaba/canal/issues/3902>>

ALTER USER 'canal'@'%' IDENTIFIED WITH mysql\_native\_password BY 'canal';

ALTER USER 'canal'@'%' IDENTIFIED BY 'canal' PASSWORD EXPIRE NEVER;

FLUSH PRIVILEGES;

## 配置 kibana 可视化看板

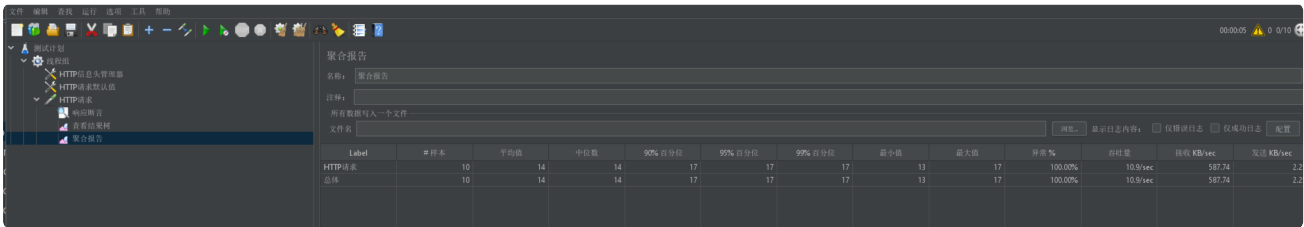
1. 创建索引
2. 导入数据
3. 创建索引模式
4. 选择图表、拖拉拽
5. 保存

# 压力测试

官方文档: <https://jmeter.apache.org/> <<https://jmeter.apache.org/>>

找到 jar 包: apache-jmeter-5.5\apache-jmeter-5.5\bin\ApacheJMeter.jar 启动

配置线程组 => 请求头 => 默认请求 => 单个请求 => 响应断言 => 聚合报告 / 结果树



| Label  | 采样本 | 平均值 | 中位数 | 90% 百分位 | 95% 百分位 | 99% 百分位 | 最小值 | 最大值 | 异常%     | 吞吐量      | 接收 KB/sec | 发送 KB/sec |
|--------|-----|-----|-----|---------|---------|---------|-----|-----|---------|----------|-----------|-----------|
| HTTP请求 | 10  | 14  | 14  | 14      | 17      | 17      | 13  | 17  | 100.00% | 10.0/sec | 587.74    | 2.2       |
| 总计     | 10  | 14  | 14  | 14      | 17      | 17      | 13  | 17  | 100.00% | 10.0/sec | 587.74    | 2.2       |

99%分位: 99%的用户都在这个响应时间内

吞吐量: 每秒处理的请求数 qps

## 更多学习

插件: <https://jmeter-plugins.org/install/Install/> <<https://jmeter-plugins.org/install/Install/>>

下载后文件为plugins-manager.jar <<https://jmeter-plugins.org/get/>> 格式, 将其放入jmeter安装目录下的lib/ext目录, 然后重启jmeter, 即可。

参考文章: [https://blog.csdn.net/weixin\\_45189665/article/details/125278218](https://blog.csdn.net/weixin_45189665/article/details/125278218)  
<[https://blog.csdn.net/weixin\\_45189665/article/details/125278218](https://blog.csdn.net/weixin_45189665/article/details/125278218)>

## 搜索建议

参考官方文档: <https://www.elastic.co/guide/en/elasticsearch/reference/7.17/search-suggesters.html> <<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/search-suggesters.html>>

示例:

```

1  GET post_v1/_search
2  {
3    "query": {
4      "match": { "content": "鱼皮不欠费" }
5    },
6    "highlight": {
7      "fields": {
8        "content": {
9          "pre_tags" : ["<h1>"], "post_tags" : ["</h1>"]
10        }
11      }
12    },
13    "suggest" : {
14      "my-suggestion" : {
15        "text" : "鱼皮不欠费",
16        "term" : {
17          "field" : "content"
18        }
19      }
20    }
21  }

```

## 搜索高亮

参考官方文档:

<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/highlighting.html>

<<https://www.elastic.co/guide/en/elasticsearch/reference/7.17/highlighting.html>>

示例:

```

1  GET post_v1/_search
2  {
3    "query": {
4      "match": { "content": "鱼皮欠费" }
5    },
6    "highlight": {
7      "fields": {
8        "content": {
9          "pre_tags" : ["<h1>"], "post_tags" : ["</h1>"]
10        }
11      }
12    }
13  }

```

高亮、建议都可以从返回值里拿到:

```
▼ searchHits = {SearchHitsImpl@12725} "SearchHits{totalHits=8, totalHitsRelation=EQUAL_TO, maxScore=3.089748, scrollId='null', sea
  totalHits = 8
  totalHitsRelation = {TotalHitsRelation@12541} "EQUAL_TO"
  maxScore = 3.089748
  scrollId = null
  searchHits = {ArrayList@12730} size = 8
    ▼ 0 = {SearchHit@12733} "SearchHit{id='1632381959341420546', score=3.089748, sortValues=[], content=PostEsDTO(id=16323
      index = "post_v1"
      id = "1632381959341420546"
      score = 3.089748
      sortValues = {Arrays$ArrayList@12745} size = 0
      content = {PostEsDTO@12746} "PostEsDTO(id=1632381959341420546, title=鱼皮是狗, content="鱼皮欠费"这句话似乎很奇怪
      highlightFields = {LinkedHashMap@12747} size = 0
      innerHits = {LinkedHashMap@12748} size = 0
      nestedMetaData = null
      routing = null
      explanation = null
      matchedQueries = {ArrayList@12749} size = 0
    > 1 = {SearchHit@12734} "SearchHit{id='1632381965322498049', score=2.8739307, sortValues=[], content=PostEsDTO(id=1632
    > 2 = {SearchHit@12735} "SearchHit{id='1632381970070450178', score=2.8739307, sortValues=[], content=PostEsDTO(id=1632
    > 3 = {SearchHit@12736} "SearchHit{id='1632381971433598978', score=2.8739307, sortValues=[], content=PostEsDTO(id=1632
    > 4 = {SearchHit@12737} "SearchHit{id='1632381972733833218', score=2.8739307, sortValues=[], content=PostEsDTO(id=1632
    > 5 = {SearchHit@12738} "SearchHit{id='1632366212904988673', score=2.8392506, sortValues=[], content=PostEsDTO(id=1632
    > 6 = {SearchHit@12739} "SearchHit{id='1634899414345404423', score=2.6907516, sortValues=[], content=PostEsDTO(id=1634
    > 7 = {SearchHit@12740} "SearchHit{id='1634900475080732678', score=2.6907516, sortValues=[], content=PostEsDTO(id=1634
  > unmodifiableSearchHits = {Lazy@12731}
  aggregations = null
  suggest = null
```

## 前端防抖节流（自行扩展）

问题：用户频繁输入、频繁点搜索按钮怎么办？

解决：使用 lodash 工具库实现防抖和节流。

节流：每段时间最多执行 x 次（比如服务器限流）

<https://www.lodashjs.com/docs/lodash.throttle>

<<https://www.lodashjs.com/docs/lodash.throttle>>

防抖：等待一段时间内没有其他操作了，才执行操作（比如输入搜索）

<https://www.lodashjs.com/docs/lodash.debounce>

<<https://www.lodashjs.com/docs/lodash.debounce>>

## 接口稳定性优化（自行扩展）

问题：调用第三方接口不稳定怎么办？（比如 bing 接口）

使用 guava-retrying 库实现自动重试：<https://github.com/rholder/guava-retrying>

<<https://github.com/rholder/guava-retrying>>

可以阅读鱼皮原创的这篇文章学习：<https://cloud.tencent.com/developer/article/1752086>  
<<https://cloud.tencent.com/developer/article/1752086>>

url=https%3A%2F%2Fwww.yuque.com%2Fu37765561%2Fak85bt%2F0e0675abc19dfb44c19789e5