



文档中心

ES服务 Elasticsearch

产品文档

目录

目录	2
概览	10
产品简介	12
产品特性	12
产品使用场景	13
创建集群	14
创建	14
集群类型选择	17
节点规格、节点数量选择以及付费方式	20
设置Kibana账号密码	23
确认支付	24
集群部署	24
集群信息	25
基本信息	25
VIP信息	26
节点配置	26
付费信息	27
集群健康	28
节点列表	30

集群操作	32
集群的启、关、删	32
集群改配	35
集群配置管理	38
集群线性重启	41
节点日志查询	44
kibana密码重置	45
上传同义词文件	47
注意事项	47
操作步骤	47
后续步骤	49
Kibana	51
Kibana代理访问	54
Kibana公网访问	55
监报告警	57
集群、集群节点监控	57
告警配置	59
数据备份	64
名词释义	64
准备工作	64
操作步骤	67

参考信息	73
数据恢复	74
准备工作	74
操作步骤	75
参考信息	79
预置插件	80
Analysis Plugins	80
Ingest Plugins	81
Elasticsearch Head	81
SQL	81
Snapshot/Restore Repository Plugins	84
插件管理	85
1.安装插件	87
2.卸载插件	90
Security插件	94
一、特别提示	94
二、插件安装	94
三、注意事项	94
四、权限管理	95
US3插件	97

一、安装插件	97
二、创建US3存储空间	97
三、创建仓库	100
四、使用方法	103
五、常用命令	103
ES测试	105
集群状态测试	105
索引一个文档	106
检索文档	107
搜索文档	107
复杂搜索	108
更新文档	108
删除	109
配置管理	110
插件管理	118
Elasticsearch-Head	118
Hadoop HDFS Repository Plugin	118
IK Analysis Plugin	122
实例运维	126
集群健康	126
单个节点监控	127

集群统计	137
索引统计	138
等待中的任务	139
cat API	139
动态变更设置	141
日志记录	142
历史数据清理	144
数据冷热分离	147
实例迁移	150
创建索引	150
数据迁移	150
故障恢复	155
yellow状态分析恢复	155
Logstash部署	158
安装Logstash	158
配置文件编写	158
功能文档	160
ES开发指南	160
ES API文档	160
Logstash	161

产品简介	162
产品特性	162
产品使用场景	162
创建Logstash实例	163
创建	163
节点规格、节点数量选择以及付费方式	165
确认支付	166
实例部署	166
集群信息	168
基本信息	168
节点配置	169
付费信息	169
节点列表	170
释放实例	171
注意事项	171
操作步骤	171
实例改配	173
注意事项	173
操作步骤	173
管道管理	177

创建管道	177
修改管道	180
删除管道	180
YML 文件配置	182
操作步骤	182
版本管理	183
版本特性	183
FAQs	185
开发时连接集群失败的原因？	185
UES支持跨地域、跨机房访问吗？	185
能不能从相同可用区的云主机通过SSH登录访问集群节点？	185
外网怎么访问UElasticsearch集群？	185
ULB产品支持UElasticsearch集群负载均衡访问吗？	186
UES内置的head插件为何连接集群失败？	186
什么时候选用主节点分离类型的集群？	186
集群是否支持扩容？节点是否支持修改配置？主节点是否支持扩容？	186
SSD磁盘的几种类型节点有什么区别？	186
索引主分片数量和副本设置几个最合理？	187
UES是否支持工具扩展，如Logstash的集成？	187
Rally压测	188

UES价格	203
1. 按配置计费	203
2. 按节点计费	204

概览

- 产品简介
- 操作指南
 - Elasticsearch
 - 创建集群
 - 集群信息
 - 集群操作
 - 同义词操作
 - Kibana
 - 监控告警
 - 数据备份与恢复
 - 数据备份
 - 数据恢复
 - 插件指南
 - 预置插件
 - 插件管理
 - Security插件
 - US3插件
 - 开发指南
 - ES测试
 - 配置管理
 - 插件管理
 - 实例运维
 - 实例迁移
 - 故障恢复
 - Logstash部署
 - 功能文档
 - Logstash
 - 产品简介

- 实例管理
 - 创建实例
 - 实例信息
 - 释放实例
- 实例变更
- 管道管理
- YML文件配置
- 版本管理
- FAQs
- Rally压测
- UES价格

产品简介

UES基于开源Elasticsearch构建,是一个分布式多用户能力的全文搜索引擎服务。UES通过创建集群的方式来提供服务,能够快速实现集群的部署,为客户提供对存储的海量数据进行实时检索和分析的能力。

产品特性

* 海量数据存储、检索、分析

对海量数据进行近实时的处理,全文检索,结构化检索,数据分析。

* 快速部署、简单易用

几分钟实现服务集群的部署、初始化,服务内置了常用的分词、文档解析和SQL插件,还支持集群线性扩容等操作。

* 性能优越、安全可靠

底层节点使用高性能的SSD磁盘,集群的运行对数据备份、数据迁移、故障恢复有最基本高效的支持。服务集群运行在私网中,保障数据的写入和检索安全。

* 性能监控、提供可视化管理

服务支持丰富的性能指标监控和内置了含有权限验证的可视化管理平台Kibana,以帮助客户更好地管理和维护自己的服务集群。

* 丰富节点配置、计费灵活

支持按年、按月、按小时计费,可依据业务需求选择合适的节点配置和计费周期,有效控制成本。

产品使用场景

* 日志管理分析

分析网址、移动设备、服务器等非结构化和半结构化日志,用于错误排查、应用程序监控、欺诈检测等分析场景

* 全文搜索

电商、O2O、企业等行业的搜索与导航服务

* 文档存储

对Json的文档的存储分析做出友好支持

* 数据存储（数据库功能）

开始一个新项目,考虑使用ES作为唯一的数据存储,以帮助保持设计尽可能简单。此场景不支持包含频繁更新、事务(transaction)的操作。

* 传统数据库数据分析

对一个在运行的复杂的系统,需求在现有系统中添加检索服务。一种相对安全的方式是就是将ES作为新的组件添加到现有系统中。

创建集群

控制台上 产品与服务 寻找 大数据 大类下 **ES服务 Elasticsearch** 进入产品页

创建

< Elasticsearch服务 / 创建集群

地域

地域与可用区

 华北一  可用区B 

基础配置

服务版本

v7.10 

 分片数默认是1，建议创建索引的时候根据实际情况增加分片数

节点告警模板

节点将绑定默认告警模板，若要调整监控项，请稍后在集群概况中调整告警配置

默认 UES 节点告警模板

节点配置

数据节点

磁盘类型

RSSD云盘型

计算规格

CPU内存比

1:21:4

o.es2m.medium

CPU2核

内存4G

o.es2m.xlarge

CPU4核

内存8G

o.es2m.2xlarge

CPU8核

内存16G

o.es2m.4xlarge

CPU16核

内存32G

o.es2m.8xlarge

CPU32核

内存64G

磁盘大小

100G

节点数量

3个

剩余节点配额:90

Kibana节点

磁盘类型	计算规格	磁盘大小	节点数量	
RSSD云盘	o.es2m.medium 2C 4G	20GB	1	
+ 添加Kibana节点				

专有主节点

磁盘类型	计算规格	磁盘大小	节点数量
+ 添加专有主节点			

协调节点

磁盘类型	计算规格	磁盘大小	节点数量
+ 添加协调节点			

网络设置

所属VPC

DefaultVPC

所属子网

DefaultNetwork()

更多设置

集群名称*

业务组

未分组

Kibana账号

elastic

Kibana密码*

设置密码

随机生成

页面会列出可创建服务的全部地域和可用区,选择需求的可用区、服务版本、磁盘类型、节点类型、磁盘大小、节点数量、填写一些初始化信息等内容后创建就可实现服务集群的快速部署。

集群类型选择

基础类型

节点配置

数据节点

磁盘类型

RSSD云盘型

计算规格

CPU内存比

1:21:4

o.es2m.medium

CPU2核

内存4G

o.es2m.xlarge

CPU4核

内存8G

o.es2m.2xlarge

CPU8核

内存16G

o.es2m.4xlarge

CPU16核

内存32G

o.es2m.8xlarge

CPU32核

内存64G

磁盘大小

100G

节点数量

3个

剩余节点配额:90

Kibana节点

磁盘类型	计算规格	磁盘大小	节点数量	
RSSD云盘	o.es2m.medium 2C 4G	20GB	1	
+ 添加Kibana节点				

专有主节点

专有主节点

磁盘类型	计算规格	磁盘大小	节点数量
------	------	------	------

+ 添加专有主节点

协调节点

协调节点

磁盘类型	计算规格	磁盘大小	节点数量
------	------	------	------

+ 添加协调节点

初始创建可以选择集群类型：

* 基础类型为集群没有单独master节点的情况,即所有节点配置为：

```
`node.master: true`
```

```
`node.data: true`
```

* 专有主节点类型为集群设置有单独master节点的情况

master节点配置为：

```
`node.master: true`
```

```
`node.data: false`
```

data节点配置为：

```
`node.master: false`
```

```
`node.data: true`
```

* 协调节点类型为集群设置有单独协调节点的情况

协调节点配置为：

```
`node.master: false`
```

```
`node.data: false`
```

如果您的集群是需求较多的节点数和较高的节点配置,推荐考虑选用主节点分离类型。这样即保证主节点选用较低的配置来避免资源浪费,又保证数据节点选用较高配置来提高数据处理能力。其中主节点分离类型主节点个数为3个,节点个数5-20个时选用的主节点配置为2核8G,节点个数大于20个时选用的主节点配置为4核8G。

值得注意的是,基础类型暂不支持添加主节点。

节点规格、节点数量选择以及付费方式

节点配置

数据节点

磁盘类型

RSSD云盘型

SSD云盘

计算规格

CPU内存比

1:2

1:4

o.es2m.medium



- CPU 2核
- 内存 4G

o.es2m.xlarge



- CPU 4核
- 内存 8G

o.es2m.2xlarge



- CPU 8核
- 内存 16G

o.es2m.4xlarge



- CPU 16核
- 内存 32G

o.es2m.8xlarge



- CPU 32核
- 内存 64G

磁盘大小

100 G

节点数量

3



剩余节点配额:83

其中可选的节点规格

按磁盘、规格分类：

磁盘类型	标识
RSSD云盘	o.es
SSD云盘	n.es

可选节点规格详细配置(最新节点参考“UES价格”导航页)：

机型	名称	配置
内存密集型	o.es4m.medium	2核 8G
内存密集型	o.es4m.xlarge	4核 16G
内存密集型	o.es4m.2xlarge	8核 32G
内存密集型	o.es4m.4xlarge	16核 64G
普通机型	o.es2m.medium	2核 4G
普通机型	o.es2m.xlarge	4核 8G
普通机型	o.es2m.2xlarge	8核 16G
普通机型	o.es2m.4xlarge	16核 32G
普通机型	o.es2m.8xlarge	32核 64G

计算规格

选择计算规格前,您需要先注意到UES节点的磁盘空间,然后关注集群主要用于存储哪些数据、数据增量等信息,选择合适的磁盘大小对应的节点规格。

磁盘大小

Elasticsearch 服务存储容量的主要因素如下：

- 副本数量:副本有利于增加数据的可靠性,但同时会增加存储成本。默认和建议的副本数量为1,对于部分可以承受异常情况导致数据丢失的场景,可考虑设置副本数量为0。
- 数据膨胀:除原始数据外,ES 需要存储索引、列存数据等,在应用编码压缩等技术后,一般膨胀10%。
- 内部开销:ES 占用约20%的磁盘空间,用于 segment 合并、ES Translog、日志等。
- 操作系统预留:默认操作系统会保留5%,用于关键流程处理、系统恢复、防止磁盘碎片化问题等。
- 安全阈值:通常至少预留15%的安全阈值。

因此,数据在 Elasticsearch 中占用的实际空间可通过下面公式估算:

$$\begin{aligned}\text{磁盘总大小} &= \text{源数据} \times (1 + \text{副本数量}) \times (1 + \text{数据膨胀}) / (1 - \text{内部任务开销}) / (1 - \text{操作系统预留}) / (1 - \text{安全阈值}) \\ &= \text{源数据} \times (1 + \text{副本数量}) \times 1.7\end{aligned}$$

节点数量

节点数量选择与索引分片、数据冗余、容错转移、故障恢复、处理效率等有密切的关系。如果您对Elasticsearch已经有了深刻的领悟和熟练的使用,想必可以做出合理的节点配置和个数选择;如果您对Elasticsearch还处于了解阶段,可以根据实际的业务需求、处理数据的量级合理分析选择。

为了保证集群的高可用性,基础类型集群默认最少3个节点,主节点分离类型集群默认主节点为3个,最少数据节点为2个。

创建成功,后续集群可以扩容节点(不支持主节点创建和扩容)。

计费选择

计费模式分为按时、月付、年付3种类型。控制台支持计费类型修改。

设置Kibana账号密码

Kibana账号

Kibana密码 *



随机生成

Kibana是一个开源的分析与可视化平台,是Elasticsearch比较常用的管理工具。您可以用kibana搜索、查看、交互集群中的索引数据,使用各种不同的图表、表格、地图等kibana能够很轻易地展示高级数据分析与可视化。

Kibana让我们理解大量数据变得很容易。

产品默认开启kibana服务,并添加了权限验证。需要在创建集群的时候设置账号密码。

确认支付

集群信息选择和填写完成后点击创建,到支付确认页面,完成支付创建。

集群部署



The screenshot shows the 'Elasticsearch服务 UES' management console. At the top, there's a navigation bar with '全部产品', 'UES', '上海二', and '可用区B'. Below this, the main content area has a header 'Elasticsearch服务 UES'. On the left, there are buttons: '创建集群', '启动', '关闭', and '重启'. On the right, there's a search bar and icons for '消息', '告警', and '帮助与支持'. The main area displays a table of clusters. The table has columns: '集群名称', '资源ID', '服务版本', '可用区', '业务组', '运行时长', '创建时间', '状态', and '操作'. There is one cluster listed with a red status indicator and the text '运行'. The '操作' column for this cluster contains buttons: '详情', '打开Kibana', and '...'. At the bottom right, there's a pagination bar showing '1' and '10 条/页'.

集群名称	资源ID	服务版本	可用区	业务组	运行时长	创建时间	状态	操作
修改集群名称		7.10	上海二可用区B	Default			运行	详情 打开Kibana ...

等待部署中,由于集群规模不同,所需要的部署时间会有所差异,部署大致时间在2分钟到5分钟。

部署中,列表的集群的集群状态会显示为"蓝色"状态创建中,创建成功后集群状态会动态变化为"绿色"状态运行。

如果长时间为创建状态,可点击刷新重新获取集群状态。

集群信息

控制台会给出集群的基本信息展示。点击详情进入详情页，我们会提供集群概况、节点信息和集群监控视图等信息。

基本信息

基本信息		
资源ID		
资源名		
服务版本	7.4.2	
可用区	北京二可用区B	
业务组	Default	
运行时长		

运行状态	 运行
VIP	
创建时间	
节点个数	3

其中资源名 (和列表中 "集群名称" 完全相同) 可以自定义, 如上图标注。请注意此资源名不是Elasticsearch配置文件中的集群名称 (配置文件中的 cluster.name 默认为 ** 资源ID **)。

VIP信息

VIP是在用户的 VPC 下提供了一个访问集群的 ULB, 通过负载均衡的方式挂载了集群内的所有节点。这样做的目的, 一方面是为了适配集群的弹性伸缩并保证其高可用性, 当集群的节点发生变化时, VIP 会自动更新节点信息; 另一方面, 简化用户的操作, 用户不用再关注集群节点IP等信息的变更。目前只使用于**RSSD**类型集群。

节点配置

节点配置

类型	STES-large
CPU	2核
内存	8G
磁盘类型	SATA
磁盘大小	500G
节点类型	SATA盘型

集群节点配置详情。

如果集群是主节点分离类型，标题为数据节点配置，主节点配置不在此栏中显示。

付费信息

付费信息

续费

创建时间	
到期时间	
付费方式	按时付

集群付费详情。续费可以到计费系统进行续费、更改计费方式等操作。

集群健康

集群信息

集群状态

节点数量	3
------	---

数据节点数量	3
--------	---

活跃主分片数量	1
---------	---

活跃分片数量	2
--------	---

正在迁移分片数量	0
----------	---

正在初始化分片数量	0
-----------	---

未分配分片数量	0
---------	---

延迟未分配分片数量	0
-----------	---

挂起的任务数量	0
---------	---

正在传送抓取数量	0
----------	---

任务最大排队等待	0
----------	---



该部分就是API `/\_cluster/health` 的展示信息,可以提供简单的集群状态。服务关闭状态,该部分显示空内容。

该部分展示了最基础的集群状态、分片数量、正在迁移分片数、未分配分片数等重要集群信息。

节点列表

节点列表

节点ID	节点IP	类型	节点配置配置	容量: 单位G	节点类型	操作
		STES-large	2核 8G	500	mdi	节点日志 重启 ...
		STES-large	2核 8G	500	mdi	节点日志 重启 ...
		STES-large	2核 8G	500	mdi	节点日志 重启 ...

10条/页 1 /0

展示集群节点IP信息、节点配置信息及节点可用操作。

节点IP为业务网地址,ues默认是通过内网访问的,为了数据安全性,强烈不建议代理到外网环境访问。

节点类型分为mi、di、mdi 3种。

mi和di为主节点分离类型集群显示,mi为主节点资格节点,di为数据节点。

mdi为基础类型集群显示,mdi为节点具备主节点资格并可以用来保存数据。

具体查看哪一个节点为主节点,可自行通过API ``/_cat/master?v`` 获取。

集群操作

集群相关操作在列表页和详情页都有入口,操作的对象为集群和集群节点。

集群的启、关、删

🔗 日志分析

集群管理

创建集群 启动 关闭 重启

搜索 设置 刷新 帮助

☐	集群名称	资源ID	可用区	运行时长	创建时间 1h	状态	操作
☐	修改集群名称		广州可用区B			● 运行	详情 打开Kibana ...

启动

关闭

重启

扩容集群

删除集群

重置Kibana密码

1 / 1

以关闭服务为例说明

关闭服务



是否关闭以下集群服务？

Elastic

● 运行

取消

确定

关闭服务



是否关闭以下集群服务？

Elastic

● 关闭中

关闭

关闭服务



是否关闭以下集群服务？

Elastic

● 关闭成功

关闭

启动、关闭、重启操作服务状态会动态改变，直至操作结束。在集群列表这种动态变化也会展现。

删除操作该集群将从列表中除去。

目前不支持批量操作，现有批量选择只会默认选中第一项进行操作。

以上操作请谨慎操作，对已在生产使用的集群可能会造成重大影响。此操作对新建或无数据的集群比较友善

集群改配

增加节点

[<](#) 编辑节点

当前节点信息

集群名称		资源ID		节点个数	3个 (2核 8G)
磁盘信息	CLOUD_RSSD 200G				

编辑节点

磁盘类型

增加节点

更改规格

修改磁盘大小

节点数量

3

↑

↓

剩余节点配额:

其中节点配置为当前节点信息。

需要注意的是目前只支持扩容,即添加集群节点,暂不支持缩减节点。

为扩容后集群节点个数,所以限制最小值为原集群节点个数。其中新增节点个数不能大于节点配额,我们分配的节点配额足够支撑你的集群使用,如果配额需求过大,请联系客户经理进行节点资源配额调整。

计费类型为初始创建集群所选计费类型,如有续费、变更计费方式,请在控制台用户中心进行查询修改。

更改规格

编辑节点

磁盘类型

升级规格期间将逐次重启所有数据节点，建议在业务低峰期操作，规格不支持降级。

[增加节点](#)[更改规格](#)[修改磁盘大小](#)

计算规格

CPU内存比

[1:2](#)[1:4](#)**o.es2m.2xlarge**

- CPU 8核
- 内存 16G

o.es2m.4xlarge

- CPU 16核
- 内存 32G

o.es2m.8xlarge

- CPU 32核
- 内存 64G

其中改配节点为当前数据节点，不包括专有主节点。

更改规格的前提是确保集群的状态为正常(绿色)、索引至少包含1个副本；并建议在业务低峰期操作。

在做节点改配操作时，会检查集群的健康状态。如果状态为正常(绿色)时，进行该节点改配、重启节点操作；如果状态不为正常(绿色)时，会持续等待检查集群状态，直至正常(绿色)。假如时间过长、或出现改配失败，请联系技术支持。

需要注意的是目前只支持升配，即节点CPU、内存升级，暂不支持减配。

修改磁盘大小

[<](#) 编辑节点

当前节点信息

集群名称

资源ID

节点个数

3个 (2核 | 8G)

磁盘信息

CLOUD_RSSD | 200G

编辑节点

磁盘类型

增加节点

更改规格

修改磁盘大小

磁盘大小

200 G

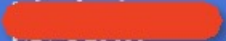
其中修改磁盘大小的节点为当前数据节点, 不包括专有主节点。

修改磁盘步长为100G, 最大磁盘是8000G。

需要注意的是修改磁盘不支持缩容。

集群配置管理

详情页集群概况包含 参数管理 操作

 产品与服务 广州 ▾

可用区B ▾



集群管理 / 集群详情

集群概况

节点信息

监控视图

打开Kibana

参数管理

启动

更多 ▾

配置管理



! 修改参数配置后需重启服务生效。其他常用配置可通过API动态设置

参数名称	参数值	使用范围	说明
discovery.zen.minimum_master_nodes	*	cluster	最小主节点资格数，可动态设置
gateway.recover_after_nodes	2	cluster	预期的节点数量达到，就可以进行恢复
gateway.expected_nodes	3	cluster	预计集群节点数量。当预期的节点数量加入到群集，分片的恢复就会开始

取消

确定

打开参数管理如上图所示。其中列出参数名称、参数值、应用范围、参数说明等项。另外提供参数名称搜索快速检索修改。

值得注意的是，其中的很多参数，我们已经根据所选集群节点配置做出了合理的初始化。除非是自己业务需求必须修改某些参数，其他情况不建议修改，另外修改后的参数值一定要符合规范，否则可能会对集群产生不友好的影响。修改参数后需重启集群服务以生效。

对于未被列出的参数修改，大多数是可以通过API动态修改的，我们也提倡尽可能使用API来动态修改配置参数。

集群线性重启

假设生产环境某种情况下需要重启集群，但是要求不能影响服务的正常使用。这时候集群的线性重启就显得尤为重要。

节点列表

节点ID	节点IP	类型	节点配置配置	容量: 单位G	节点类型	操作
		STES-large	2核 8G	500	mdi	节点日志 重启 ...
		STES-large	2核 8G	500	mdi	节点日志 重启
		STES-large	2核 8G	500	mdi	节点日志 重启 ...

10条/页 1 /0

详情节点列表中提供了单节点重启的操作，如上图。进行重启单节点操作

重启服务



是否重启以下节点服务?

取消

确定

重启服务



是否重启以下节点服务？



! 操作成功，集群重启操作要求集群状态恢复后再进行下一个节点操作

关闭

重启成功会有上图提示。

值得注意的是,单节点重启成功后，集群需要一定时间恢复到健康状态，之后再进行下一个节点的操作

检测集群是否恢复健康的方法：

* 集群概况中刷新集群信息模块,一定要刷新来重新获取集群状态

集群信息

集群状态	
节点数量	3
数据节点数量	3

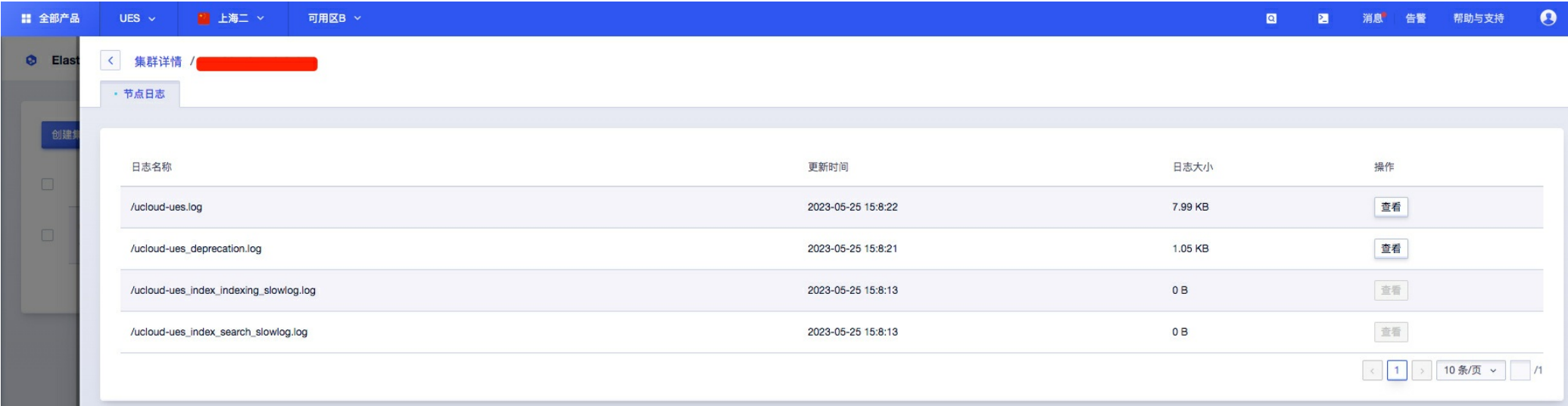


刷新一次后,如果状态变为黄色,可停止刷新,等待集群自动恢复到绿色健康状态;如果状态为绿色,说明集群已恢复健康状态;如果状态显示消失,请持续间隔刷新,直到集群状态显示。当集群为绿色状态时,可进行下一个节点操作。

* 通过API ``/_cluster/health`` 获取实时集群状态,直到集群状态恢复健康状态。

节点日志查询

详情页节点列表包含 节点日志 操作



点击查看操作, 查看对应日志。

kibana密码重置

重置Kibana密码

Kibana密码 *

.....

确认密码 *

.....

· 密码长度限8-30个字符
· 必须同时包含大写字母, 小写字母和数字。无特殊字符

取消

确定

重置Kibana密码

Kibana密码 *

确认密码 *

操作成功

关闭

忘记kibana密码可以在控制台来重置密码,kibana账号目前不支持修改。如果忘记账号,请联系客户经理找回。

上传同义词文件

在通过同义词文件方式使用同义词时,您需要先上传同义词文件。本文介绍上传同义词文件的注意事项和操作步骤。

注意事项

- 1.单个字典文件最大支持4MB,最多支持5个字典文件,文件名支持大小写字母、数字、下划线,长度不要超过30个字符
- 2.一次支持单个文件上传
- 3.同义词文件要求每行一个同义词表达式(表达式支持 Solr 规则 和 WordNet 规则),并且文件需要为utf-8编码,扩展名为.txt
- 4.上传、更新同义词词典之后,需要执行重启集群操作生效

操作步骤

1. 在控制台上进入Elasticsearch目标实例
2. 进入同义词管理页面,单击上传文件,选择上传同义词文件并确认。

[同义词管理页](#)

<

Elasticsearch服务 / 集群详情

集群概况

节点信息

监控视图

插件管理

数据备份

同义词管理

1

1.单个字典文件最大支持4MB，最多支持5个字典文件，文件名支持大小写字母、数字、下划线，长度不要超过30个字符

2

2.一次支持单个文件上传

3

3.同义词文件要求每行一个同义词表达式（表达式支持 Solr 规则 和 WordNet 规则），并且文件需要为utf-8编码，扩展名为.txt

4

4.上传、更新同义词词典之后，需要执行重启集群操作生效

详情请参考文档 [\(🔗\)](#)

上传文件

文件名称📄	状态	更新时间🕒	操作
1 数据为空			

单击上传文件

synonyms

搜索

名称	大小	种类	添加日期
 ucloud_synonyms.txt	238 字节	纯文本文稿	今天 上午11:24

选项

取消

打开

后续步骤

等待实例的状态变为正常后,登录Kibana控制台创建索引、校验同义词,并上传测试数据进行搜索测试。创建索引时需要配置 settings和 mapping,并且需要在 settings中配置 "synonyms_path": "analysis/your_dict_name.txt"。

Kibana

Kibana是一个开源的分析与可视化平台,设计出来用于和Elasticsearch一起使用的。你可以用kibana搜索、查看、交互存放在Elasticsearch索引里的数据,使用各种不同的图表、表格、地图等kibana能够很轻易地展示高级数据分析与可视化。

Kibana让我们理解大量数据变得很容易。

产品默认打开kibana服务,控制台产品页上可以点击`打开Kibana` 打开kibana服务窗口。集群的kibana地址已经做了域名代理访问,服务地址不唯一,需要使用时通过控制台来打开。另外我们对kibana做了权限验证,所以不必担心开放的kibana地址会被恶意登陆操作,也请妥善保存您的账号,可以不定期修改密码确保安全。

需要进行身份验证

http://

您与此网站的连接不是私密连接

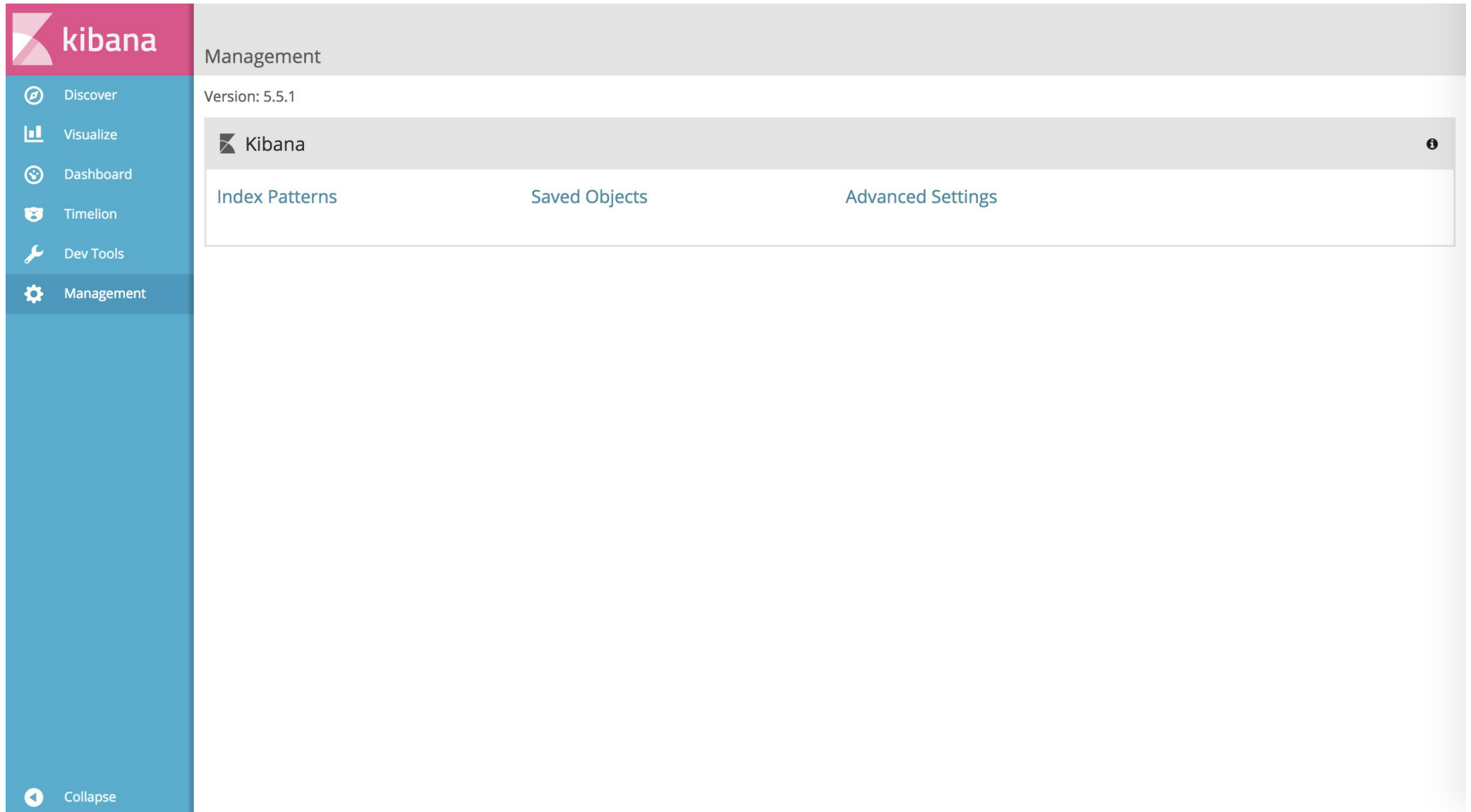
用户名

密码

取消

登录

kibana打开了,你可以使用它了。



官网Kibana

Kibana代理访问

可以通过在控制台开启或关闭Kibana代理访问, 设置Kibana在公网访问是否能访问。(默认为开启状态)

开启Kibana代理访问

☐	集群名称	资源ID	服务版本	可用区	业务组	运行时长	创建时间 ↕	状态	操作
☐	修改集群名称			上海二可用区A	Default			● 运行	详情 打开Kibana ... 启动 关闭 重启 扩容集群 升级节点配置 删除集群 重置Kibana密码 自定义IK分词词库 Kibana服务操作 告警配置 插件管理 开启kibana代理访问

关闭Kibana代理访问

☐	集群名称	资源ID	服务版本	可用区	业务组	运行时长	创建时间	状态	操作
☐	修改集群名称			上海二可用区A	Default			● 运行	详情 打开Kibana ...

启动

关闭

重启

扩容集群

升级节点配置

删除集群

重置Kibana密码

自定义IK分词词库

Kibana服务操作

告警配置

插件管理

关闭Kibana代理访问

10 条/页

Kibana公网访问

应用型负载均衡ALB

通过ALB设置防火墙白名单策略,保证公网访问Kibana安全性。

新建集群

2024年8月1日之后,创建UES集群时,Kibana节点独立部署,不再提供公网Kibana代理访问功能。

若客户有公网访问Kibana需求,则可以通过创建应用型负载均衡ALB, 绑定Kibana节点的IP地址、端口为5601,实现Kibana公网访问。

原有集群

原有集群Kibana服务部署在数据节点上,由于公网访问安全的问题,建议客户关闭Kibana代理访问,通过创建应用型负载均衡ALB, 绑定第一个数据节点的IP地址、端口为5601,实现Kibana公网访问。

监控告警

集群、集群节点监控

UES开通对集群和节点相关性能的监控,在控制台集群详情中视图展示

< 集群管理 / 集群详情

集群概况

节点信息

监控视图

集群

节点

监控信息

请选择

1小时

6小时

12小时

1天

7天

自定义

(%)

☒ CPU使用率

可用内存百分比(%)

☒ 可用内存百分比

磁盘使用率(%)

☒ 磁盘使用率

磁盘读吞吐(Bps)

☒ 磁盘读吞吐

告警配置

控制台上打开 监控 **UMon** 产品面板页(可用区需选择和UES一样)。告警模板 选择 创建模板 (因为UES没有默认告警模板,需要创建自定义模板)

创建告警模板

类型:

云主机

从现有模板导入:

ES集群

ES集群节点

跨域通道

容器服务

容器节点

单机版memcache

主备版Redis

云内存存储

模板名称:

备注:

建

如上图,创建两个自定义告警模板,集群监控项告警(ES集群)和节点监控告警项(ES集群节点),也可根据需求只创建需要的模板。

创建成功后,可在列表中看到自定义的UES告警模板

监控	创建模板	删除	什么是告警模板?	全部	
资源监控	☐	模板名	类型	绑定资源数	操作
告警模板	☐	默认分布式消息系统节点告警模板 默认推荐模板	分布式消息系统节点	0	编辑规则 查看资源
告警记录	☐	默认云数据仓库告警模板 默认推荐模板	云数据仓库	0	编辑规则 查看资源
通知人管理	☐	默认云数据仓库节点告警模板 默认推荐模板	云数据仓库节点	0	编辑规则 查看资源
站点监控	☐	默认单机memcache告警模板 默认推荐模板	单机版memcache	0	编辑规则 查看资源
监控代理	☐	ES集群监控告警	ES集群	0	编辑规则 查看资源 重命名 删除
消息订阅	☐	ES集群节点监控告警	ES集群节点	0	编辑规则 查看资源 重命名 删除
网络质量监控					

以节点为例,编辑模版告警规则

← ES集群节点监报告警

添加规则

删除规则

保存

另存为

节点cpu利用率

已分配内存百分比

磁盘使用率

节点磁盘读吞吐量

点磁盘写吞吐量

网络入口

网络出口

节点磁盘读次数

节点磁盘写次数

节点网卡入包量

节点网卡出包量

对比方式	告警阈值	探测周期(s)	触发周期(次)	收敛策略	通知对象

其中, 可选添加规则和监控指标项相同。添加规则后, 可以设置每一个告警项的告警阈值、触发周期、收敛策略、通知对象等。

模版设置完成后, 可以在 资源监控 里, 将集群节点资源配置上告警模板。

监控

资源类型 ES集群节点

刷新

帮助

设置

资源列表

告警管理

全部

搜索

	名称	告警模板	操作
☐			
☐		无告警	数据视图 告警管理

告警管理

?

您正在为以下1个资源设置告警模板

无告警

×

取消

设置

告警管理

通过模板您可以进行更方便的告警设置。我们提供了一些默认告警模板，您可以直接应用到您的主机或其他云资源。您也可以自定义模板，创建或编辑模板请移步至[告警模板列表](#)

无告警

ES集群节点监控告警

告警模板详情

监控对象	告警阈值	通知对象
节点cpu利用率(%)	大于等于:95	默认组
已分配内存百分比(%)	大于等于:98	默认组
磁盘使用率(%)	大于等于:85	默认组

取消

确定

确定,资源添加告警成功!

Copyright © 2012-2021 UCloud 优刻得

63/206

数据备份

UES的数据备份功能,使用Elasticsearch(以下简称ES)的 snapshot API,通过elasticsearch-repository-ufire插件(以下简称US3插件)的支持,按照预先设定的备份计划,定时自动生成快照,保存在对象存储 US3当中,从而实现对UES索引数据的有效备份。

名词释义

快照 (snapshot) :一个快照是对于运行中的ES集群在快照生成时间点的一个备份,可以针对单个索引或者所有索引生成快照。快照采用增量机制生成,这意味着后续快照中仅包含之前快照中不存在的那部分数据。

仓库 (repository) :仓库是用于存放快照文件的存储空间,必须先创建一个仓库才能进行快照的备份操作。通过安装插件,ES可以支持使用不同类型的存储系统作为快照仓库。

对象存储 US3:对象存储 US3是UCloud提供的非结构化文件云存储的服务,通过安装US3插件,UES使用US3作为存放快照的仓库。

准备工作

1.创建US3存储空间

登录US3控制台,创建用于保存UES快照数据的存储空间。要求:(1)存储空间所在地域必须与您需要快照备份的UES集群所在地域保持一致。(2)空间类型请选择**"私有空间"**,以确保存储空间内的数据需要得到密钥授权才能访问。

创建存储空间

选择地域

北京

香港

广州

上海二

洛杉矶

新加坡

雅加达

拉各斯

圣保罗

迪拜

胡志明市

台北

孟买

华盛顿

法兰克福

首尔

空间类型 ?

私有空间

公开空间

业务组

未分组

存储空间域名 *?

ues-repository

.cn-bj.ufileos.com

数据分析 ?

数据湖分析USQL

未授权

取消

确定

2.创建US3令牌

在令牌管理页面,创建用于授权访问上述存储空间的令牌,合理设置有效时长,授权访问上述存储空间,并赋予对所有文件的上传、下载、删除、文件列表权限。

创建令牌

令牌名称 *

token_of_ues_repository

有效时长 *

1

年

授权存储空间 ⑦

选择存储空间

已选择: ues-repository

授权文件 ⑦

所有文件

设置前缀

令牌权限

☒ 上传 ☒ 下载 ☒ 删除 ☒ 文件列表 ☐ 图片处理

取消

确定

3.为需要备份的UES集群安装US3插件

登录UES控制台,在目标UES集群的插件管理页面进行US3插件的安装,具体操作方法可以参考文档插件管理章节。

[<](#) [集群管理](#) / 集群详情

集群概况

节点信息

监控视图

插件管理

数据备份

☐	analysis-ukrainian	default	Provides stemming for Ukrainian.	● 已安装
☐	elasticsearch-repository-ufile	default	该插件支持将Elasticsearch数据的快照备份至UFile，并可从已备份至UFile的快照中将数据恢复至Elasticsearch。	● 已安装
☐	ingest-attachment	default	使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。	● 已安装

操作步骤

数据备份功能入口：登录UES控制台，点击集群列表中目标集群的“详情”按钮，进入集群详情页面，切换至“数据备份”标签页。

1.注册仓库

在**“仓库管理”子页面中，点击“注册仓库”**按钮，在弹出的对话框中，按照页面提示完成相关信息的填写，点击确认，系统将为该集群创建一个存放快照的仓库。

注册仓库



❗ UES使用对象存储UFile作为保存快照的仓库，您需要先开通UFile服务并新建一个UFile存储空间，且此存储空间所在地域必须与您需要快照备份的UES集群所在地域保持一致。

UFile插件	● 已安装	
UFile存储空间 *	ues-repository	创建新存储空间请前往UFile页面
UFile令牌 *	token_of_ues_repository	创建新令牌请前往UFile令牌管理页面。
仓库名称 *	my_first_repository	
基础路径 ①	ues_backup	
只读	OFF 对于同一个仓库，只应该设置一个集群对其具有写入权限，其他集群应该设置为只读权限。	
快照压缩 ①	ON	
分块大小 ②	60mb	
快照生成最大速度 ①	40mb /秒	
快照恢复最大速度 ①	40mb /秒	
快照自动清理	ON 自动删除 10 天之前的快照	

取消

确认

参数	说明
US3存储空间	用于存放UES集群的快照。 注意:建立仓库之后不可以删除该仓库绑定的US3存储空间,否则已经生成的快照文件将会丢失,新的快照也将无法生成。

US3令牌	用于授权访问上述存储空间,必须具有该存储空间的上传、下载、删除、文件列表权限,缺一不可。 注意:建立仓库之后不可以删除该仓库绑定的US3令牌,也不可以取消任何一项必备权限,同时还必须确保该US3令牌未过有效期,否则快照生成、快照恢复等操作将会失败。
仓库名称	用于标识一个仓库,在一个UES集群内部不允许出现多个同名仓库。
基础路径	快照文件在存储空间中的路径,默认为ues_backup,生成的快照文件会存放在US3存储空间中的该目录下。
只读	控制UES集群对于该仓库的读写权限。 注意:在多个UES集群中注册指向同一个路径(US3存储空间 + 基础路径)的仓库时,只应该设置一个集群对其具有写入权限,其他集群应该设置为只读权限,以确保快照文件的完整性和一致性。
快照压缩	启用压缩后,将对快照中的索引映射和设置文件进行压缩,但数据文件不会被压缩。
分块大小	用于限制快照过程中单个文件分块的大小。
快照生成最大速度	用于限制每个节点生成快照的最大速度。
快照恢复最大速度	用于限制每个节点恢复快照的最大速度。
快照自动清理	如果开启快照自动清理,系统将会按照设定的时间长度,定期删除该仓库中过期的快照文件。

2. 设定备份规则

在**“备份策略”子页面中, 点击“新建规则”**按钮,在弹出的对话框中,按照页面提示完成相关信息的填写,点击确认,系统将为该集群创建一条自动生成快照的规则。

新建规则

规则名称 *

my_first_policy

快照名称 * ①

daily_backup

所在仓库 * ①

my_first_repository

执行计划 *

每天 2 时 0 分

备份全部索引 ①

OFF

需要备份的索引

index_2,index_3

多个索引名称之间使用英文逗号分隔

忽略不可用索引 ①

OFF

允许不完全索引 ①

OFF

包含全局状态 ①

ON

取消

确认

参数	说明
规则名称	用于标识一条备份规则。
快照名称	指定该备份规则所生成快照的名称,系统将在快照名称后添加唯一标识符以区分不同时间生成的快照。
所在仓库	选择用于存放快照的仓库,只能选择注册仓库时设置为具有写入权限的仓库。
执行计划	设定进行自动备份的频率和时间点,目前支持在每天的指定时间进行一次自动备份。

备份索引	默认备份全部索引,也可以关闭此开关,手动输入索引名称以备份指定的索引,多个索引名称之间使用英文逗号分隔,支持索引表达式。
忽略不可用索引	用于控制对快照备份期间不可用索引的处理方式。此选项开启时,如果遇到不可用索引,将会继续生成不包含这些索引的快照。此选项关闭时,如果遇到不可用索引,快照备份将会失败。
允许不完全索引	用于控制对快照备份期间主分片不可用的索引的处理方式。此选项开启时,如果遇到此类情况,将会允许使用主分片不可用的索引生成快照。此选项关闭时,如果遇到此类情况,快照备份将会失败。
包含全局状态	用于控制是否将集群的全局状态保存在快照中。

3.执行备份规则

备份规则的执行有两种方式:自动定时执行和手动执行。

(1) 自动定时执行

设置好备份规则后,系统会在规则设定的时间,按照规则中配置的参数,定时自动执行快照备份操作。

(2) 手动执行

在**“备份策略”子页面中, 点击一条规则右侧的“立即执行”**按钮,可以按照规则中配置的参数,立即执行一次快照备份操作。

4.管理生成的快照

在**“快照管理”**子页面中,可以查看或删除已经生成的快照。

(1) 查看快照详情

点击一个快照右侧的**“详情”**按钮,可以查看这个快照的详细信息,包括快照的版本、备份状态、包含的索引等。

快照详情

基本信息

快照名称	daily_backup_202003281600
快照版本	6.5.4
备份状态	成功
包含全局状态	是
快照开始时间	2020-03-28 16:57:22
快照结束时间	2020-03-28 16:57:23

成功索引(2)

index_3
index_2

失败索引(0)

没有备份失败索引

关闭

(2) 删除快照

点击一个快照右侧的“删除”按钮, 可以从仓库中删除这个快照。如果在注册仓库时, 开启了快照自动清理功能, 系统将会按照设定的时间长度, 定期删除该仓库中过期的快照文件。

索引表达式

索引表达式用于指定多个索引,常用的表达式类型包括:

(1) 英文逗号分隔的多个索引名称

例如:index_1,index_2,index_3 指定了3个索引。

(2) 带通配符 * 的索引名称

例如:index_* 可以表示所有以 index_ 开头的索引,如 index_1,index_2020.02.02 等。

(3) 使用通配符 * 时,利用 - 排除不需要的索引。

例如:index_*,-index_1 表示所有以 index_ 开头,但不包含 index_1 在内的索引。

参考信息

关于创建快照、从快照恢复数据、查询快照信息等操作的更多信息,请参考Elasticsearch官方文档:

Snapshot And Restore

数据恢复

如果使用UES的数据备份功能生成了集群的快照,则可以使用Elasticsearch (以下简称ES) 的 snapshot API, 将快照中的数据恢复到指定的UES集群中。(下文中描述的数据恢复方式, 对于通过其他途径生成的UES集群快照也适用。)

准备工作

1. 确认快照仓库能否正常访问

如果需要从使用UES的数据备份功能生成的快照中恢复数据, 必须确保以下几点:

- (1) 注册仓库时绑定的US3存储空间可以正常访问, 且存放在其中的快照文件完好。
- (2) 注册仓库时绑定的用于授权访问上述US3存储空间的令牌状态正常, 即自仓库创建之后该令牌未被删除过, 具有上传、下载、删除、文件列表等各项必备权限, 且当前该令牌处于有效期内。

2. 确认快照的版本与目标**UES**集群的版本是否兼容

由于不同版本UES集群的索引数据结构可能存在差异, 因此快照中的索引数据只能被恢复到能够正常识别和读取这些索引的UES集群中。关于快照版本和集群版本的兼容性, Elastic官方给出的基本规则是: 版本相同时可以正常恢复, 跨版本情况下低版本集群创建的快照可以被恢复到大版本号差异不超过 1 的高版本集群中, 即:

- (1) 在 6.x 版本集群中创建的快照可以被恢复到 7.x 版本的集群中。
- (2) 在 5.x 版本集群中创建的快照可以被恢复到 6.x 版本的集群中。

(3) 在 2.x 版本集群中创建的快照可以被恢复到 5.x 版本的集群中。

(4) 在 1.x 版本集群中创建的快照可以被恢复到 2.x 版本的集群中。

除此以外,无法将高版本集群创建的快照恢复到低版本集群中,也无法将低版本集群创建的快照恢复到跨 2 个或者更多的大版本号的高版本集群中。

使用UES的数据备份功能生成的快照的版本信息,可以通过UES控制台:数据备份 -> 快照管理 -> 快照详情 查看,或者接下来介绍的**“查看指定仓库中的所有快照”**命令查看。

操作步骤

下文中描述的操作步骤,均在UES集群的**Kibana**控制台中进行。功能入口:登录UES控制台,点击集群列表中目标集群右侧的“打开Kibana”按钮,进入**Kibana**控制台,点击左侧导航栏的**Dev Tools**,在 **Console** 框中输入并执行相应命令。

注意: 下列命令样例中尖括号 <> 里的内容,需要根据自身情况,替换成相应内容(去掉尖括号)。

1.查看当前集群的所有快照仓库

命令样例:

```
GET /_snapshot/_all
```

命令返回结果样例:

```
{
  "my_first_repository": {
    "type": "ufile",
    "settings": {
      "bucket": "ues-repository",
      "public_key": " ",
      "chunk_size": "60mb",
      "endpoint": "ues-repository.ufile.cn-north-04.ucloud.cn",
      "max_restore_bytes_per_sec": "40mb",
      "readonly": "false",
      "compress": "true",
      "base_path": "ues_backup/",
      "private_key": " ",
      "max_snapshot_bytes_per_sec": "40mb"
    }
  }
}
```

参数	说明
my_first_repository	仓库名称,将在后续的命令中用到。
type	仓库存储系统的类型,“ufile”表示使用的是对象存储US3。
bucket	仓库所在的US3存储空间名称。
base_path	快照在上述US3存储空间中的存放路径。

2.查看指定仓库中的所有快照

命令样例:

```
GET /_snapshot/<my_first_repository>/_all
```

命令返回结果样例:

```
{
  "snapshots" : [
    {
      "snapshot" : "daily_backup_202003281600",
      "uuid" : "W7PxpXNJRe6jUcWjYaDtPQ",
      "version_id" : 6050499,
      "version" : "6.5.4",
      "indices" : [
        "index_3",
        "index_2"
      ],
      "include_global_state" : true,
      "state" : "SUCCESS",
      "start_time" : "2020-03-28T08:57:22.327Z",
      "start_time_in_millis" : 1585385842327,
      "end_time" : "2020-03-28T08:57:23.085Z",
      "end_time_in_millis" : 1585385843085,
      "duration_in_millis" : 758,
      "failures" : [ ],
      "shards" : {
        "total" : 2,
        "failed" : 0,
        "successful" : 2
      }
    }
  ]
}
```

参数	说明
snapshot	快照名称,将在后续的命令中用到。
version	快照版本,可用于判断与本集群的兼容性。
indices	快照中包含的索引,恢复时可以指定只恢复其中的部分索引。

3.从指定快照中恢复索引数据

注意：如果待恢复的索引在目标集群中已经存在，需要确保该索引已经被关闭，且与快照中的同名索引具有相同的分片数，否则恢复操作无法执行。

命令样例：

```
POST /_snapshot/<my_first_repository>/<daily_backup_202003281600>/_restore
```

命令返回结果样例：

```
{
  "acknowledged" : true
}
```

由于该命令为异步命令,返回上述信息时,数据恢复任务将在后台执行。如果希望监测数据恢复进程,可以在上述命令后增加 **wait_for_completion** 参数,这将阻塞客户端直至操作完成后返回信息。

带有更多参数的命令样例：

```
POST /_snapshot/<my_first_repository>/<daily_backup_202003281600>/_restore?wait_for_completion=true
{
  "indices": "index_2,index_3",
  "ignore_unavailable": true,
  "include_global_state": true,
  "rename_pattern": "index_(.+)",
  "rename_replacement": "restored_index_$1"
}
```

参数	说明
indices	指定需要恢复的快照名称,支持多索引语法。
ignore_unavailable	控制是否忽略不可用的索引。

include_global_state	控制是否恢复集群的全局状态。
rename_pattern	使用正则表达式匹配需要重命名的索引。
rename_replacement	指定重命名规则。

该命令返回结果样例：

```
{
  "snapshot" : {
    "snapshot" : "daily_backup_202003281600",
    "indices" : [
      "restored_index_3",
      "restored_index_2"
    ],
    "shards" : {
      "total" : 2,
      "failed" : 0,
      "successful" : 2
    }
  }
}
```

参考信息

关于创建快照、从快照恢复数据、查询快照信息等操作的更多信息，请参考Elasticsearch官方文档：

Snapshot And Restore

预置插件

UES 提供了插件机制,当前我们的 UES 服务内置了以下常用插件,暂不支持用户自己上传安装插件。

[_ \[参考链接\]\(https://www.elastic.co/guide/en/elasticsearch/plugins/current/index.html\)_](https://www.elastic.co/guide/en/elasticsearch/plugins/current/index.html)

Analysis Plugins

- **analysis-ik**

IK为常用的中文分词插件。需要注意的是,IK插件提供的中文分词仅对新生成并指定IK作为分词器的index有效,对之前创建的index无效。

- **analysis-pinyin**

拼音分析插件,用来做汉字和拼音之间的转换。

- **analysis-icu**

使用 ICU 实现的一个针对亚洲语言的分词器插件。

- **analysis-kuromoji**

使用 Kuromoji analyzer 实现的一个日语的分词器插件。

- **analysis-smartcn**

针对中文或者中英文混合的文本的分词器插件。

- **analysis-ukrainian**

Provides stemming for Ukrainian.

Ingest Plugins

- **ingest-attachment**

使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。

- **ingest-geoip**

GeoIP处理器根据Maxmind数据库的数据添加有关IP地址位置的信息。该处理器默认在地理位置字段下添加此信息。地理位置处理器可以解析IPv4和IPv6地址。

- **ingest-user-agent**

该插件默认使用由uap-java提供的具有Apache 2.0许可证的regexes.yaml。 插件处理器从浏览器发送的用户代理字符串中提取其Web请求的详细信息。该处理器默认在user_agent字段下添加此信息。

Elasticsearch Head

elasticsearch-head是比较常用的集群管理工具。

官方ES在5.0以上版本中已不再提供直接安装head插件。内置的head插件做为一个独立的工程运行,head插件默认端口为9100。

SQL

可以使用户以类SQL语法查询、分析存储在Elasticsearch中的数据。

Basic Usage

* Simple query

```
/_sql?sql=select * from indexName limit 10
```

* Explain SQL to elasticsearch query DSL

```
/_sql/_explain?sql=select * from indexName limit 10
```

SQL Usage

* Query

```
SELECT * FROM bank WHERE age > 30 AND gender = 'm'
```

* Aggregation

```
select COUNT(*),SUM(age),MIN(age) as m, MAX(age),AVG(age) FROM bank GROUP BY gender ORDER BY SUM(age), m DESC
```

* Delete

```
DELETE FROM bank WHERE age > 30 AND gender = 'm'
```

Beyond sql

* Search

```
SELECT address FROM bank WHERE address = matchQuery('880 Holmes Lane') ORDER BY _score DESC LIMIT 3
```

* Aggregations

```
# range age group 20-25,25-30,30-35,35-40
```

```
SELECT COUNT(age) FROM bank GROUP BY range(age, 20,25,30,35,40)
```

```
# range date group by day
```

```
SELECT online FROM online GROUP BY date_histogram(field='insert_time','interval'='1d')
```

```
# range date group by your config
```

```
SELECT online FROM online GROUP BY date_range(field='insert_time','format'='yyyy-MM-dd' , '2014-08-18','2014-08-17','now-8d','now-7d','now-6d','now')
```

* ES Geographic

```
SELECT * FROM locations WHERE GEO_BOUNGING_BOX(fieldname,100.0,1.0,101,0.0)
```

* Select type

```
SELECT * FROM indexName/type
```

该插件使用来自开源项目, [参考链接](https://github.com/NLPchina/elasticsearch-sql)\

Snapshot/Restore Repository Plugins

- **Hadoop HDFS Repository Plugin**

该插件是针对最新的Apache Hadoop 2.x构建的,为使用HDFS文件系统作为快照/恢复存储库提供支持。

插件管理

在UES控制台集群管理页面, 点击“更多”按钮打开下拉菜单, 从中选择“插件管理”菜单项, 可以进入插件管理界面。



插件管理界面列出了UES预置的所有插件及其安装状态。

[<](#) [集群管理](#) / [集群详情](#)

集群概况

节点信息

监控视图

插件管理

☐	analysis-ukrainian	default	Provides stemming for Ukrainian.	● 已安装	卸载
☐	elasticsearch-repository-ufie	default	该插件支持将Elasticsearch数据的快照备份至UFile，并可从已备份至UFile的快照中将数据恢复至Elasticsearch。	● 未安装	安装
☐	ingest-attachment	default	使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。	● 已安装	卸载
☐	repository-hdfs	default	该插件是针对最新的Apache Hadoop 2.x构建的，为使用HDFS文件系统作为快照/恢复存储库提供支持。	● 已安装	卸载
☐	ucloud-alerting	default	监控告警插件，支持监控ES中的数据，并根据预设的规则通过邮件发送告警信息。	● 未安装	安装

[<](#) [1](#) [2](#) [>](#) 10 条/页 [v](#) [1/2](#)

1. 安装插件

在插件管理界面,选中处于“未安装”状态的插件,点击右侧“安装”按钮可以进行所选择插件的安装。

[←](#) [集群管理](#) / [集群详情](#)

集群概况

节点信息

监控视图

插件管理

☐	analysis-ukrainian	default	Provides stemming for Ukrainian.	● 已安装	卸载
☐	elasticsearch-repository-ufie	default	该插件支持将Elasticsearch数据的快照备份至UFile,并可从已备份至UFile的快照中将数据恢复至Elasticsearch。	● 未安装	安装
☐	ingest-attachment	default	使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。	● 已安装	卸载
☐	repository-hdfs	default	该插件是针对最新的Apache Hadoop 2.x构建的,为使用HDFS文件系统作为快照/恢复存储库提供支持。	● 已安装	卸载
☒	ucloud-alerting	default	监控告警插件,支持监控ES中的数据,并根据预设的规则通过邮件发送告警信息。	● 未安装	安装

[←](#) [1](#) [2](#) [>](#) 10 条/页 [▾](#) [/2](#)

插件的安装采用逐个节点安装并重启的方式进行。



插件安装过程中，状态显示为“安装中”，待状态显示为“已安装”时，表示该插件已安装完成。

[集群管理](#) / 集群详情

集群概况

节点信息

监控视图

插件管理

☐	analysis-ukrainian	default	Provides stemming for Ukrainian.	● 已安装	卸载
☐	elasticsearch-repository-ufie	default	该插件支持将Elasticsearch数据的快照备份至UFile，并可从已备份至UFile的快照中将数据恢复至Elasticsearch。	● 未安装	安装
☐	ingest-attachment	default	使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。	● 已安装	卸载
☐	repository-hdfs	default	该插件是针对最新的Apache Hadoop 2.x构建的，为使用HDFS文件系统作为快照/恢复存储库提供支持。	● 已安装	卸载
☐	ucloud-alerting	default	监控告警插件，支持监控ES中的数据，并根据预设的规则通过邮件发送告警信息。	● 安装中	卸载

[<](#) [1](#) [2](#) [>](#) 10 条/页 [v](#) [1/2](#)

2. 卸载插件

在插件管理界面,选中处于“已安装”状态的插件,点击右侧“卸载”按钮可以进行所选择插件的卸载。

[←](#) [集群管理](#) / [集群详情](#)

集群概况

节点信息

监控视图

插件管理

☐	analysis-ukrainian	default	Provides stemming for Ukrainian.	● 已安装	卸载
☐	elasticsearch-repository-ufie	default	该插件支持将Elasticsearch数据的快照备份至UFile,并可从已备份至UFile的快照中将数据恢复至Elasticsearch。	● 未安装	安装
☐	ingest-attachment	default	使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。	● 已安装	卸载
☐	repository-hdfs	default	该插件是针对最新的Apache Hadoop 2.x构建的,为使用HDFS文件系统作为快照/恢复存储库提供支持。	● 已安装	卸载
☒	ucloud-alerting	default	监控告警插件,支持监控ES中的数据,并根据预设的规则通过邮件发送告警信息。	● 已安装	卸载

[←](#) [1](#) [2](#) [>](#) 10 条/页 [▾](#) [/2](#)

插件的卸载采用逐个节点卸载并重启的方式进行。



插件卸载过程中, 状态显示为“卸载中”, 待状态显示为“未安装”时, 表示该插件已卸载完成。

[< 集群管理 / 集群详情](#)

集群概况

节点信息

监控视图

插件管理

☐	analysis-ukrainian	default	Provides stemming for Ukrainian.	● 已安装	卸载
☐	elasticsearch-repository-ufie	default	该插件支持将Elasticsearch数据的快照备份至UFile，并可从已备份至UFile的快照中将数据恢复至Elasticsearch。	● 未安装	安装
☐	ingest-attachment	default	使用 Apache Tika 实现的用来解析 PPT、XLS、PDF 和 Microsoft Word 等文档的插件。	● 已安装	卸载
☐	repository-hdfs	default	该插件是针对最新的Apache Hadoop 2.x构建的，为使用HDFS文件系统作为快照/恢复存储库提供支持。	● 已安装	卸载
☐	ucloud-alerting	default	监控告警插件，支持监控ES中的数据，并根据预设的规则通过邮件发送告警信息。	● 卸载中	安装

[<](#) [1](#) [2](#) [>](#) 10 条/页 [v](#) [/2](#)

Security插件

关于安全插件的详细使用方法, 请参考UES安全插件使用说明。

一、特别提示

- 1) 安装安全插件将会触发**UES**集群重启, 对**ES**服务造成影响。请在安装之前做好必要的准备工作, 选择合适的时间进行安装。
- 2) 安全插件安装完成之后, 访问**ES**服务的代码需要修改, 加上账号密码信息方可正常使用。

二、插件安装

1.支持的版本

当前支持安全插件的UES服务版本为6.5.4, 其他版本的UES暂不支持安装安全插件。

2.安全插件安装方法

安全插件的安装方法与其他插件类似, 具体方式可参见“插件管理”小节的内容。

三、注意事项

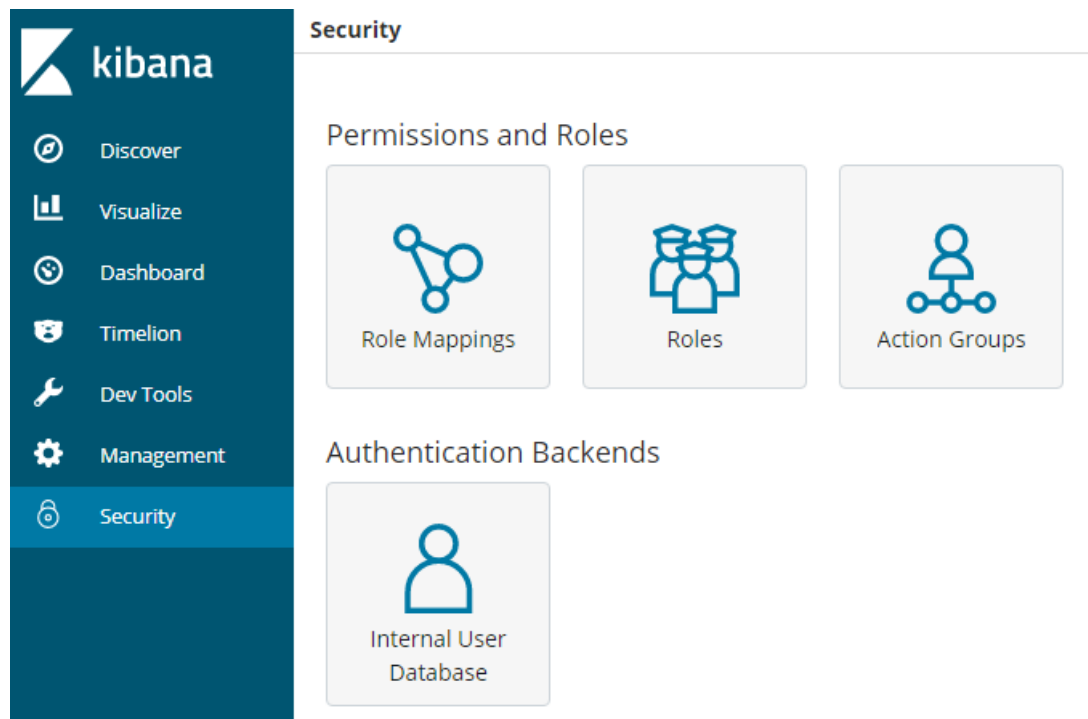
1) 安全插件的安装采用逐个节点进行安装并重启的方式来进行,一个节点的安全插件安装完成后,由于安全机制发生变化,将无法加入原集群。需要等待所有节点的安全插件都安装完成之后,集群才能恢复提供服务,一般情况下,拥有3个节点的集群安装安全插件大约需要15分钟时间。

2) 安全插件安装完成之后,访问ES服务的代码需要修改,加上账号密码信息方可正常使用(默认沿用创建集群时设置的Kibana的账号密码作为启用安全机制后的账号密码)。示例如下:

```
curl -H "Content-Type: application/json" -u admin:admin -XGET http://localhost:9200/_cat/health?v
```

四、权限管理

在安全插件安装完成之后,使用创建集群时设置的Kibana账号密码(也是启用安全机制后的账号密码)登录Kibana,在Kibana的菜单栏中,选择“Security”菜单项进入权限管理界面。



角色是控制集群访问的核心方式。角色包含集群范围权限, 特定于索引的权限, 文档和字段级安全性以及租户的任意组合。通过将用户映射到这些角色, 可以让用户获得这些权限。

Security插件附带了许多预定义的操作组, 角色, 映射和用户, 可以根据需要从中选用。

Roles

Role 	Cluster permissions	Indices	Tenants
all_access  RESERVED	UNLIMITED	*	admin_tenant
kibana_read_only  RESERVED			
kibana_server  RESERVED	CLUSTER_COMPOSITE_OPS CLUSTER_MONITOR cluster:admin/xpack/monito... indices:admin/template* indices:data/read/scroll*	* ?kibana ?kibana-6 ?kibana_* ?management-beats* ?monitoring* ?reporting* ?tasks	
kibana_user  RESERVED	CLUSTER_COMPOSITE_OPS INDICES_MONITOR	* ?kibana ?kibana-6 ?kibana_* ?management-beats ?tasks	

US3插件

US3插件支持将Elasticsearch数据的快照备份至US3,并可从已备份至US3的快照中将数据恢复至Elasticsearch。

一、安装插件

可以通过UES控制台的“插件管理”功能进行US3插件的安装,插件名称为elasticsearch-repository-ufire,具体操作方法请参考文档插件管理章节。

二、创建US3存储空间

1.创建存储空间

在US3控制台创建用于UES数据备份的存储空间,为了确保网络传输质量,请选择与目标**UES**集群相同的地域进行创建。空间类型请选择**“私有空间”**。

创建存储空间

选择地域

北京

香港

广州

上海二

洛杉矶

新加坡

雅加达

拉各斯

圣保罗

迪拜

台北

空间类型 [?]

私有空间

公开空间

业务组

未分组 ▼

存储空间域名 ^{*} [?]

ues-backup

.uae-dubai.ufileos.com

数据分析 [?]

数据湖分析USQL

未授权

取消

确定

2.创建令牌

在US3控制台创建用于访问上述存储空间的令牌, 请注意选择上述创建的存储进行授权, 并在“令牌权限”中勾选“上传”、“下载”、“删除”、“文件列表”等。

创建令牌

令牌名称 *

ues-backup

有效时长 *

1

年

授权存储空间 ①

选择存储空间

已选择: ues-backup

授权文件 ①

所有文件

设置前缀

令牌权限

☒ 上传 ☒ 下载 ☒ 删除 ☒ 文件列表 ☐ 图片处理

取消

确定

3. 获取存储空间内网域名及密钥信息

获取并记录以下信息以备使用：

(1) 存储空间内网域名

请注意US3控制台展示的存储空间域名是外网域名，您需要查看US3文档获取其内网域名，详见US3文档的“地域和域名”一节。

以上述创建的示例存储空间为例，其外网域名为：ues-backup.uae-dubai.ufileos.com，查询得到US3迪拜的内网域名为：www.internal-uae-dubai.ufileos.com，故示例存储空间的内网域名为：**ues-backup.internal-uae-dubai.ufileos.com**

[创建存储空间](#)[更改业务组](#)[删除存储空间](#)

☐	存储空间域名 1	地域 ▼	业务组 1	空间类型 ▼
☐		上海二	未分组	私有空间
☐	ues-backup <u>uae-dubai.ufileos.com</u>	迪拜	未分组	私有空间

需要将此处外网域名替换为内网域名

(2) 拥有上述存储空间访问权限的令牌密钥

在US3控制台“令牌管理”页面查看令牌的公钥和私钥。

[创建令牌](#)

令牌名称 1	密钥
MyToken	公钥 : *****06a55 私钥 : *****91220
ues-backup	公钥 : *****828ba 私钥 : *****c9dee

三、创建仓库

在对ES数据创建快照或者从快照恢复数据之前,需要先创建一个仓库。

```
PUT /_snapshot/<1>
{
  "type": "ufile",
  "settings": {
    "endpoint": <2>,
    "public_key": <3>,
    "private_key": <4>,
    "bucket": <5>,
    "compress": <6>,
    "chunk_size": <7>,
    "base_path": <8>,
    "max_snapshot_bytes_per_sec": <9>,
    "max_restore_bytes_per_sec": <10>
  }
}
```

<1>:ES备份仓库的名称

<2> endpoint:上述US3存储空间的内网域名

<3> public_key:上述US3令牌的公钥

<4> private_key:上述US3令牌的私钥

<5> bucket:上述US3存储空间的名称

<6> compress:是否对备份的索引进行压缩(true / false)

<7> chunk_size:上传文件的分片大小,默认为64MB

<8> base_path:备份文件在bucket中的路径(前缀名称),默认为根路径(无前缀)

<9> max_snapshot_bytes_per_sec:生成快照时的上传速度,默认40MB/s

<10> max_restore_bytes_per_sec:从快照恢复时的下载速度,默认40MB/s

一个典型的仓库创建请求示例如下:

```
PUT /_snapshot/ues_snapshot
{
  "type": "ufile",
  "settings": {
    "endpoint": "ues-backup.internal-uae-dubai.ufileos.com",
    "public_key": "TOKEN_XXXXXXXX-XXXX-XXXX-XXXX-XXXXXXXXXXXX",
    "private_key": "XXXXXXXX-XXXX-XXXX-XXXX-XXXXXXXXXXXX",
    "bucket": "ues-backup",
    "compress": true,
    "chunk_size": "50mb",
    "base_path": "ues",
    "max_snapshot_bytes_per_sec": "20mb",
    "max_restore_bytes_per_sec": "20mb"
  }
}
```

四、使用方法

关于创建快照、从快照恢复数据、查询快照信息等操作的具体方法, 请参考Elastic官方文档:

Snapshot And Restore

五、常用命令

查询所有仓库信息:

```
GET /_snapshot?pretty
```

查询一个仓库信息:

```
GET /_snapshot/<repository>?pretty
```

查询一个仓库中所有快照:

```
GET /_snapshot/<repository>/_all?pretty
```

查询一个指定快照:

```
GET /_snapshot/<repository>/<snapshot>?pretty
```

创建一个快照：

```
PUT /_snapshot/<repository>/<snapshot>
```

创建指定索引的一个快照：

```
PUT /_snapshot/<repository>/<snapshot>
{
  "indices": "<index1>[,index2]..."
}
```

恢复一个指定快照：

```
POST /_snapshot/<repository>/<snapshot>/_restore
```

删除一个指定快照：

```
DELETE /_snapshot/<repository>/<snapshot>
```

ES测试

ES集群创建成功后(创建参考操作指南),可以进行集群可用性的初步测试。这里示例通过一台可以访问 ues 的 uhost 进行 curl 调用命令测试。

集群状态测试

```
curl -s -XGET 'http://<host>:9200/_cluster/health?pretty'
```

正常情况,系统返回结果:

```
{
  "cluster_name" : "ues-qwerty",
  "status" : "green",
  "timed_out" : false,
  "number_of_nodes" : 3,
  "number_of_data_nodes" : 3,
  "active_primary_shards" : 1,
  "active_shards" : 2,
  "relocating_shards" : 0,
  "initializing_shards" : 0,
  "unassigned_shards" : 0,
```

```
"delayed_unassigned_shards" : 0,
"number_of_pending_tasks" : 0,
"number_of_in_flight_fetch" : 0,
"task_max_waiting_in_queue_millis" : 0,
"active_shards_percent_as_number" : 100.0
}
```

说明集群是正常的运行状态了,之后进行测试数据写入。

索引一个文档

例,索引为 "ucloud",类型为 "information",选择"1"作为ID的编号。这时,请求就是这样的:

```
curl -X PUT \
http://<host>:9200/ucloud/information/1 \
-H 'Content-Type: application/json' \
-d '{
  "name": "UCloud",
  "age": 5,
  "about": "at shanghai",
  "interests": ["service"]
}'
```

创建成功返回

```
{"_index":"ucloud","_type":"information","_id":"1","_version":1,"result":"created","_shards":{"total":2,"successful":2,"failed":0},"created":true}
```

检索文档

```
curl -X GET 'http://<host>:9200/ucloud/information/1'
```

返回

```
{"_index":"ucloud","_type":"information","_id":"1","_version":1,"found":true,"_source":{"name": "UCloud",  
"age": 5,  
"about": "at shanghai",  
"interests": ["service"]  
}}
```

搜索文档

```
curl -X GET 'http://<host>:9200/_search'  
curl -X GET 'http://<host>:9200/ucloud/_search'  
curl -X GET 'http://<host>:9200/ucloud/information/_search'  
curl -X GET "http://<host>:9200/_sql" -H 'Content-Type: application/json' -d'select * from ucloud limit 10'
```

复杂搜索

复杂搜索需使用POST方法, 详情请参考 [功能文档](#) [API文档](#)

更新文档

```
curl -X PUT \  
http://<host>:9200/ucloud/information/1 \  
-H 'Content-Type: application/json' \  
-d '{  
  "name" : "UCloud",  
  "age": 5,  
  "about": "at shanghai",  
  "interests": ["service", "professional"]  
}'
```

更新成功返回

```
{"_index":"ucloud","_type":"information","_id":"1","_version":2,"result":"updated","_shards":{"total":2,"successful":2,"failed":0},"created":false}
```

删除

删除文档

```
curl -X DELETE 'http://<host>:9200/ucloud/information/1'
```

删除指定类型文档

```
curl -X DELETE 'http://<host>:9200/ucloud/information'
```

删除索引

```
curl -X DELETE 'http://<host>:9200/ucloud'
```

配置管理

本章将列出可在控制台修改的静态配置项和可**API**动态更新的配置项。

控制台支持修改常见静态配置参数到elasticsearch.yml配置文件。这些配置不支持动态更新,修改后需重启集群。

配置项	默认值	配置描述
cluster.name	instance_id	集群名称
discovery.zen.ping.unicast. .hosts	[]	单播地址,默认主节点IP或前三个节点IP,示例[127.0.0.1,127.0.0.2,127.0.0.3]
gateway.recover_after_n odes	N - 2	预期的节点数量达到,就可以进行恢复
gateway.expected_nodes	N	预计集群节点数量。当预期的节点数量加入到群集,分片的恢复就会开始
gateway.recover_after_ti me	5m	如果未达到预期的节点数量,则不会尝试恢复,恢复过程都会等待配置的时间量
http.cors.enabled	true	允许跨域访问
http.cors.allow-origin	*	允许跨域访问的域,默认支持所有域
rest.action.multi.allow_ex plicit_index	true	允许Body中的index参数覆盖URL中的index参数

node.attr.tag		节点tag,默认不设置
gateway.expected_master_nodes	0	预计集群主节点数量。当预期的主节点数量加入集群,分片的恢复就会开始
gateway.expected_data_nodes	0	预计集群数据节点数量。当预期的数据节点数量加入集群,分片的恢复就会开始
gateway.recover_after_master_nodes	0	预期的主节点数量达到,就可以进行恢复
gateway.recover_after_data_nodes	0	预期的数据节点数量达到,就可以进行恢复
http.max_content_length	100mb	HTTP请求的最大内容,默认100mb。如果设置大于Integer.MAX_VALUE,将被重置为100mb
http.max_initial_line_length	4kb	HTTP URL的最大长度。默认4kb
http.max_header_size	8kb	允许的headers的最大值。默认8kb
http.compression	true	尽可能支持压缩(使用Accept-Encoding)
http.compression_level	3	定义HTTP响应的压缩级别。有效值在1(最小压缩)和9(最大压缩)的范围内
http.cors.max-age	1728000	浏览器发送"preflight"请求来确定CORS设置。max-age定义了结果应该被缓存的时间。默认为20天
http.cors.allow-methods	OPTIONS,HEAD,GET,POST,PUT,DELETE	允许的方法

http.cors.allow-headers	X-Requested-With,Content-Type,Content-Length	允许的headers
http.cors.allow-credentials	false	是否应该返回Access-Control-Allow-Credentials header。注意:只有在设置为true的情况下才会返回此标题
http.detailed_errors.enabled	true	在响应输出中启用或禁用详细错误消息和堆栈跟踪的输出。注意:如果设置为false,并指定了error_trace请求参数,则会返回错误;当没有指定error_trace时,会返回一个简单的消息
http.pipelining	true	启用或禁用HTTP流水线
http.pipelining.max_events	10000	在HTTP连接关闭之前在内存中排队的最大事件数量
indices.fielddata.cache.size	unlimited	字段数据高速缓存的最大值,例如节点堆空间的30%,或绝对值,如12GB
indices.queries.cache.size	10%	控制过滤器缓存的内存大小,默认10%。接受百分比值(如5%)或精确值(如512mb)
index.queries.cache.enabled	true	控制是否启用查询缓存,设置是可以基于每个索引配置的索引设置
indices.memory.index_buffer_size	10%	接受百分比或字节大小值。默认10%,这意味着分配给节点的总堆的10%将被用作在所有分片之间共享的索引缓冲区
indices.memory.min_index_buffer_size	48mb	如果将index_buffer_size指定为百分比,则可以使用此设置指定绝对最小值
indices.memory.max_index_buffer_size	unlimited	如果将index_buffer_size指定为百分比,则可以使用此设置指定绝对最大值

indices.requests.cache.size	1%	缓存在节点级别进行管理, 默认最大值为堆的1%
node.ingest	true	负载器资格
search.remote.connect	true	跨集群搜索
index.number_of_shards	5	索引应该具有的主分片的数量, 此设置可在索引创建时动态设置
index.shard.check_on_startup	false	在分片打开之前是否检查有无损坏现象。当检测到损坏时, 它将防止分片被打开
index.routing_partition_size	1	自定义路由值可以转到的分片数量。默认1, 只能在索引创建时设置。此值必须小于index.number_of_shards, 除非index.number_of_shards值也为1

- 注: N为节点个数

上面没有列出的配置项大多都可以通过 `**_[cluster-update-settings](https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster-update-settings.html)\_*` 动态修改。动态配置项详细描述请参考 `_[Modules](https://www.elastic.co/guide/en/elasticsearch/reference/current/modules.html)\_*`

下面列出常见的可 动态更新 的配置参数。

配置项	默认值	配置描述
cluster.routing.allocation.enable	all	启用或禁用分配特定种类的分片
cluster.routing.allocation.node_concurrent_incoming_recoveries	2	允许在一个节点上发生多少个并发传入分片恢复

cluster.routing.allocation.node_concurrent_outgoing_recoveries	2	允许在一个节点上发生多少个并行传出分片恢复
cluster.routing.allocation.node_concurrent_recoveries	2	node_concurrent_incoming_recoveries和node_concurrent_outgoing_recoveries快捷设置
cluster.routing.allocation.node_initial_primaries_recoveries	4	启动节点后恢复未分配的主节点将使用本地磁盘中的数据
cluster.routing.allocation.same_shard.host	false	允许执行检查单个主机上分配同一分片的多个实例
cluster.routing.rebalance.enable	all	启用或禁用特定种类的分片重新平衡
cluster.routing.allocation.allow_rebalance	indices_all_active	指定何时允许分片重新平衡
cluster.routing.allocation.cluster_concurrent_rebalance	2	允许控制集群范围允许多少个并发分片重新平衡
cluster.routing.allocation.balance.shard	0.45f	定义节点上分配的分片总数(浮点数)的权重因子
cluster.routing.allocation.balance.index	0.55f	定义在特定节点上分配的每个索引的分片数(浮点数)的权重因子
cluster.routing.allocation.balance.threshold	1.0f	应该执行的操作的最小优化值(非负浮点数)
cluster.routing.allocation.disk.threshold_enabled	true	设置为false可禁用磁盘分配决策程序
cluster.routing.allocation.disk.watermark.low	85%	控制磁盘使用的低水位, 默认85%。超出ES就不会为节点分配新的分片

cluster.routing.allocation.disk.watermark.high	90%	控制高水位, 默认90%。超出ES将尝试将分片重定位到另一个节点
cluster.routing.allocation.disk.watermark.flood_stage	95%	控制洪水阶段, 默认95%, 超出ES强制在至少一个磁盘超出的节点上分配一个或多个分片的每个索引上的只读索引块。一旦有足够的磁盘空间可用于索引操作继续, 索引块必须手动释放
cluster.info.update.interval	30s	ES应该检查集群中每个节点的磁盘使用情况
cluster.routing.allocation.disk.include_relocations	true	ES将在计算节点的磁盘使用情况时考虑当前正在重定位到目标节点的分片
cluster.routing.allocation.awareness.attributes.*		分片分配感知设置允许您告诉ES您的硬件配置
cluster.routing.allocation.include.{attribute}		将分片分配给{属性}至少有一个逗号分隔值的节点, attribute 可能是_name、_ip、_host
cluster.routing.allocation.require.{attribute}		只将分片分配给{属性}具有所有逗号分隔值的节点, attribute 可能是_name、_ip、_host
cluster.routing.allocation.exclude.{attribute}		不要将分片分配给{属性}没有逗号分隔值的节点, attribute 可能是_name、_ip、_host
cluster.blocks.read_only	false	使整个群集只读(索引不接受写操作), 数据不允许被修改(创建或删除索引)
cluster.blocks.read_only_allow_delete	false	与cluster.blocks.read_only相同, 但允许删除索引以释放资源
cluster.indices.tombstones.size	500	集群状态维护索引墓碑以明确表示已被删除的索引
logger.org.elasticsearch.indices.recovery	INFO	日志记录级别

discovery.zen.ping.unicast.hosts.resolve_timeout	5s	在每轮ping之前等待DNS查找的时间
discovery.zen.ping_timeout	3s	ping超时时间
discovery.zen.join_timeout	60s	默认超时时间是ping超时的20倍
discovery.zen.minimum_master_nodes	$N/2 + 1$	最小主节点资格数, N为主节点个数
discovery.zen.no_master_block	write	设置控制在没有活动的主节点时应该拒绝哪些操作, 默认write
indices.breaker.total.limit	70%	总体的parent breaker的起始限制, 默认JVM堆的70%
indices.breaker fielddata.limit	60%	fielddata breaker的限制, 默认JVM堆的60%
indices.breaker fielddata.overhead	1.03	一个常数, 所有的fielddata估计相乘决定最后的估算
indices.breaker.request.limit	60%	request breaker的限制, 默认JVM堆的60%
indices.breaker.request.overhead	1	一个常数, 所有的请求估计相乘决定最后的估算
network.breaker.inflight_requests.limit	100%	inflight_requests breaker的限制, 默认JVM堆的100%
network.breaker.inflight_requests.overhead	1	一个常数, 所有的飞行中的请求估计相乘决定最后的估算
script.max_compilations_rate	75/5m	一定间隔内允许编译的唯一动态脚本的数量限制
index.requests.cache.enable	true	启用或禁用索引缓存
indices.recovery.max_bytes_per_sec	40mb	数据在节点间传输最大带宽
indices.store.throttle.max_bytes_per_sec	20mb	写磁盘最大带宽

index.merge.scheduler.max_thread_count		索引merge最大线程数,默认为 $\max(1, \min(4, \text{availableProcessors} / 2))$
index.number_of_replicas	1	每个主分片具有的副本数量
index.auto_expand_replicas	默认false (即禁用)	根据可用节点的数量自动扩展副本的数量。设置为划定短划线的下限和上限 (例如0-5)
index.refresh_interval	1s	多久执行一次刷新操作,这会使最近对索引进行的更改可见,从而进行搜索。可以设置为-1来禁用刷新
index.max_result_window	10000	搜索索引的from + size的最大值
index.max_inner_result_window	100	内部匹配定义和顶部的from + size的最大值将聚合到此索引
index.max_rescore_window	10000	用于搜索此索引的重新调用请求的window_size的最大值。默认index.max_result_window
index.max_docvalue_fields_search	100	查询中允许的最大docvalue_fields数量
index.max_script_fields	32	查询中允许的script_fields的最大数
index.max_ngram_diff	1	NGramTokenizer和NGramTokenFilter的min_gram和max_gram的最大允许差值
index.max_shingle_diff	3	max_shingle_size和min_shingle_size的最大允许差值
index.blocks.read_only	false	设置为true以使索引和索引数据只读,否则允许写入和数据更改
index.blocks.read_only_allow_delete	false	与index.blocks.read_only相同,但允许删除索引以释放资源
index.blocks.read	false	设置为true以禁止对索引执行读取操作
index.blocks.write	false	设置为true以禁止对索引执行写入操作
index.blocks.metadata	false	设置为true以禁用索引数据读取和写入
index.max_refresh_listeners		索引的每个分片上可用的刷新监听器的最大数量。这些监听器用于实现refresh = wait_for

插件管理

Elasticsearch-Head

最新elasticsearch版本(5.0以上)已经不再支持使用内部 **elasticsearch-plugin** 命令来安装head插件,内置的head插件本质是调用es服务REST API包装而形成的图形化管理工具,所以使用head插件需要head管理平台网络和输入的es连接地址是相通的。

Hadoop HDFS Repository Plugin

* 通过REST API定义hdfs存储库的配置:

```
PUT /_snapshot/my_hdfs_repository
{
  "type": "hdfs",
  "settings": {
    "uri": "hdfs://namenode:8020/",
    "path": "elasticsearch/respositories/my_hdfs_repository"
  }
}

# uri: hdfs的地址,例如:"hdfs://<host>:<port>/"
```

```
# path: 数据存储/加载的文件系统中的文件路径,例如:"path/to/file"
```

创建成功,可以获取hdfs仓库信息:

```
GET /_snapshot/my_hdfs_repository
```

示例

```
{
  "my_hdfs_repository": {
    "type": "hdfs",
    "settings": {
      "path": "elasticsearch/respositories/my_hdfs_repository",
      "uri": "hdfs://namenode:8020/"
    }
  }
}
```

* 创建快照:

```
PUT /_snapshot/my_hdfs_repository/snapshot_1
{
  "indices": "index_1,index_2",
  "ignore_unavailable": true,
  "include_global_state": false
}
```

```
}
```

创建成功, 可以获取快照信息:

```
GET /_snapshot/my_hdfs_repository/snapshot_1
```

示例

```
{
  "snapshots": [
    {
      "snapshot": "snapshot_1",
      "uuid": "yr9T6jtLTCEVFRoNGN-9Lw",
      "version_id": 5050199,
      "version": "5.5.1",
      "indices": [
        ".kibana",
        "ucloud"
      ],
      "state": "SUCCESS",
      "start_time": "2018-02-01T08:13:26.128Z",
      "start_time_in_millis": 1517472806128,
      "end_time": "2018-02-01T08:13:28.870Z",
      "end_time_in_millis": 1517472808870,
      "duration_in_millis": 2742,
```

```
"failures": [],  
"shards": {  
  "total": 6,  
  "failed": 0,  
  "successful": 6  
}  
}  
]  
}
```

* 删除快照：

```
DELETE /_snapshot/my_hdfs_repository/snapshot_1
```

* 快照恢复：

```
POST /_snapshot/my_hdfs_repository/snapshot_1/_restore
```

注意： 若集群中已经存在需要快照恢复的索引并且为 open 状态,需要使用 `\_close` API先关闭该索引,示例：

```
POST /.kibana/_close
```

更详细插件使用请参考 [_ \[Hadoop HDFS Repository Plugin\]\(https://www.elastic.co/guide/en/elasticsearch/reference/current/modules-snapshots.html#_repository_plugins\)](https://www.elastic.co/guide/en/elasticsearch/reference/current/modules-snapshots.html#_repository_plugins)

IK Analysis Plugin

自定义分词词库操作

通过在 IK 配置文件中提到的如下配置：

```
<?xml version="1.0" encoding="UTF-8"?>
<!DOCTYPE properties SYSTEM "http://java.sun.com/dtd/properties.dtd">
<properties>
<comment>IK Analyzer 扩展配置</comment>
<!--用户可以在这里配置自己的扩展字典 -->
<entry key="ext_dict">custom/mydict.dic;custom/single_word_low_freq.dic</entry>
<!--用户可以在这里配置自己的扩展停止词字典-->
<entry key="ext_stopwords">custom/ext_stopword.dic</entry>
<!--用户可以在这里配置远程扩展字典 -->
<entry key="remote_ext_dict">location</entry>
<!--用户可以在这里配置远程扩展停止词字典-->
<entry key="remote_ext_stopwords">http://xxx.com/xxx.dic</entry>
</properties>
```

IK分词支持本地自定义词库,并且支持远程热更新词库,UES通过 词库信息保存在特定索引中 来完成自定义词库的更新。下面给出在 **Kibana** 中的具体操作流程,其中 **API** 要保持完全一致。

* 本地扩展字典

```
PUT /custom_ik/analyzer/1
{
  "ext_dict": [
  ]
}
```

示例：

```
PUT /custom_ik/analyzer/1
{
  "ext_dict": [
    "中华人民共和国",
    "美利坚合众国",
    "大不列颠和北爱尔兰联合王国"
  ]
}
```

* 本地扩展停止词字典

```
PUT /custom_ik/analyzer/2
{
  "ext_stopwords": [
  ]
}
```

示例：

```
PUT /custom_ik/analyzer/2
{
  "ext_stopwords": [
    "中华人民共和国",
    "美利坚合众国",
    "大不列颠和北爱尔兰联合王国"
  ]
}
```

* 远程扩展字典

```
PUT /custom_ik/analyzer/3
{
  "remote_ext_dict": ""
}
```

示例：

```
PUT /custom_ik/analyzer/3
{
  "remote_ext_dict": "http://localhost:8080/my_dict.dic"
}
```

* 远程扩展停止词字典

```
PUT /custom_ik/analyzer/4
{
  "remote_ext_stopwords": ""
}
```

示例：

```
PUT /custom_ik/analyzer/4
{
  "remote_ext_stopwords": "http://localhost:8080/my_stopwords.dic"
}
```

检测索引数据是否成功

```
GET /custom_ik/analyzer/1
GET /custom_ik/analyzer/2
GET /custom_ik/analyzer/3
GET /custom_ik/analyzer/4
```

* 最后还需一步，到控制台点击需要更新的分词词库项完成配置文件的更新，并重启集群节点使之生效。

实例运维

生产环境的集群要健康运行及故障处理需要合理、高效的运营经验。本章将提供比较实用的在线运营实例。

集群健康

在 Elasticsearch 集群中可以监控统计很多信息,其中最重要的就是:集群健康(**Cluster Health**)。它的 status 有 `Green`、`Yellow`、`Red` 三种。

```
GET /_cluster/health
```

系统返回:

```
{
  "cluster_name" : "ues-qwerty",
  "status" : "green",
  "timed_out" : false,
  "number_of_nodes" : 3,
  "number_of_data_nodes" : 3,
  "active_primary_shards" : 1,
  "active_shards" : 2,
  "relocating_shards" : 0,
  "initializing_shards" : 0,
```

```
"unassigned_shards" : 0,
"delayed_unassigned_shards" : 0,
"number_of_pending_tasks" : 0,
"number_of_in_flight_fetch" : 0,
"task_max_waiting_in_queue_millis" : 0,
"active_shards_percent_as_number" : 100.0
}
```

响应信息中的一块就是`status`，这是我们最应该关注的字段，它告诉我们当前集群是否处于一个可用的状态。其中三种颜色分别代表：

状态	描述	备注
`green`	所有主分片和从分片都可用	所有的主分片和副本分片都已分配。集群是 100% 可用的。
`yellow`	所有主分片可用,但存在不可用的从分片,即存在未分配(unassigned)的从分片	所有的主分片已经分片了,但至少还有一个副本是缺失的。不会有数据丢失,所以搜索结果依然是完整的。不过,你的高可用性在某种程度上被弱化。如果更多的分片消失,会丢数据了。把`yellow`想象成一个需要及时调查的警告。
`red`	存在不可用的主要分片,即存在未分配(unassigned)的主分片	至少一个主分片(以及它的全部副本)都在缺失中。这意味着你在缺少数据:搜索只能返回部分数据,而分配到这个分片上的写入请求会返回一个异常。

为保障集群为`green`状态,即所有主从分片都可用。强烈建议索引主分片设置小于数据节点个数的**3倍**，副本设置小于数据节点个数。

单个节点监控

集群健康是对集群的所有信息进行高度概述,而节点统计值 API 提供一个让人眼花缭乱的统计数据的数组,包含集群的每一个节点统计值。

节点统计值 API 可以通过如下命令执行：

```
GET /_nodes/stats
```

在输出内容的开头,我们可以看到集群名称和我们的第一个节点:

```
{
  "_nodes": {
    "total": 3,
    "successful": 3,
    "failed": 0
  },
  "cluster_name": "ues-qwerty",
  "nodes": {
    "cZabQfdFTVC0dh9Zxrjocg": {
      "timestamp": 1515997798238,
      "name": "ues-qwerty-01",
      "transport_address": "192.168.1.1:9300",
      "host": "192.168.1.1",
      "ip": "192.168.1.1:9300",
      "roles": [
        "master",
        "data",
        "ingest"
      ],
      "indices": {
```

```
...
```

节点是排列在一个哈希里,以节点的 UUID 作为键名。还显示了节点网络属性的一些信息,这些值对调试诸如节点未加入集群这类自动发现问题很有用。通常你会发现是端口用错了,或者节点绑定在错误的 IP 地址/网络接口上了。

索引部分

索引(indices)部分列出了这个节点上所有索引的聚合过的统计值：

```
"indices": {  
  "docs": {  
    "count": 0,  
    "deleted": 0  
  },  
  "store": {  
    "size_in_bytes": 0,  
    "throttle_time_in_millis": 0  
  },  
}
```

* docs 展示节点内存有多少文档,包括还没有从段里清除的已删除文档数量。

* store 部分显示节点耗用了多少物理存储。这个指标包括主分片和副本分片在内。如果限流时间很大,那可能表明你的磁盘限流设置得过低。

```
"indexing": {  
  "index_total": 2,  
  "index_time_in_millis": 146,  
}
```

```
"index_current": 0,  
"index_failed": 0,  
"delete_total": 1,  
"delete_time_in_millis": 6,  
"delete_current": 0,  
"noop_update_total": 0,  
"is_throttled": false,  
"throttle_time_in_millis": 0  
},  
"get": {  
"total": 2,  
"time_in_millis": 7,  
"exists_total": 2,  
"exists_time_in_millis": 7,  
"missing_total": 0,  
"missing_time_in_millis": 0,  
"current": 0  
},  
"search": {  
"open_contexts": 0,  
"query_total": 50,  
"query_time_in_millis": 74,  
"query_current": 0,  
"fetch_total": 40,  
"fetch_time_in_millis": 25,
```

```
"fetch_current": 0,
"scroll_total": 0,
"scroll_time_in_millis": 0,
"scroll_current": 0,
"suggest_total": 0,
"suggest_time_in_millis": 0,
"suggest_current": 0
},
"merges": {
"current": 0,
"current_docs": 0,
"current_size_in_bytes": 0,
"total": 0,
"total_time_in_millis": 0,
"total_docs": 0,
"total_size_in_bytes": 0,
"total_stopped_time_in_millis": 0,
"total_throttled_time_in_millis": 0,
"total_auto_throttle_in_bytes": 104857600
},
```

* indexing 显示已经索引了多少文档。这个值是一个累加计数器,在文档被删除时,数值不会下降,在发生文档更新等内部索引操作时,值会增加。此外还列出了索引操作耗费的时间,正在索引的文档数量,以及删除操作的类似统计值。

* get 显示通过 ID 获取文档的接口相关的统计值。包括对单个文档的 GET 和 HEAD 请求。

* `search` 描述在活跃中的搜索 (`open_contexts`) 数量、查询的总数量、以及自节点启动以来在查询上消耗的总时间。用 `query\_time\_in\_millis / query\_total` 计算的比值, 可以用来粗略的评价你的查询有多高效。比值越大, 每个查询花费的时间越多, 你应该要考虑调优了。

`fetch` 统计值展示了查询处理的后半流程 (query-then-fetch 里的 `fetch`)。如果 `fetch` 耗时比 `query` 还多, 说明磁盘较慢, 或者获取了太多文档, 或者可能搜索请求设置了太大的分页 (比如, `size: 10000`)。

* `merges` 包括了 Lucene 段合并相关的信息。它会告诉你目前在运行几个合并, 合并涉及的文档数量, 正在合并的段的总大小, 以及在合并操作上消耗的总时间。

在你的集群写入压力很大时, 合并统计值非常重要。合并要消耗大量的磁盘 I/O 和 CPU 资源。如果你的索引有大量的写入, 同时又发现大量的合并数, 一定要去阅读 [**\[索引性能技巧\]](https://www.elastic.co/guide/cn/elasticsearch/guide/current/indexing-performance.html) (<https://www.elastic.co/guide/cn/elasticsearch/guide/current/indexing-performance.html>)[**](#)。

```
"filter_cache": {
  "memory_size_in_bytes": 48,
  "evictions": 0
},
"fielddata": {
  "memory_size_in_bytes": 0,
  "evictions": 0
},
"segments": {
  "count": 319,
  "memory_in_bytes": 65812120
},
...
```

* `filter\_cache` 展示了已缓存的过滤器位集合所用的内存数量,以及过滤器被驱逐出内存的次数。过多的驱逐数 可能 说明你需要加大过滤器缓存的大小,或者你的过滤器不太适合缓存(比如它们因为高基数而在大量产生,就像是缓存一个 `now` 时间表达式)。

不过,驱逐数是一个很难评定的指标。过滤器是在每个段的基础上缓存的,而从一个小的段里驱逐过滤器,代价比从一个大的段里要廉价的多。有可能你有很大的驱逐数,但是它们都发生在小段上,也就意味着这些对查询性能只有很小的影响。把驱逐数指标作为一个粗略的参考。如果你看到数字很大,检查一下你的过滤器,确保他们都是正常缓存的。不断驱逐着的过滤器,哪怕都发生在很小的段上,效果也比正确缓存住了的过滤器差很多。

* `field\_data` 显示 `fielddata` 使用的内存,用以聚合、排序等等。这里也有一个驱逐计数。和 `filter\_cache` 不同的是,这里的驱逐计数是很有用的:这个数应该或者至少是接近于 0。因为 `fielddata` 不是缓存,任何驱逐都消耗巨大,应该避免掉。如果你在这里看到驱逐数,你需要重新评估你的内存情况,`fielddata` 限制,请求语句,或者这三者。

* `segments` 会展示这个节点目前正在服务中的 Lucene 段的数量。这是一个重要的数字。大多数索引会有大概 50-150 个段,哪怕它们存有 TB 级别的数十亿条文档。段数量过大表明合并出现了问题(比如,合并速度跟不上段的创建)。注意这个统计值是节点上所有索引的汇聚总数。

`memory` 统计值展示了 Lucene 段自己用掉的内存大小。这里包括底层数据结构,比如倒排表,字典,和布隆过滤器等。太大的段数量会增加这些数据结构带来的开销,这个内存使用量就是一个方便用来衡量开销的度量值。

操作系统和进程部分

`OS` 和 `Process` 部分基本是自描述的,不会在细节中展开讲解。它们列出来基础的资源统计值,比如 CPU 和负载。OS 部分描述了整个操作系统,而 Process 部分只显示 Elasticsearch 的 JVM 进程使用的资源情况。

这些都是非常有用的指标,不过通常在你的监控技术栈里已经都测量好了。统计值包括下面这些:

* CPU

* 负载

* 内存使用率

* Swap 使用率

* 打开的文件描述符

JVM部分

`jvm` 部分包括了运行 Elasticsearch 的 JVM 进程一些很关键的信息。最重要的,它包括了垃圾回收的细节,这对你的 Elasticsearch 集群的稳定性有着重大影响。

```
"jvm": {  
  "timestamp": 1515997798245,  
  "uptime_in_millis": 1704518268,  
  "mem": {  
    "heap_used_in_bytes": 143523960,  
    "heap_used_percent": 3,  
    "heap_committed_in_bytes": 4277534720,  
    "heap_max_in_bytes": 4277534720,  
    "non_heap_used_in_bytes": 95476824,  
    "non_heap_committed_in_bytes": 102416384,
```

线程池部分

Elasticsearch 在内部维护了线程池。这些线程池相互协作完成任务,有必要的话相互间还会传递任务。通常来说,你不需要配置或者调优线程池,不过查看它们的统计值有时候还是有用的,可以洞察你的集群表现如何。

这有一系列的线程池,但以相同的格式输出:

```
"index": {  
  "threads": 1,  
  "queue": 0,  
  "active": 0,  
  "rejected": 0,  
  "largest": 1,  
  "completed": 1  
}
```

每个线程池会列出已配置的线程数量 (threads)，当前在处理任务的线程数量 (active)，以及在队列中等待处理的任务单元数量 (queue)。

如果队列中任务单元数达到了极限，新的任务单元会开始被拒绝，你会在 rejected 统计值上看到它反映出来。这通常是你的集群在某些资源上碰到瓶颈的信号。因为队列满意味着你的节点或集群在用较高的速度运行，但依然跟不上工作的蜂拥而入。

值得关注的线程池部分有：

- * indexing

普通的索引请求的线程池

- * bulk

批量请求，和单条的索引请求不同的线程池

- * get

Get-by-ID 操作

- * search

所有的搜索和查询请求

* merging

专用于管理 Lucene 合并的线程池

文件系统和网络部分

继续向下阅读 `/_node/stats` API, 你会看到一串和你的文件系统相关的统计值: 可用空间, 数据目录路径, 磁盘 I/O 统计值, 等等。如果你没有监控磁盘可用空间的话, 可以从这里获取这些统计值。磁盘 I/O 统计值也很方便, 不过通常那些更专门的命令行工具 (比如 `iostat`) 会更有用些。

显然, Elasticsearch 在磁盘空间满的时候很难运行——所以请确保不会这样。

还有两个跟 网络统计值 相关的部分:

```
"transport": {
  "server_open": 26,
  "rx_count": 8834703,
  "rx_size_in_bytes": 19706396482,
  "tx_count": 8834702,
  "tx_size_in_bytes": 18212792172
},
"http": {
  "current_open": 3,
  "total_opened": 153
},
```

* transport 显示和 传输地址 相关的一些基础统计值。包括节点间的通信 (通常是 9300 端口) 以及任意传输客户端或者节点客户端的连接。如果看到这里有很多连接数不要担心; Elasticsearch 在节点之间维护了大量的连接。

* http 显示 HTTP 端口(通常是 9200)的统计值。如果你看到 total_opened 数很大而且还在一直上涨,这是一个明确信号,说明你的 HTTP 客户端里有没启用 keep-alive 长连接的。持续的 keep-alive 长连接对性能很重要,因为连接、断开套接字是很昂贵的(而且浪费文件描述符)。请确认你的客户端都配置正确。

断路器

fielddata 断路器相关的统计值:

```
"fielddata": {  
  "limit_size_in_bytes": 2566520832,  
  "limit_size": "2.3gb",  
  "estimated_size_in_bytes": 0,  
  "estimated_size": "0b",  
  "overhead": 1.03,  
  "tripped": 0  
},
```

这里你可以看到断路器的最大值(比如,一个请求申请更多的内存时会触发断路器)。这个部分还会让你知道断路器被触发了多少次,以及当前配置的间接开销。间接开销用来铺垫评估,因为有些请求比其他请求更难评估。

主要需要关注的是 tripped 指标。如果这个数字很大或者持续上涨,这是一个信号,说明你的请求需要优化,或者你需要添加更多内存(单机上添加,或者通过添加新节点的方式)。

集群统计

集群统计 API 提供了和 节点统计 相似的输出。但有一个重要的区别:节点统计显示的是每个节点上的统计值,而 集群统计 展示的是对于单个指标,所有节点的总和值。

这里面提供一些很值得一看的统计值。比如说你可以看到, 整个集群用了 50% 的堆内存, 或者说过滤器缓存的驱逐情况不严重。这个接口主要用途是提供一个比 **集群健康** 更详细、但又没有 **节点统计** 那么详细的快速概览。对于非常大的集群来说也很有用, 因为那时候 **节点统计** 的输出已经非常难于阅读了。

这个 API 可以像下面这样调用:

```
GET /_cluster/stats
```

索引统计

有时候我们想从 **索引为中心** 的角度看统计值: 索引 收到了多少个搜索请求? 索引 获取文档耗费了多少时间?

要做到这点, 选择你感兴趣的索引 (或者多个索引) 然后执行一个索引级别的 **统计 API**:

```
# 统计 my_index 索引
GET /my_index/_stats

# 使用逗号分隔索引名可以请求多个索引统计值
GET /my_index,another_index/_stats

# 使用特定的 _all 可以请求全部索引的统计值
GET /_all/_stats
```

返回的统计信息和 **节点统计** 的输出很相似: search 、 fetch 、 get 、 index 、 bulk 、 segment counts 等等。

索引为中心的统计在有些时候很有用, 比如辨别或验证集群中的 **热门 索引**, 或者试图找出某些索引比其他索引更快或者更慢的原因。

实践中,节点为中心的统计还是显得更有用些。瓶颈往往是针对整个节点而言,而不是对于单个索引。因为索引一般是分布在多个节点上的,这导致以索引为中心的统计值通常不是很有用,因为它们是从不同环境的物理机器上汇聚的数据。

索引为中心的统计作为一个有用的工具可以保留在你的技能表里,但是通常它不会是第一个用的上的工具。

等待中的任务

有一些任务只能由主节点去处理,比如创建一个新的索引或者在集群中移动分片。由于一个集群中只能有一个主节点,所以只有这一节点可以处理集群级别的元数据变动。在 99.9999% 的时间里,这不会有什么问题。元数据变动的队列基本上保持为零。

在一些 罕见 的集群里,元数据变动的次数比主节点能处理的还快。这会导致等待中的操作会累积成队列。

等待中的任务 API 会给你展示队列中 (如果有的话) 等待的集群级别的元数据变更操作:

```
GET /_cluster/pending_tasks
```

通常,响应都是像这样的:

```
{
  "tasks": []
}
```

cat API

cat API 对一些集群信息获取会非常有用。

通过 GET 请求发送 cat 命名可以列出所有可用的 API:

GET /_cat

```
=^.^=  
/_cat/allocation  
/_cat/shards  
/_cat/shards/{index}  
/_cat/master  
/_cat/nodes  
/_cat/tasks  
/_cat/indices  
/_cat/indices/{index}  
/_cat/segments  
/_cat/segments/{index}  
/_cat/count  
/_cat/count/{index}  
/_cat/recovery  
/_cat/recovery/{index}  
/_cat/health  
/_cat/pending_tasks  
/_cat/aliases  
/_cat/aliases/{alias}  
/_cat/thread_pool
```

```
/_cat/thread_pool/{thread_pools}  
/_cat/plugins  
/_cat/fielddata  
/_cat/fielddata/{fields}  
/_cat/nodeattrs  
/_cat/repositories  
/_cat/snapshots/{repository}  
/_cat/templates
```

动态变更设置

Elasticsearch 里很多设置都是动态的,可以通过 API 修改。需要强制重启节点(或者集群)的配置修改都要极力避免。而且虽然通过静态配置项也可以完成这些变更,我们建议你还是用 API 来实现。

集群更新 API 有两种工作模式:

* 临时(Transient)

这些变更在集群重启之前一直会生效。一旦整个集群重启,这些配置就被清除。

* 永久(Persistent)

这些变更会永久存在直到被显式修改。即使全集群重启它们也会存活下来并覆盖掉静态配置文件里的选项。

临时或永久配置需要在 JSON 体里分别指定:

```
PUT /_cluster/settings
```

```
{
  "persistent" : {
    "discovery.zen.minimum_master_nodes" : 2
  },
  "transient" : {
    "indices.store.throttle.max_bytes_per_sec" : "50mb"
  }
}
```

这个永久设置会在全集群重启时存活下来

这个临时设置会在第一次全集群重启后被移除

日志记录

Elasticsearch 会输出很多日志, 默认的日志记录等级是 INFO 。它提供了适度的信息, 但是又设计好了不至于让你的日志太过庞大。

当调试问题的时候, 特别是节点发现相关的问题 (因为这个经常依赖于各式过于繁琐的网络配置), 提高日志记录等级到 DEBUG 是很有帮助的。

调高节点发现的日志记录级别:

```
PUT /_cluster/settings
{
  "transient" : {
    "logger.discovery" : "DEBUG"
  }
}
```

```
}
```

慢日志

还有另一个日志叫 慢日志 。这个日志的目的是捕获那些超过指定时间阈值的查询和索引请求。这个日志用来追踪由用户产生的很慢的请求很有用。

默认情况,慢日志是不开启的。要开启它,需要定义具体动作(query,fetch 还是 index),你期望的事件记录等级(WARN 、 DEBUG 等),以及时间阈值。

这是一个索引级别的设置,也就是说可以独立应用给单个索引:

```
PUT /my_index/_settings
{
  "index.search.slowlog.threshold.query.warn": "10s",
  "index.search.slowlog.threshold.fetch.debug": "500ms",
  "index.indexing.slowlog.threshold.index.info": "5s"
}

# 查询慢于 10 秒输出一个 WARN 日志
# 获取慢于 500 毫秒输出一个 DEBUG 日志
# 索引慢于 5 秒输出一个 INFO 日志
```

一旦阈值设置过了,你可以和其他日志器一样切换日志记录等级:

```
PUT /_cluster/settings
{
  "transient" : {
```

```
"logger.index.search.slowlog" : "DEBUG",  
"logger.index.indexing.slowlog" : "WARN"  
}  
}
```

设置搜索慢日志为 DEBUG 级别

设置索引慢日志为 WARN 级别

历史数据清理

随着数据的不断写入,节点的磁盘使用率就会不断攀升,超过一定阈值就会出现不能分配分片的现象。一种解决方法就是删除部分索引数据来使集群节点的磁盘使用率保持在一种合理状态。

删除历史数据,可以通过网络互通的 云主机UHost 或 UES的 Kibana 来实现。通过API可以删除通配符的一类索引或特定索引,也可以借助脚本来定时清理。

获取集群索引信息:

```
# UHost  
curl -s -XGET 'http://<host>:9200/_cat/indices?v'  
  
# Kibana  
GET /_cat/indices?v
```

删除索引:

```
# UHost
curl -s -XDELETE 'http://<host>:9200/index'

# Kibana
DELETE /index
```

定时清理脚本 (此脚本适用于logstash写入索引有时间标记的情况, 其他情况可参考修改):

```
#!/bin/bash
#author: UCloud
#created at 2017/08/30 19:00

if test ! -f "/var/log/elkDailyDel.log"; then
touch /var/log/elkDailyDel.log
fi

# 获取索引信息
# 请将下行当中的localhost改成自己elasticsearch服务对应的Ip
indices=$(curl -s -XGET "localhost:9200/_cat/indices?v" |grep 'logstash' |awk '{print $3}')

# 设置保留索引的时间线
thirtyDaysAgo=$(date -d "$(date +%Y%m%d) -30 days" "+%s")
function DelOrNot(){
if [ $((($1-$2)) -ge 0 )]; then
echo 1
```

```
else
echo 0
fi
}

for index in ${indices}
do
indexDate=`echo ${index##*-} |sed 's/\./-/g`
indexTime=`date -d "${indexDate}" "+%s"`
if [ `DelOrNot ${indexTime} ${thirtyDaysAgo}` -eq 0 ]; then
# 索引时间在时间线之前执行删除操作
# 请将下行当中的localhost改成你自己elasticsearch服务对应的ip
result=`curl -s -XDELETE "localhost:9200/${index}"`
echo "delResult is ${result}" >> /var/log/elkDailyDel.log
if [ `echo ${result} |grep 'acknowledged' |wc -l` -eq 1 ]; then
echo "${index} had already been deleted!" >> /var/log/elkDailyDel.log
else
echo "there is something wrong happend when deleted ${index}" >> /var/log/elkDailyDel.log
fi
fi
done

echo $?
```

数据冷热分离

SSD磁盘类型的集群随着数据的不断写入,节点的磁盘数据逐渐增大。可能对于一些较老的数据,不需要太多的查询或不再做查询使用,那么这些数据就会占用大量的SSD性能资源和存储空间。数据冷热分离通过配置使最新的数据保存在ssd磁盘节点上,较老的数据自动迁移到廉价sata节点,使用有限的ssd节点资源来实现同时支持高并发读写和大数据量的存储。

UES也给出了数据冷热分离的策略,过程如下:

* 准备

集群实例:

SSD (愿集群):ues-qwerty,前三个节点ip分别为 ues-qwerty-ip1、ues-qwerty-ip2、ues-qwerty-ip3

SATA (新集群):ues-asdfgh

1、愿集群(ues-qwerty)修改配置

控制台上修改 **参数管理** 中 **node.attr.tag** 项为 **hot**。

2、线性重启原集群(ues-qwerty)

3、固定原索引在原集群节点上(ues-qwerty)

```
# UHost
curl -XPUT localhost:9200/*/_settings -d '{"settings": {"index.routing.allocation.require.tag": "hot"}}'

# Kibana
PUT /*/_settings
{
  "settings": {
```

```
"index.routing.allocation.require.tag": "hot"
}
}
```

4、新集群(ues-asdfgh) 修改配置,同1

cluster.name: ues-qwerty

discovery.zen.ping.unicast.hosts: [ues-qwerty-ip1,ues-qwerty-ip2,ues-qwerty-ip3]

node.attr.tag: cold

5、重启新集群(ues-asdfgh)

6、老数据迁移到 cold 节点

```
# UHost
curl -XPUT localhost:9200/cold_index/_settings -d '{"settings": {"index.routing.allocation.require.tag": "cold"}}'

# Kibana
PUT /cold_index/_settings
{
  "settings": {
    "index.routing.allocation.require.tag": "cold"
  }
}
```

定时迁移脚本示例：

```
#!/bin/bash
#author: UCloud
#created at 2017/08/30 19:00

time=`date -d last-day "+%Y.%m.%d"`

# 请将下行当中的localhost改成自己elasticsearch服务对应的Ip
curl -XPUT http://localhost:9200/*-${time}/_settings -d'
{
  "index.routing.allocation.require.tag": "cold"
}'
```

本章主要参考：

[《Elasticsearch: 权威指南》](<https://www.elastic.co/guide/cn/elasticsearch/guide/current/administration.html>)

实例迁移

创建索引

数据迁移本质是索引的重建,重建索引不会尝试设置目标索引,它不会复制源索引的设置。所以在操作之前设置目标索引,包括设置映射,分片数,副本等。

数据迁移

1. Reindex from Remoteedit

Reindex支持从远程Elasticsearch集群重建索引:

```
POST _reindex
{
  "source": {
    "remote": {
      "host": "http://otherhost:9200",
      "username": "user",
      "password": "pass"
    },
    "index": "source",
```

```
"query": {
  "match": {
    "test": "data"
  }
},
"dest": {
  "index": "dest"
}

# host参数必须包含scheme、host和port (例如https:// otherhost:9200)
# username和password参数可选
```

使用时需要在elasticsearch.yml中配置 `reindex.remote.whitelist` 属性。可以设置多组 (例如, `otherhost:9200`, `another:9200`, `127.0.10.*:9200`, `localhost:*`)。

具体使用可参考 **[Reindex from Remoteedit]**(<https://www.elastic.co/guide/en/elasticsearch/reference/current/docs-reindex.html#reindex-from-remote>)_

2. Elasticsearch-Dump

Elasticsearch-Dump是一个elasticsearch数据导入导出开源工具包。安装、迁移相关执行可以在相同可用区的云主机上进行,使用方便。

* ** Installing **

需要node环境, npm安装elasticsearch-dump

```
npm install elasticsearch -g
elasticsearch
```

**** Use ****

Copy an index from production to staging with analyzer and mapping:

```
elasticsearch \
--input=http://production.es.com:9200/my_index \
--output=http://staging.es.com:9200/my_index \
--type=analyzer
elasticsearch \
--input=http://production.es.com:9200/my_index \
--output=http://staging.es.com:9200/my_index \
--type=mapping
elasticsearch \
--input=http://production.es.com:9200/my_index \
--output=http://staging.es.com:9200/my_index \
--type=data
```

Copy a single shard data:

```
elasticsearch \
--input=http://es.com:9200/api \
--output=http://es.com:9200/api2 \
--params='{"preference" : "_shards:0"}'
```

elasticdump命令其他参数使用参考 [_ \[Elasticdump Options\]\(https://github.com/taskrabbit/elasticsearch-dump#options\)_](https://github.com/taskrabbit/elasticsearch-dump#options)

3. Elasticsearch-Migration

Elasticsearch-Migration是基于Go语言开源的工具包。和Elasticsearch-Dump一样, 相关操作可以在相同可用区的云主机上进行。

* ** Download **

[_ \[下载地址\]\(https://github.com/medcl/esm/releases\)_](https://github.com/medcl/esm/releases)

* ** Use **

```
./bin/esm -s http://192.168.1.x:9200/ -d http://192.168.1.y:9200/ -x src_index -y dest_index -w=5 -b=100
```

Options

-s, --source= source elasticsearch instance

-d, --dest= destination elasticsearch instance

-q, --query= query against source elasticsearch instance, filter data before migrate, ie: name:medcl

-m, --source_auth basic auth of source elasticsearch instance, ie: user:pass

-n, --dest_auth basic auth of target elasticsearch instance, ie: user:pass

-c, --count= number of documents at a time: ie "size" in the scroll request (10000)

--sliced_scroll_size= size of sliced scroll, to make it work, the size should be > 1, default:"1"

-t, --time= scroll time (1m)

--shards= set a number of shards on newly created indexes

--copy_settings copy index settings from source

--copy_mappings copy mappings mappings from source

-f, --force delete destination index before copying, default:false

```
-x, --src_indexes= list of indexes to copy, comma separated (_all), support wildcard match(*)
-y, --dest_index= indexes name to save, allow only one indexname, original indexname will be used if not specified
-a, --all copy indexes starting with . and _ (false)
-w, --workers= concurrency number for bulk workers, default is: "1"
-b --bulk_size bulk size in MB" default:5
-v --log setting log level,options:trace,debug,info,warn,error
-i --input_file indexing from local dump file, file format: {"_id":"xxx","_index":"xxx","_source":{"xxx":"xxx"},"_type":"xxx" }
-o --output_file output documents of source index into local file, file format same as input_file.
--source_proxy set proxy to source http connections, ie: http://127.0.0.1:8080
--dest_proxy set proxy to destination http connections, ie: http://127.0.0.1:8080
--refresh refresh after migration finished
```

Elasticsearch-Migration注意事项 [\[参考链接\]\(https://github.com/medcl/esm\)](https://github.com/medcl/esm)

故障恢复

yellow状态分析恢复

`yellow` 状态表明存在未分配(unassigned)的从分片。

查询索引状态

```
curl -s -XGET 'http://<host>:9200/_cat/indices?v'  
curl -s -XGET 'http://<host>:9200/_cluster/health?level=indices'
```

查询未分配的分片

```
curl -s -XGET 'http://<host>:9200/_cat/shards?v' | grep UNASSIGNED  
curl -s -XGET 'http://<host>:9200/_cluster/health?level=shards'
```

* 索引副本设置不合理

如果是索引副本数设置大于数据节点个数而导致集群处于 `yellow` 状态, 修改副本数来修复集群状态

```
curl -XPUT \  
http://<host>:9200/unassigned_index/_settings \  
-H 'Content-Type: application/json' \  

```

```
-d '{
  "index": {
    "number_of_replicas": replicasCount
  }
}'

# unassigned_index为未分配的分片索引
# replicasCount为新的索引副本数
```

正常情况下,未分配的从分片会自动得到分配,集群状态也会恢复`green`。某些特殊情况可能需要手动分配掉未分配的从分片

```
curl -XPOST \
  http://<host>:9200/_cluster/reroute \
  -H 'Content-Type: application/json' \
  -d '{
    "commands": [{
      "allocate_replica": {
        "index": "unassigned_index",
        "shard": num,
        "node": "nodeName"
      }
    }]
  }'

# unassigned_index为未分配的分片索引
```

```
# num为未分配的分片编号  
# nodeName为节点名称,也可以为节点编号ID,如kVWVil1PQt2Bk2rP7PlrbQ
```

在放弃和离开分片之前,集群将尝试在一行中分配一个最大的index.allocation.max_retries时间片(默认为5)。这种情况可能是由结构性问题引起的,例如有一个分析器引用了所有节点上不存在的停用词文件。一旦问题得到解决,可以通过retry_failed调用 [_cluster/reroute API](https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster-reroute.html)(<https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster-reroute.html>)来手动重试分配,这将试图对这些分片进行一次重试。

```
POST /_cluster/reroute?retry_failed=true
```

集群更糟糕的情况可能出现未分配的主分片,需手动分配请参考 [_cluster/reroute API](https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster-reroute.html)(<https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster-reroute.html>)

* 节点磁盘使用率超出阈值

节点磁盘使用率对分片分配的影响,参考 [_Disk-based Shard Allocation](https://www.elastic.co/guide/en/elasticsearch/reference/6.2/disk-allocator.html)(<https://www.elastic.co/guide/en/elasticsearch/reference/6.2/disk-allocator.html>)

如果是磁盘使用率高导致集群出现未分配分片,可以选择修改磁盘使用策略暂时缓解 或 扩容节点个数。

另外,如果确定可以永久关闭部分历史索引数据,也可以通过删除此类索引来恢复集群状态。删除历史索引数据可参考 [_实例运维 - 历史数据清理](https://docs.ucloud.cn/ues/develop/online)(<https://docs.ucloud.cn/ues/develop/online>)

Logstash部署

安装Logstash

Logstash5.5需要Java 8,不支持Java 9。使用官方的Oracle发行版或OpenJDK等开源发行版。

安装参考 [_ \[Installing-Logstash\]\(https://www.elastic.co/guide/en/logstash/current/installing-logstash.html\)_](https://www.elastic.co/guide/en/logstash/current/installing-logstash.html)

配置文件编写

要配置Logstash,需要创建一个配置文件,指定要使用的插件以及每个插件的设置。您可以引用配置中的事件字段,并使用条件来处理符合特定条件的事件。当运行logstash时,使用-f来指定配置文件。

创建一个名为`logstash-simple.conf`的文件,并将其保存在与Logstash相同的目录中。

```
input { stdin { } }  
output {  
  elasticsearch { hosts => ["<host>:9200"] }  
  stdout { codec => rubydebug }  
}
```

运行logstash并使用-f标志指定配置文件。

```
bin/logstash -f logstash-simple.conf
```

更多配置示例参考 [_【Logstash-Config-Examples】\(https://www.elastic.co/guide/en/logstash/current/config-examples.html\)_](https://www.elastic.co/guide/en/logstash/current/config-examples.html)

功能文档

UES基于开源Elasticsearch版本构建, 相关文档参考如下。

ES开发指南

中文文档请参考:[《Elasticsearch: 权威指南》](https://www.elastic.co/guide/cn/elasticsearch/guide/cn/index.html)

英文文档请参考:[《Elasticsearch Reference》](https://www.elastic.co/guide/en/elasticsearch/reference/current/index.html)

ES API文档

[API Conventions](https://www.elastic.co/guide/en/elasticsearch/reference/current/api-conventions.html)

[Document APIs](https://www.elastic.co/guide/en/elasticsearch/reference/current/docs.html)

[Search APIs](https://www.elastic.co/guide/en/elasticsearch/reference/current/search.html)

[Indices APIs](https://www.elastic.co/guide/en/elasticsearch/reference/current/indices.html)

[cat APIs](https://www.elastic.co/guide/en/elasticsearch/reference/current/cat.html)

[Cluster APIs](https://www.elastic.co/guide/en/elasticsearch/reference/current/cluster.html)

Logstash

英文文档请参考:[《Logstash Reference》](https://www.elastic.co/guide/en/logstash/current/index.html)

[Configuring Logstash](https://www.elastic.co/guide/en/logstash/current/configuration.html)

[Input plugins](https://www.elastic.co/guide/en/logstash/current/input-plugins.html)

[Output plugins](https://www.elastic.co/guide/en/logstash/current/output-plugins.html)

[Filter plugins](https://www.elastic.co/guide/en/logstash/current/filter-plugins.html)

产品简介

ULogstash 是基于开源数据收集引擎 Logstash 构建的云端托管服务, 它是一个服务器端的数据处理管道, 支持动态的从不同来源采集和转换数据, 并将数据标准化到目标位置。

产品特性

- 快速部署、简单易用。
- 支持弹性扩展节点数量。
- 集成官方所有 Input、Output、Filter 插件。

产品使用场景

* 数据采集

支持多种的数据来源, 通过输入插件方便的采集日志、指标、应用、数据库、消息队列等来源数据。

* 数据迁移

自建Elasticsearch数据迁移

* 数据转化

通过过滤插件清理和转换数据, 如将非结构化数据解析导出结构、解析IP地址、标准化日期等。

创建Logstash实例

可点击 **ULogstash** 控制台 进入产品页

创建

< [ULogstash服务](#) / 创建集群

地域

地域可用区

 华北一  可用区B 

基础配置

服务版本

7.10.2 

节点告警模版

节点将绑定默认告警模板，若要调整监控项，请稍后在集群概况中调整告警配置。

默认 ULogstash 节点告警模板

节点配置

磁盘类型

RSSD云盘

计算规格

CPU内存比

1:2

1:4

o.logstash2m.medium

- CPU 2核
- 内存 4G

o.logstash2m.xlarge

- CPU 4核
- 内存 8G

o.logstash2m.2xlarge

- CPU 8核
- 内存 16G

o.logstash2m.4xlarge

- CPU 16核
- 内存 32G

o.logstash2m.8xlarge

- CPU 32核
- 内存 64G

磁盘大小

20 GB

节点数量

1 台

网络设置

所属VPC *

DefaultVPC

所属子网 *

DefaultNetwork

可用IP数:

更多设置

集群名称 *

业务组

未分组

页面会列出可创建服务的全部地域和可用区,选择需求的可用区、服务版本、磁盘类型、节点类型、磁盘大小、节点数量、填写一些初始化信息等内容后创建就可实现服务实例的快速部署。

节点规格、节点数量选择以及付费方式

其中可选的节点规格

按磁盘、规格分类：

磁盘类型	标识
RSSD云盘	o.logstash

可选节点规格详细配置(最新节点参考“UES价格”导航页)：

机型	名称	配置
内存密集型	o.logstash4m.medium	2核 8G
内存密集型	o.logstash4m.xlarge	4核 16G
内存密集型	o.logstash4m.2xlarge	8核 32G
内存密集型	o.logstash4m.4xlarge	16核 64G
普通机型	o.logstash2m.medium	2核 4G
普通机型	o.logstash2m.xlarge	4核 8G
普通机型	o.logstash2m.2xlarge	8核 16G
普通机型	o.logstash2m.4xlarge	16核 32G
普通机型	o.logstash2m.8xlarge	32核 64G

计费选择

计费模式分为按时、月付、年付3种类型。控制台支持计费类型修改。

确认支付

集群信息选择和填写完成后点击创建,到支付确认页面,完成支付创建。

实例部署

全部产品

UES

华北一

可用区C

消息

告警

帮助与支持

Elasticsearch服务 UES

Elasticsearch管理

ULogstash管理

创建

☐	集群名称	资源ID	服务版本	可用区▼	业务组▼	运行时长	创建时间⌵	状态▼	操作
☐			7.10.2	华北一可用区C	未分组			运行	详情编辑节点...

<1>

10 条/页

/1

等待部署中, 由于实例规模不同, 所需要的部署时间会有所差异, 部署大致时间在2分钟到5分钟。

部署中, 列表的实例的实例状态会显示为"蓝色"状态创建中, 创建成功后实例状态会动态变化为"绿色"状态运行。

如果长时间为创建状态, 可点击刷新重新获取实例状态。

集群信息

控制台会给出集群的基本信息展示。点击详情进入详情页，我们会提供集群概况、节点信息和集群监控视图等信息。

基本信息

基本信息

资源ID

资源名

服务版本

7.10.2

可用区

华北一可用区B

业务组

未分组

运行时长

1天5小时9分钟

运行状态

运行

所属VPC

所属子网

集群基本信息详情。

节点配置

节点信息	
节点数量	2
计算规格	o.logstash4m.medium
磁盘类型	RSSD云盘
磁盘大小	30G

集群节点配置详情。

付费信息

付费信息		续费
创建时间	2023-09-19 10:09:29	
到期时间	2023-09-20 16:00:00	
付费方式	1.51元/小时 按时	

集群付费详情。续费可以到计费系统进行续费、更改计费方式等操作。

节点列表

节点列表

节点IP	节点机型	计算规格	磁盘信息	状态▼	操作
	o.logstash4m.medium	2核 8G	RSSD云盘 30GB	运行	删除节点
	o.logstash4m.medium	2核 8G	RSSD云盘 30GB	运行	删除节点

<

1

>

10 条/页

/1

展示集群节点IP信息、节点配置信息及节点可用操作。

节点IP为业务网地址,默认是通过内网访问的，为了数据安全性，强烈不建议代理到外网环境访问。

释放实例

当 Logstash 实例不再需要时或无法满足您的需求,您可以在 Elasticsearch服务 控制台对实例进行释放,以避免服务继续运行而产生费用。

注意事项

实例释放后数据无法恢复,实例中所包含的管道会被删除,请谨慎操作。

操作步骤

1. 登录 ULogstash 控制台,进入ULogstash 实例列表页。

Elasticsearch服务 UES

Elasticsearch管理

ULogstash管理

创建

🔍

🔄

📄

☐	集群名称	资源ID	服务版本	可用区▼	业务组▼	运行时长	创建时间🕒	状态▼	操作
☐			7.10.2	华北一可用区B	未分组	1天5小时43分钟	2023-09-19	● 运行	详情编辑节点...
☐			7.10.2	华北一可用区B	未分组	1天0小时19分钟	2023-09-19	● 运行	详情编辑节点...

<1>

10条/页

/1

2. 在实例列表页, 选择需要释放的实例, 选择操作 > ... > 删除集群; 或单击实例 ID/名称进入实例基本信息页, 选择左上角删除集群。

Elasticsearch服务 UES

Elasticsearch管理 ULogstash管理

创建

☐	集群名称	资源ID	服务版本	可用区 ▼	业务组 ▼	运行时长	创建时间 升	状态 ▼	操作
☐	[REDACTED]	[REDACTED]	7.10.2	华北一可用区B	未分组	1天5小时43分钟	2023-09-19	● 运行	详情 编辑节点 ...
☐	[REDACTED]	[REDACTED]	7.10.2	华北一可用区B	未分组	1天0小时19分钟	2023-09-19	● 运行	详情 编辑节点 删除集群 更改业务组

< 1 > 10 条/页 /1

3. 在删除集群页面中, 点击确定, 系统将清空实例数据, 并回收资源, 数据清空后, 无法恢复。

实例改配

随着业务的发展,您可能对集群配置的要求有变化。当集群的配置无法满足业务需求时,可通过集群改配功能,更改集群配置。

注意事项

改配集群会触发集群重启,重启时间与集群规格、数据结构和大小等因素有关,建议在业务低峰期操作。

操作步骤

1. 登录 ULogstash 控制台。
2. 在顶部菜单栏处,选择地域。
3. 单击ULogstash管理,然后在ULogstash集群列表中单击目标集群名称。
4. 在节点信息页面,单击右侧的编辑节点。
5. 在编辑节点页面,选择相应的修改类型。

(1)增加节点,新增相同配置的节点。

当前节点信息

集群名称 [REDACTED] 资源ID [REDACTED] 节点个数 1个 (2核 | 8G)

磁盘信息 RSSD云盘 | 30GB

编辑节点

修改类型

增加节点

调整规格

修改磁盘大小

1 台

(2) 调整规格, 修改当前节点配置、节点重启生效。

当前节点信息

集群名称 [REDACTED] 资源ID [REDACTED] 节点个数 1个 (2核 | 8G)

磁盘信息 RSSD云盘 | 30GB

编辑节点

修改类型

调整规格将按节点顺序依次进行，调整规格节点期间会重启机器，建议在业务低峰期操作。

增加节点 调整规格 修改磁盘大小

计算规格

CPU内存比

1:2 1:4

o.logstash2m.medium

- CPU 2核
- 内存 4G

o.logstash2m.xlarge

- CPU 4核
- 内存 8G

o.logstash2m.2xlarge

- CPU 8核
- 内存 16G

o.logstash2m.4xlarge

- CPU 16核
- 内存 32G

o.logstash2m.8xlarge

- CPU 32核
- 内存 64G

(3) 修改磁盘大小, 磁盘只允许扩大。

当前节点信息

集群名称		资源ID		节点个数	1个 (2核 8G)
磁盘信息	RSSD云盘 30GB				

编辑节点

修改类型

磁盘扩容将对集群内所有节点逐个进行操作。

增加节点 调整规格 **修改磁盘大小**

30 GB

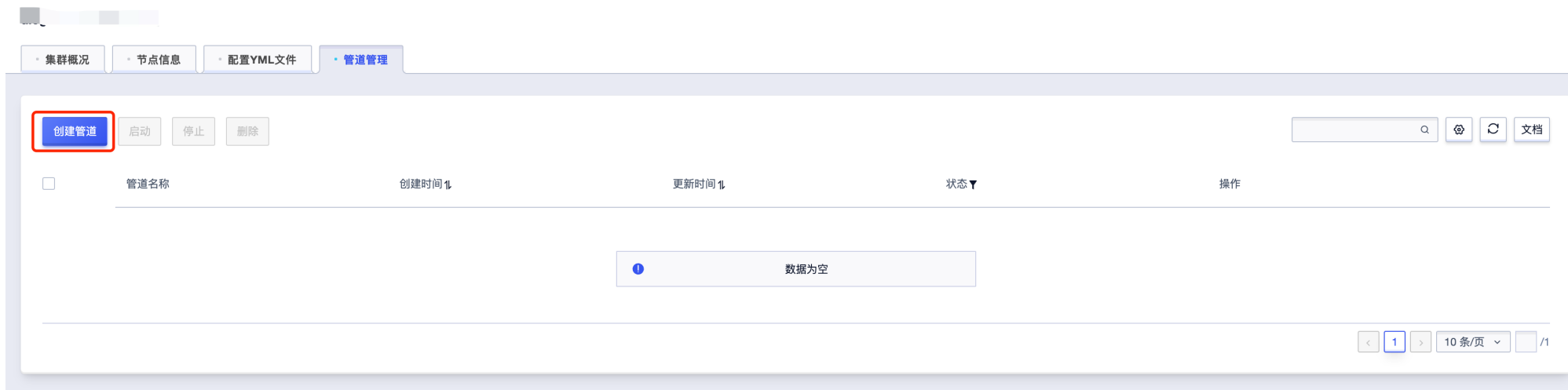
6. 选中对应修改配置, 单击立即购买。购买后, 集群会完成相应的配置更改。

管道管理

Logstash 通过管道来实现数据的采集处理,它包含必选的 input 和 output 插件,以及可选的 filter 插件,并支持多管道并行运行,。本文介绍如何通过配置文件管理管道,包括创建管道、修改管道和删除管道。

创建管道

1. 登录 ULogstash 控制台,进入 Logstash 实例列表页。
2. 在实例列表页,单击实例名称进入实例详情页,然后进入管道管理页签,单击创建管道。



3. 填写管道名称(支持字母、数字、_)

Config配置

管道名称 *

请输入，例：test_pipeline

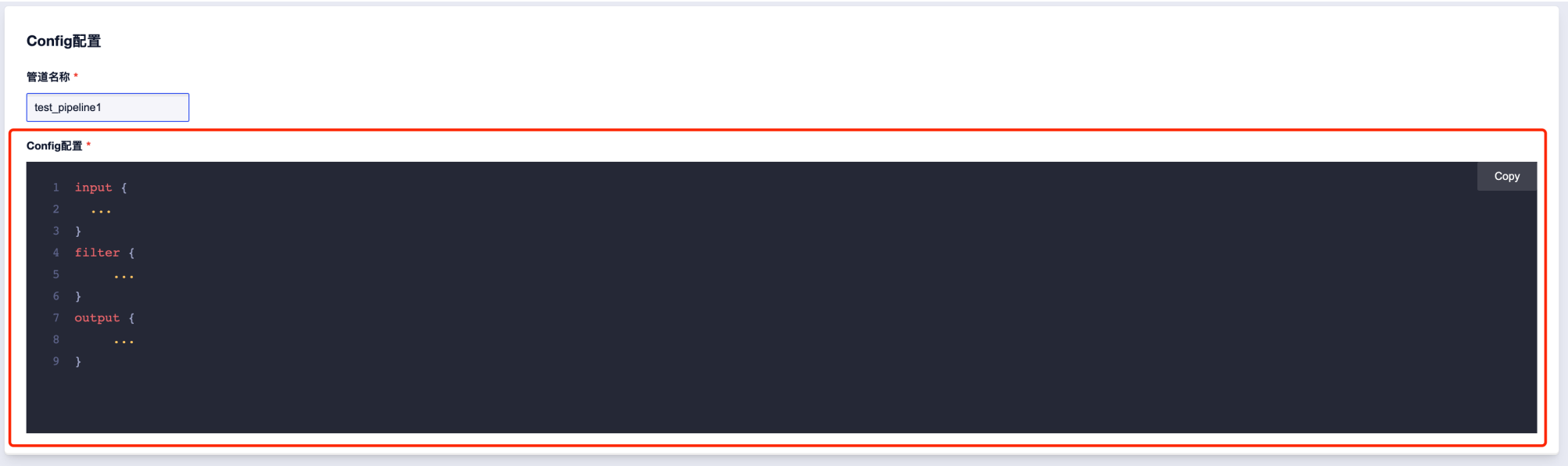
Config配置 *

```
1 input {  
2   ...  
3 }  
4 filter {  
5   ...  
6 }  
7 output {  
8   ...  
9 }
```

Copy

4. Config配置

用户根据任务内容进行修改。



参数说明 | 参数 | 说明 | | ----- | ----- | | input | 输入数据源配置。Logstash 支持的输入数据源类型,可参考 [Input plugins](#) | | filter | 对数据进行过滤或者预处理的配置。Logstash 支持的 filter 插件类型,可参考 [Filter plugins](#) | | output | 输出数据源配置。Logstash 支持的输出数据源类型,可参考 [Output plugins](#) |

管道任务配置详情可参考 [配置文件结构](#)。

5. 管道参数配置

参数	说明	默认值
管道名称	pipeline.id, 管道的唯一标识	-
管道工作线程	pipeline.workers, 管道的工作线程数量,也是并行执行管道的 filter 和 output 的工作线程数量	实例单节点的 CPU 核数
管道批处理大小	pipeline.batch.size, 每个批次处理的最大事件数量	125
管道批处理延迟	pipeline.batch.delay, 当管道批处理大小不满足时,每个批次最大的等待时间,单位为毫秒	50ms

队列类型	queue.type,用于事件缓冲的排队模型,可选值为 memory (基于内存的内列)。暂不支持修改	memory
队列最大字节数	queue.max_bytes,当选择 persisted 队列类型时,队列中可存放的最大字节数量,需确保该值小于实例单节点的磁盘容量。暂不支持修改	1024MB
队列检查点写入数	queue.checkpoint.acks,当选择 persisted 队列类型时,在强制执行检查点时已写入的最大的事件数量,若设置为0,则表示无限制。暂不支持修改	1024

修改管道

管道修改后需要进行重启才能生效

1. 在管道列表中,点击要修改的管道右侧的修改按钮,进入管道修改页面。
2. 在管道修改页面修改管道的任务配置和参数配置。
3. 点击保存,完成管道修改,修改后的内容将在管道下次启动时生效。

删除管道

注:

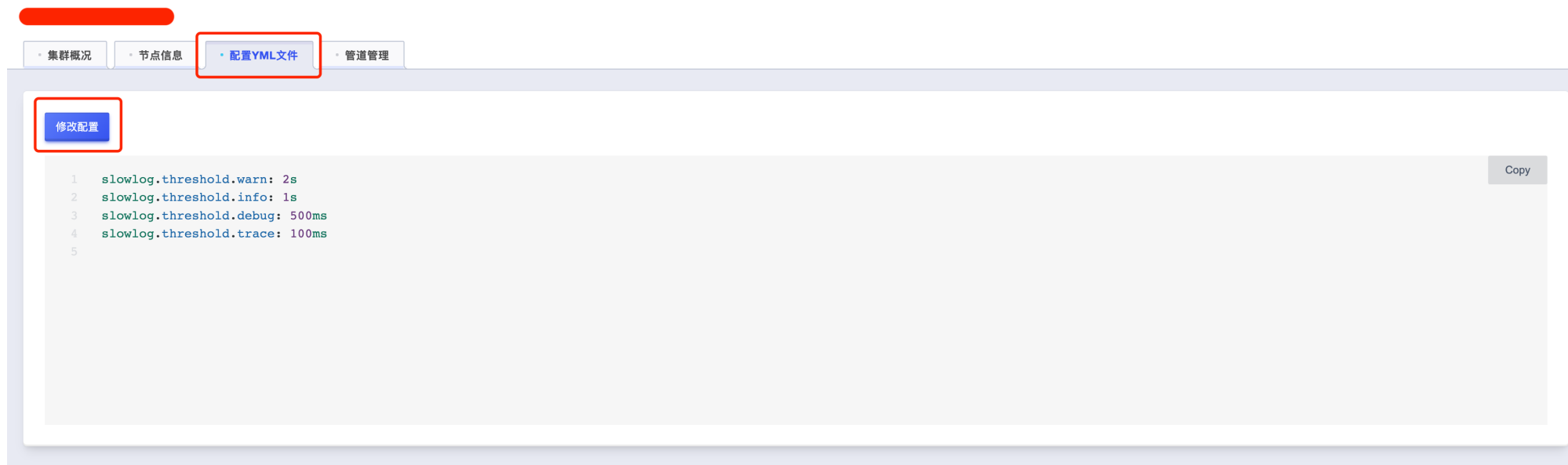
- 管道删除后无法恢复
- 运行中的管道需要先停止后才能执行删除

1. 在管道列表中,找到需要删除的管道,在右侧操作列表中点击删除。

YML 文件配置

操作步骤

1. 登录 ULogstash 控制台,在上侧导航栏单击 Logstash 管理,进入 Logstash 管理列表页。
2. 选择要修改 YML 参数配置的实例,单击 ID/名称,进入实例基本信息页。
3. 在实例详情信息页面,切换到配置YML文件签,单击修改配置,根据业务需求修改 YML 参数。详细参数说明,可参见 Logstash 配置文件。



4. YML 参数配置完成后,单击保存。修改 YML 参数需要重启管道任务才能生效。

版本管理

UES 提供了不同的产品版本,根据官网更新会定时发布新版本。用户可根据实际需求和习惯来选择要使用的版本,建议使用最新版本。

现在发布上线的版本有:

*** V5.5.1 * V6.2.1 * V6.5.4 * V7.1.1**

版本特性

*** V6.0版本亮点**

使用序列号更快地重启和还原

6.0 版本中最大的一个新特性就是序列 ID,它允许基于操作的分片恢复。以前,如果由于网络问题或节点重启而从集群断开连接的节点,则节点上的每个分区都必须通过将分段文件与主分片进行比较并复制任何不同的分段来重新同步。这可能是一个漫长而昂贵的过程,甚至使节点的滚动重新启动非常缓慢。使用序列 ID,每个分片将只能重放该分片中缺少的操作,使恢复过程更加高效。

使用排序索引更快查询

通过索引排序,只要收集到足够的命中,搜索就可以终止。它对通常用作过滤器的低基数字段(例如 age, gender, is_published)进行排序时可以更高效的搜索,因为所有潜在的匹配文档都被分组在一起。

稀疏区域改进

以前,每个列中的每个字段都预留了一个存储空间。如果只有少数文档出现很多字段,则可能会导致磁盘空间的巨大浪费。现在,你付出你使用的东西。密集字段将使用与以前相同的空间量,但稀疏字段将显着减小。这不仅可以减少磁盘空间使用量,还可以减少合并时间并提高查询吞吐量,因为可以更好地利用文件系统缓存。

类型去除

6.0 版本中一个 index 不再支持多个 type ,基于 type 的 parent-child 关系将通过单独的 join 字段来实现,并且 type 会在 7.0 版本彻底移除。

索引版本继承

6.0 版本为了避免索引模板所有匹配的合并而造成冲突,将会只匹配一个,索引创建时进行验证。

基于负载的请求路由

以前,搜索请求是全节点轮询,性能慢的节点往往会造成整体的延迟。6.0 版本实现方式将基于队列的耗费时间自动调节队列长度,搜索和索引都将基于这种机制。

FAQs

开发时连接集群失败的原因?

1. 开发与集群网络是否相通; 2. 集群名称是否配置正确, 配置文件cluster.name默认为集群** 资源ID **。

UES支持跨地域、跨机房访问吗?

UElasticsearch不支持跨地域用户网访问和不同VPC下的互访, 支持相同地域不同机房的访问。

能不能从相同可用区的云主机通过SSH登录访问集群节点?

不能, 集群节点不支持ssh登录。

外网怎么访问UElasticsearch集群?

暂不提供外网访问集群功能。

ULB产品支持UElasticsearch集群负载均衡访问吗？

暂不提供ULB访问。ULB七层协议只提供了外网访问，四层协议内网暂不支持UElasticsearch产品转发。

UES内置的head插件为何连接集群失败？

必须保证head插件管理平台的网络可以访问输入的es连接地址。强烈不建议为了使用head插件而搭建外网代理访问UElasticsearch集群。

什么时候选用主节点分离类型的集群？

节点个数大于5时建议选用主节点分离类型集群（集群包含独立的主节点）。主节点选择默认最小配置即可，当节点个数大于20时可考虑高一级的配置。

集群是否支持扩容？节点是否支持修改配置？主节点是否支持扩容？

集群支持增加节点，暂不支持控制台删除节点；节点支持升级配置，暂不支持降低配置；主节点不支持扩容，不支持升降配置。集群扩容节点默认配置和原数据节点配置相同，不支持选择其他配置。

SSD磁盘的几种类型节点有什么区别？

SSD磁盘的节点类型有：内存密集机型、普通机型、存储密集机型。三个大类上没有本质的区别，只是为了方便管理分成了三类，可以根据实际的业务需求，来选择合适的CPU、内存和存储容

量。

索引主分片数量和副本设置几个最合理?

UElasticsearch集群默认索引5个主分片和1个副本。建议的主分片个数 $\geq$ 数据节点个数,但不要超过数据节点个数的三倍,要求副本设置小于数据节点个数。

UES是否支持工具扩展,如Logstash的集成?

支持ELK的集成使用,UElasticsearch内部已集成有Kibana。Logstash的版本没有特殊要求,但是建议选用一样的版本。

Rally压测

我们利用 Elasticsearch 官方提供的 benchmark 压测工具 [rally](https://github.com/elastic/rally) 对华北一可用区B的Elasticsearch集群 (V5.5.1) 进行测试, 测试结果如下:

8核16G 3个节点

Metric	Operation	Value	Unit
Indexing time		31.7303	min
Merge time		9.5995	min
Refresh time		1.93053	min
Flush time		0.0881167	min
Merge throttle time		1.54493	min
Total Young Gen GC		58.079	s
Total Old Gen GC		0.196	s
Heap used for segments		18.3168	MB
Heap used for doc values		0.103638	MB
Heap used for terms		17.1259	MB
Heap used for norms		0.0717773	MB
Heap used for points		0.217482	MB

Heap used for stored fields		0.797966	MB
Segment count		94	
Min Throughput	index-stats	100.045	ops/s
Median Throughput	index-stats	100.068	ops/s
Max Throughput	index-stats	100.125	ops/s
50th percentile latency	index-stats	3.61395	ms
90th percentile latency	index-stats	4.04719	ms
99th percentile latency	index-stats	4.96298	ms
99.9th percentile latency	index-stats	16.8928	ms
100th percentile latency	index-stats	23.5338	ms
Min Throughput	node-stats	100.031	ops/s
Median Throughput	node-stats	100.09	ops/s
Max Throughput	node-stats	100.506	ops/s
50th percentile latency	node-stats	4.48539	ms
90th percentile latency	node-stats	5.0753	ms
99th percentile latency	node-stats	7.10029	ms
99.9th percentile latency	node-stats	17.5065	ms

100th percentile latency	node-stats	23.1605	ms
Min Throughput	default	49.9519	ops/s
Median Throughput	default	50.013	ops/s
Max Throughput	default	50.0216	ops/s
50th percentile latency	default	15.2663	ms
90th percentile latency	default	16.0737	ms
99th percentile latency	default	18.9459	ms
99.9th percentile latency	default	40.3196	ms
100th percentile latency	default	45.0553	ms
Min Throughput	term	200.063	ops/s
Median Throughput	term	200.091	ops/s
Max Throughput	term	200.174	ops/s
50th percentile latency	term	2.59966	ms
90th percentile latency	term	3.12504	ms
99th percentile latency	term	31.4645	ms
99.9th percentile latency	term	43.7732	ms
100th percentile latency	term	45.5124	ms
Min Throughput	phrase	200.03	ops/s

Median Throughput	phrase	200.049	ops/s
Max Throughput	phrase	200.075	ops/s
50th percentile latency	phrase	3.88241	ms
90th percentile latency	phrase	5.3865	ms
99th percentile latency	phrase	44.4824	ms
99.9th percentile latency	phrase	54.733	ms
100th percentile latency	phrase	55.2758	ms
Min Throughput	country_agg_uncached	2.75134	ops/s
Median Throughput	country_agg_uncached	2.76574	ops/s
Max Throughput	country_agg_uncached	2.8003	ops/s
50th percentile latency	country_agg_uncached	162393	ms
90th percentile latency	country_agg_uncached	222884	ms
99th percentile latency	country_agg_uncached	234288	ms
99.9th percentile latency	country_agg_uncached	235730	ms
100th percentile latency	country_agg_uncached	235940	ms
Min Throughput	country_agg_cached	100.051	ops/s
Median Throughput	country_agg_cached	100.073	ops/s
Max Throughput	country_agg_cached	100.138	ops/s

50th percentile latency	country_agg_cached	2.87947	ms
90th percentile latency	country_agg_cached	3.31356	ms
99th percentile latency	country_agg_cached	6.88291	ms
99.9th percentile latency	country_agg_cached	30.7471	ms
100th percentile latency	country_agg_cached	37.4828	ms
Min Throughput	scroll	35.6468	ops/s
Median Throughput	scroll	35.81	ops/s
Max Throughput	scroll	35.8374	ops/s
50th percentile latency	scroll	296588	ms
90th percentile latency	scroll	427849	ms
99th percentile latency	scroll	457200	ms
100th percentile latency	scroll	460396	ms
Min Throughput	expression	1.28785	ops/s
Median Throughput	expression	1.31216	ops/s
Max Throughput	expression	1.31935	ops/s
50th percentile latency	expression	78301.4	ms
90th percentile latency	expression	99426.7	ms
99th percentile latency	expression	103425	ms

100th percentile latency	expression	103679	ms
--------------------------	------------	--------	----

2核8G 3个节点

Metric	Operation	Value	Unit
Indexing time		43.8383	min
Merge time		22.5664	min
Refresh time		9.2856	min
Flush time		0.0502833	min
Merge throttle time		2.30762	min
Total Young Gen GC		418.967	s
Total Old Gen GC		3.078	s
Heap used for segments		17.7271	MB
Heap used for doc values		0.102936	MB
Heap used for terms		16.5306	MB
Heap used for norms		0.071167	MB
Heap used for points		0.209009	MB
Heap used for stored fields		0.813393	MB
Segment count		96	

Min Throughput	index-stats	100.035	ops/s
Median Throughput	index-stats	100.064	ops/s
Max Throughput	index-stats	100.099	ops/s
50th percentile latency	index-stats	3.98357	ms
90th percentile latency	index-stats	5.42446	ms
99th percentile latency	index-stats	17.3357	ms
99.9th percentile latency	index-stats	38.0944	ms
100th percentile latency	index-stats	44.4198	ms
Min Throughput	node-stats	100.011	ops/s
Median Throughput	node-stats	100.089	ops/s
Max Throughput	node-stats	100.539	ops/s
50th percentile latency	node-stats	4.45035	ms
90th percentile latency	node-stats	6.89122	ms
99th percentile latency	node-stats	19.1262	ms
99.9th percentile latency	node-stats	39.2981	ms
100th percentile latency	node-stats	44.4663	ms
Min Throughput	default	25.5403	ops/s
Median Throughput	default	26.1976	ops/s

Max Throughput	default	28.0422	ops/s
50th percentile latency	default	18200.4	ms
90th percentile latency	default	26643.4	ms
99th percentile latency	default	28551	ms
99.9th percentile latency	default	28734.5	ms
100th percentile latency	default	28751.2	ms
Min Throughput	term	198.796	ops/s
Median Throughput	term	200.081	ops/s
Max Throughput	term	200.118	ops/s
50th percentile latency	term	2.69169	ms
90th percentile latency	term	15.5542	ms
99th percentile latency	term	52.1687	ms
99.9th percentile latency	term	62.2191	ms
100th percentile latency	term	63.3586	ms
Min Throughput	phrase	196.082	ops/s
Median Throughput	phrase	200.032	ops/s
Max Throughput	phrase	200.049	ops/s
50th percentile latency	phrase	4.70513	ms

90th percentile latency	phrase	41.5631	ms
99th percentile latency	phrase	61.2904	ms
99.9th percentile latency	phrase	64.1431	ms
100th percentile latency	phrase	64.8171	ms
Min Throughput	country_agg_uncached	2.53125	ops/s
Median Throughput	country_agg_uncached	2.55542	ops/s
Max Throughput	country_agg_uncached	2.56809	ops/s
50th percentile latency	country_agg_uncached	191703	ms
90th percentile latency	country_agg_uncached	265937	ms
99th percentile latency	country_agg_uncached	282614	ms
99.9th percentile latency	country_agg_uncached	284102	ms
100th percentile latency	country_agg_uncached	284297	ms
Min Throughput	country_agg_cached	100.046	ops/s
Median Throughput	country_agg_cached	100.072	ops/s
Max Throughput	country_agg_cached	100.134	ops/s
50th percentile latency	country_agg_cached	3.24472	ms
90th percentile latency	country_agg_cached	4.16849	ms
99th percentile latency	country_agg_cached	25.1214	ms

99.9th percentile latency	country_agg_cached	43.7322	ms
100th percentile latency	country_agg_cached	50.3143	ms
Min Throughput	scroll	31.3553	ops/s
Median Throughput	scroll	31.6427	ops/s
Max Throughput	scroll	31.7946	ops/s
50th percentile latency	scroll	337930	ms
90th percentile latency	scroll	485487	ms
99th percentile latency	scroll	518748	ms
100th percentile latency	scroll	522454	ms
Min Throughput	expression	1.09803	ops/s
Median Throughput	expression	1.11245	ops/s
Max Throughput	expression	1.11853	ops/s
50th percentile latency	expression	120302	ms
90th percentile latency	expression	150630	ms
99th percentile latency	expression	157380	ms
100th percentile latency	expression	158111	ms

4核8G 3个节点

Metric	Operation	Value	Unit
--------	-----------	-------	------

Indexing time		30.2766	min
Merge time		11.1005	min
Refresh time		3.1837	min
Flush time		0.0554	min
Merge throttle time		1.49132	min
Total Young Gen GC		87.16	s
Total Old Gen GC		2.738	s
Heap used for segments		17.9447	MB
Heap used for doc values		0.0973434	MB
Heap used for terms		16.7564	MB
Heap used for norms		0.0637817	MB
Heap used for points		0.220229	MB
Heap used for stored fields		0.806961	MB
Segment count		84	
Min Throughput	index-stats	99.9043	ops/s
Median Throughput	index-stats	100.055	ops/s
Max Throughput	index-stats	100.111	ops/s

50th percentile latency	index-stats	4.31123	ms
90th percentile latency	index-stats	5.01966	ms
99th percentile latency	index-stats	6.73182	ms
99.9th percentile latency	index-stats	16.3152	ms
100th percentile latency	index-stats	16.7211	ms
Min Throughput	node-stats	99.9997	ops/s
Median Throughput	node-stats	100.068	ops/s
Max Throughput	node-stats	100.441	ops/s
50th percentile latency	node-stats	5.41895	ms
90th percentile latency	node-stats	6.62956	ms
99th percentile latency	node-stats	10.4876	ms
99.9th percentile latency	node-stats	16.8256	ms
100th percentile latency	node-stats	22.3502	ms
Min Throughput	default	49.9808	ops/s
Median Throughput	default	50.016	ops/s
Max Throughput	default	50.0348	ops/s
50th percentile latency	default	14.0526	ms
90th percentile latency	default	16.2503	ms

99th percentile latency	default	21.8674	ms
99.9th percentile latency	default	29.5538	ms
100th percentile latency	default	31.5155	ms
Min Throughput	term	199.745	ops/s
Median Throughput	term	200.06	ops/s
Max Throughput	term	200.087	ops/s
50th percentile latency	term	3.06953	ms
90th percentile latency	term	4.04719	ms
99th percentile latency	term	24.0726	ms
99.9th percentile latency	term	37.8865	ms
100th percentile latency	term	38.5456	ms
Min Throughput	phrase	199.35	ops/s
Median Throughput	phrase	199.992	ops/s
Max Throughput	phrase	200.071	ops/s
50th percentile latency	phrase	4.59871	ms
90th percentile latency	phrase	22.3203	ms
99th percentile latency	phrase	48.3212	ms
99.9th percentile latency	phrase	52.8152	ms

100th percentile latency	phrase	52.8472	ms
Min Throughput	country_agg_uncached	4.49216	ops/s
Median Throughput	country_agg_uncached	4.67822	ops/s
Max Throughput	country_agg_uncached	4.7315	ops/s
50th percentile latency	country_agg_uncached	14575.8	ms
90th percentile latency	country_agg_uncached	17213.5	ms
99th percentile latency	country_agg_uncached	18765.9	ms
99.9th percentile latency	country_agg_uncached	18810.1	ms
100th percentile latency	country_agg_uncached	18813.6	ms
Min Throughput	country_agg_cached	100.046	ops/s
Median Throughput	country_agg_cached	100.06	ops/s
Max Throughput	country_agg_cached	100.129	ops/s
50th percentile latency	country_agg_cached	3.49425	ms
90th percentile latency	country_agg_cached	4.68459	ms
99th percentile latency	country_agg_cached	6.65144	ms
99.9th percentile latency	country_agg_cached	17.7776	ms
100th percentile latency	country_agg_cached	24.6469	ms
Min Throughput	scroll	32.6045	ops/s

Median Throughput	scroll	32.725	ops/s
Max Throughput	scroll	32.7822	ops/s
50th percentile latency	scroll	326129	ms
90th percentile latency	scroll	470066	ms
99th percentile latency	scroll	502268	ms
100th percentile latency	scroll	505901	ms
Min Throughput	expression	1.89188	ops/s
Median Throughput	expression	1.90472	ops/s
Max Throughput	expression	1.94701	ops/s
50th percentile latency	expression	7336.16	ms
90th percentile latency	expression	8155.22	ms
99th percentile latency	expression	8272.59	ms
100th percentile latency	expression	8322.61	ms

UES价格

UES产品包含两种计费模式：

1. 按配置计费:按照CPU、内存、存储3个维度进行计费；
2. 按节点计费:按照节点维度根据节点配置进行整体计费；

其中按节点计费方式仅提供部分存量实例使用,新实例均采用按配置计费方式。

说明:文档中的价格以华北一可用区E为例,其他可用区价格可查看UES控制台。

1. 按配置计费

产品价格按照CPU、内存、存储3个维度进行计费,详细价格如下：

1.1 CPU与内存

计费项	华北一可用区E 月单价
CPU(元/核/月)	82
内存(元/GB/月)	22

常用机型(不含存储)价格如下：

CPU内存比	名称	配置	华北一可用区E 月单价（元/月）
1:2	o.es2m.medium	2核 4G	252
1:2	o.es2m.xlarge	4核 8G	504
1:2	o.es2m.2xlarge	8核 16G	1008
1:2	o.es2m.4xlarge	16核 32G	2016
1:2	o.es2m.8xlarge	32核64G	4032
1:4	o.es4m.medium	2核 8G	340
1:4	o.es4m.xlarge	4核 16G	680
1:4	o.es4m.2xlarge	8核 32G	1360
1:4	o.es4m.4xlarge	16核 64G	2720

1.2 存储

计费项	华北一可用区E 月单价（元/GB/月）
存储	0.75

2. 按节点计费

产品价格根据节点类型及配置不同, 详细价格如下：

机型	名称	配置	华北一可用区E 月单价（元/月）
RSSD云盘	R1ES-large	2核 8G 200G(SSD)	405
RSSD云盘	R1ES-xlarge	4核 16G 400G(SSD)	810
RSSD云盘	R1ES-2xlarge	8核 32G 800G(SSD)	1620
RSSD云盘	R1ES-4xlarge	16核 64G 1600G(SSD)	3240
RSSD云盘	R2ES-large	2核 8G 400G(SSD)	535
RSSD云盘	R2ES-xlarge	4核 16G 800G(SSD)	1070
RSSD云盘	R2ES-2xlarge	8核 32G 1600G(SSD)	2140
RSSD云盘	R2ES-4xlarge	16核 64G 3200(SSD)	4280
内存密集机型	M1ES-large	2核 8G 150G(SSD)	450
内存密集机型	M1ES-xlarge	4核 16G 300G(SSD)	900
内存密集机型	M1ES-2xlarge	8核 32G 600G(SSD)	1800
内存密集机型	M1ES-4xlarge	16核 64G 1200G(SSD)	3600
普通机型	C1ES-large	4核 8G 300G(SSD)	700
普通机型	C1ES-xlarge	8核 16G 600G(SSD)	1400
普通机型	C1ES-2xlarge	16核 32G 1200G(SSD)	2800
普通机型	C1ES-4xlarge	32核 64G 1200G(SSD)	4700
存储密集机型	S1ES-xlarge	16核 64G 2000G(SSD)	4500

SATA盘机型	STES-large	2核 8G 500G(SATA)	500
SATA盘机型	STES-xlarge	4核 16G 1000G(SATA)	950
SATA盘机型	STES-2xlarge	8核 32G 2000G(SATA)	1900
SATA盘机型	STES-4xlarge	16核 64G 4000G(SATA)	3800