

容器服务

常见问题

产品文档



腾讯云

【版权声明】

©2013-2024 腾讯云版权所有

本文档著作权归腾讯云单独所有，未经腾讯云事先书面许可，任何主体不得以任何形式复制、修改、抄袭、传播全部或部分本文档内容。

【商标声明】

及其它腾讯云服务相关的商标均为腾讯云计算（北京）有限责任公司及其关联公司所有。本文档涉及的第三方主体的商标，依法由权利人所有。

【服务声明】

本文档意在向客户介绍腾讯云全部或部分产品、服务的当时的整体概况，部分产品、服务的内容可能有所调整。您所购买的腾讯云产品、服务的种类、服务标准等应由您与腾讯云之间的商业合同约定，除非双方另有约定，否则，腾讯云对本文档内容不做任何明示或默示的承诺或保证。

文档目录

常见问题

运维类

基础监控相关

日志采集相关

TKE 标准集群

扩容缩容相关

集群相关

如何选择 Containerd 和 Docker

自建 DNS 导致节点初始化异常

服务类

TKE Serverless 集群

TKE Serverless 集群相关

负载均衡相关

超级节点相关

镜像仓库相关

镜像仓库类

远程终端类

事件类

常见问题

运维类

基础监控相关

最近更新时间：2022-10-12 16:05:10

基础监控常见问题

节点 cpu/memory 分配量为什么会超出节点资源规格？

原因：node 层级的 cpu/memory 分配量指标依赖节点上各个 pod 的 cpu/memory request 来计算，在计算时没有把 failed 的 pod 排除。

示例：节点规格是 4c8g，节点上目前运行 3 个 pod（资源 request 用量如下）：

- pod1 正常运行其 request 为 2c4g；
- pod2 正常运行其 request 为 1c2g；
- pod3 状态为 failed 其 request 为 0.5c1g；

此时节点剩余可调度资源为 4-2-1=1c、8-4-2=2g，pod4 request 为 0.8c1.5g，满足调度器筛选，正常被调度到该节点上。此时节点上共 4 个 pod，3 个正常 1 个异常，此时 node 层级的分配量为 4.3c8.5g（因计算时没有把 failed 的 pod 排除，因此超过了节点规格）。

该问题已在 5 月新版本中修复，即计算 node 资源分配量已把异常 pod 排除。

为什么 pod 状态显示正常，但监控指标 k8s_workload_abnormal 展示异常？

原因：该指标根据 workload 下 pod 是否异常来判断，pod 是否异常取决于 `pod.status.condition` 下这四个 Type 来确定。当这四个指标同时为 `True` 时 `k8s_workload_abnormal` 才会认为是正常，否则认为是异常。

- PodScheduled：Pod 已经被调度到某节点。
- ContainersReady：Pod 中所有容器都已就绪。
- Initialized：所有的 Init 容器 都已成功完成。
- Ready：Pod 可以为请求提供服务，并且应该被添加到对应服务的负载均衡池中。

daemonSet tke-monitor-agent 报错原因总结

现象	原因	解决措施
----	----	------

现象	原因	解决措施
无法解析域名 `receiver.barad.tencentyun.com`，指标上报失败，导致用户集群没有监控数据	节点 dns 被修改	给 DaemonSet tke-monitor-agent 加上 hostAlias。参考代码如下： <pre>hostAliases: - hostnames: - receiver.barad.tencentyun.com ip: 169.254.0.4</pre>

日志采集相关

最近更新时间：2023-03-27 11:08:16

日志采集通用问题

集群配置日志采集后，为什么在日志服务控制台查看不到日志？

发生日志查看不到或者缺失的情况，请检查是否存在以下问题：

检查所选的日志 topic 是否开启了索引。索引配置是使用日志服务进行检索分析的必要条件。若未开启，则无法查看日志。配置索引的详细操作，请参见 [日志服务配置索引](#)。



日志、审计、事件是否选用同一个 topic。若选用同一个 topic 则会发生日志覆盖，造成缺失。

检查所选的 topic 是否选用了两种提取模式。新的提取模式会覆盖旧的提取模式。

检查是否存在软链接。如果选择采集类型为“容器文件路径”时，对应的“容器文件路径”不能为软链接，否则会导致软链接的实际路径在采集器的容器内不存在，导致采集日志失败。

若采用环境变量的方式开启 TKE Serverless 日志采集，并选择角色授权，在**创建角色时需要选择 CVM 载体**，选择**容器服务**则无法完成授权。

如果您的问题仍未解决，请 [提交工单](#) 联系我们。

配置好日志规则后，日志在哪查看？

1. 登录 [日志服务控制台](#)，选择左侧**检索分析**。
2. 进入“检索分析”页面，选择地域、需要查看日志的日志集和日志主题，开启全文索引，即可检索分析日志。如下图所示：



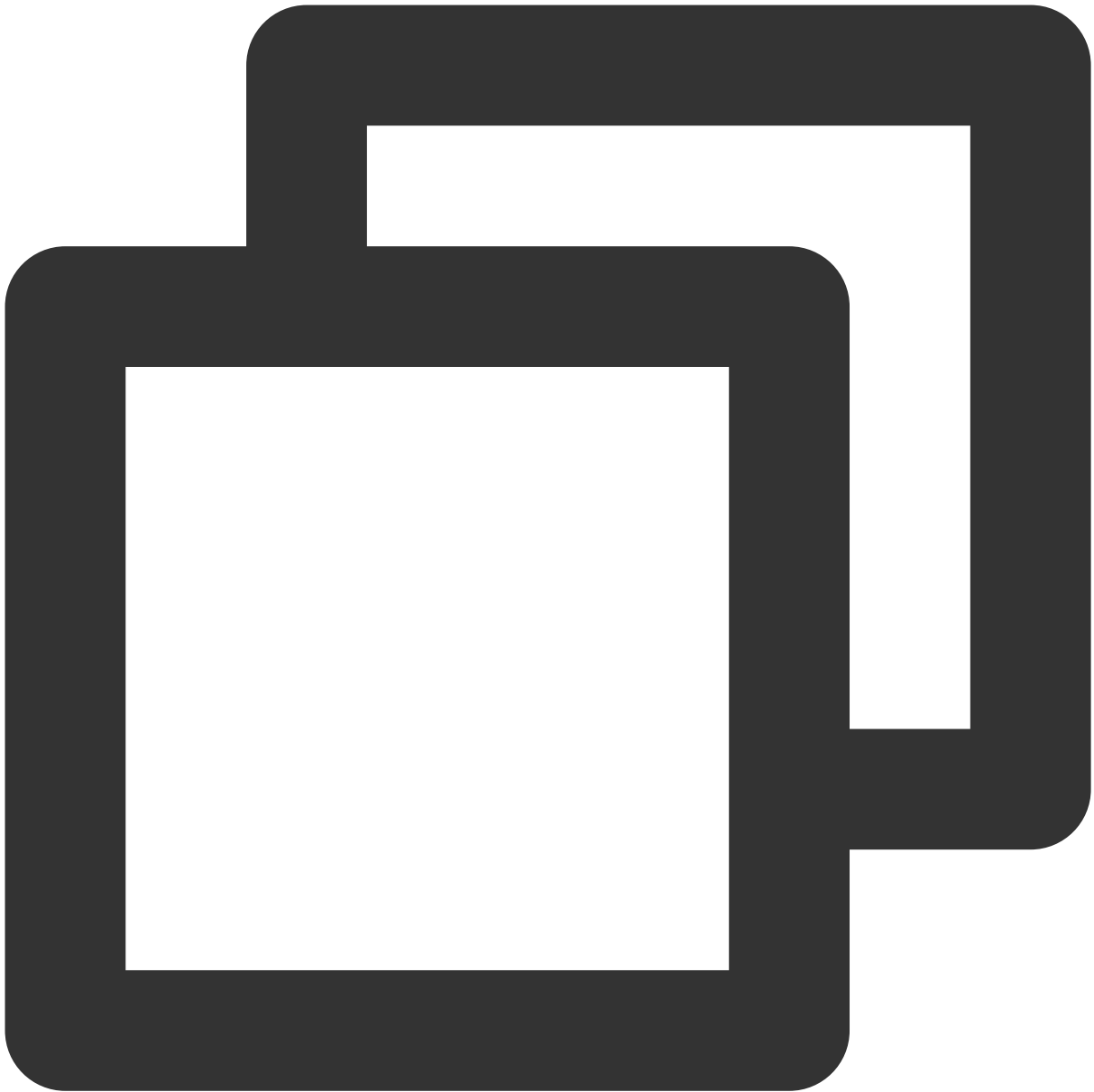
使用环境变量开启 TKE Serverless 日志采集时

Java 类应用如何实现多行日志合并？

当客户程序的日志数据跨占多行时（例如 Java 程序日志），不能使用换行符 `\\n` 作为日志的结束标识符。为了让日志系统明确区分开每条日志，需要配置具有首行正则表达式的 `configmap`，当某行日志匹配了预先设置的正则表达式，则认为是一条日志的开头，而在下一个行首出现则作为该条日志的结束标识符。详情可参见 [TKE Serverless 日志采集实现多行日志合并](#)。

如何调整日志采集配置，以适应不同的日志输出速率？

当使用环境变量的方式开启 TKE Serverless 日志采集时，TKE Serverless 会在 `pod sandbox` 里额外启动一个日志采集组件来进行日志收集并上报。由于 TKE Serverless 限制了日志采集组件的内存使用量，当日志的输出速率过高时，日志采集组件可能会发生 OOM 的情况，此时可以按需调整日志采集的配置。具体方法为手动修改以下 `pod` annotation，将日志采集组件使用的内存缓冲区调小，以减少内存占用。



```
internal.eks.tke.cloud.tencent.com/tail-buffer-chunk-size: "2M"
internal.eks.tke.cloud.tencent.com/tail-buffer-max-size: "2M"
```

annotation 的含义如下所示，详情可参见 [Fluent Bit](#)。

参数	含义	默认值
Buffer_Chunk_Size	设置初始缓冲区大小以采集容器文件日志，也可用于增加缓冲区的大小。	2M
Buffer_Max_Size	设置每个受监视文件的缓冲区大小的限制。当需要增加缓冲区时（例如长	2M

日志），此值可以限制缓冲区的增加量。如果读取的文件超过此限制，将从监视列表中删除此文件。
--

Serverless 集群容器内应用输出日志的标准是什么？

在应用输出日志时，应注意容器内应用尽量输出至 `stdout`。如果您的程序日志输出至容器的文件，需要有相应的手段来管理日志文件的滚动和清理，或者挂载 `volume` 持久化存储，否则20G临时空间会被占满，详情可参见 [Pod 临时存储](#)。

若无管理日志文件的滚动和清理手段，推荐您进行以下操作：

挂载 `volume` 持久化存储，详情可参见 [存储管理](#)。

开启 TKE Serverless 日志采集，详情可参见 [使用环境变量配置日志采集](#)。

TKE 标准集群

扩容缩容相关

最近更新时间：2020-10-10 11:29:58

Cluster Autoscaler 与基于监控指标的弹性伸缩的节点扩缩容有什么不同？

Cluster Autoscaler 确保集群中的所有 Pod 都可调度，不管具体的负载；而基于监控指标的节点弹性伸缩在自动扩缩时不关心 Pod，可能会添加一个没有任何 Pod 的节点，或者删除一个有一些系统关键 Pod 的节点，例如 kube-dns。Kubernetes 不鼓励这种自动缩容机制，故 Cluster Autoscaler 与基于监控指标的弹性伸缩的节点互相冲突，请不要同时启用。

CA 和伸缩组的对应关系是什么？

启用 CA 的集群会根据选择的节点配置，创建一个启动配置和绑定此启动配置的伸缩组。绑定后，将会在此伸缩组内进行扩缩容，扩容后的 CVM 自动加入集群。自动扩缩容的节点都是按量计费的。伸缩组的相关文档请参见 [弹性伸缩文档](#)。

容器服务控制台手动添加的节点是否会 CA 缩容？

不会，CA 缩容的节点只限于伸缩组内的节点。在 [容器服务控制台](#) 添加的节点不会加入到伸缩组中。

弹性伸缩控制台是否可以添加或者移出云服务器？

不可以，不建议您在 [弹性伸缩控制台](#) 进行任何修改操作。

扩缩容会继承所选节点的哪些配置？

创建伸缩组时，需要选择集群内的一个节点作为参考来创建 [启动配置](#)，参考的节点配置包括：

- vCPU
- 内存
- 系统盘大小
- 数据盘大小
- 磁盘类型
- 带宽
- 带宽计费模式
- 是否分配公网IP
- 安全组
- 私有网络
- 子网

如何使用多个伸缩组？

根据服务的重要级别、类型等特点，您可以通过创建多个伸缩组，为伸缩组设置不同的 label，从而指定伸缩组扩容出节点的 label，来对服务进行分类。

扩缩容最大值可以设置为多少？

目前腾讯云用户每个可用区均有30个按量计费类型 CVM 配额，如果希望伸缩组有超过 30 台按量计费的 CVM，请[提交工单](#)申请。

具体配额请参见您当前可用区的云服务器 [实例数及配额](#)。另外弹性伸缩也有最大值的限制，其最大值为200。如果希望弹性伸缩超过最大值，请[提交工单](#)申请。

集群启用缩容是否安全？

由于在缩容节点时会发生 Pod 重新调度的情况，所以服务必须可以容忍重新调度和短时的中断后再启用缩容。建议您为您的服务设置 [PDB](#)。PDB 可以在任何时候指定一个处于运行状态的 Pod 集合副本的最小数量或者最小百分比。有了 PodDisruptionBudget，应用部署者可以保证同一时间内主动移除 Pod 的集群操作不会销毁过多 Pod，避免了因销毁过多 Pod 导致数据丢失、服务中断或者无法接受的服务降级等影响。

节点上有哪些类型的 Pod 时不会被缩容？

- 当您设置了严格的 PodDisruptionBudget 的 Pod 不满足 PDB 时，不会缩容。
- Kube-system 下的 Pod。
- 节点上有非 deployment, replica set, job, stateful set 等控制器创建的 Pod。
- Pod 有本地存储。
- Pod 不能被调度到其他节点上。

节点满足缩容条件后多长时间会触发缩容？

10分钟。

节点 Not Ready 后多长时间会触发缩容？

20分钟。

多长时间扫描一次是否需要扩缩容？

10秒。

需要多长时间才可以扩容出 CVM？

一般在10分钟内，相关弹性伸缩的说明文档请参见 [弹性伸缩](#)。

为什么有 Unschedulable 的 Pod，却未进行扩容？

请确认以下原因：

- Pod 的请求资源是否过大。
- 是否设置了 node selector。

- 伸缩组的最大值是否已经达到。
- 账号余额是否充足（账号余额不足，弹性伸缩无法扩容），以及配额不足等其他原因，请参见 [弹性伸缩故障处理](#)。

如何防止 Cluster Autoscaler 缩容特定节点？

```
# 可以在节点的annotations中设置如下信息
kubectl annotate node <nodename> cluster-autoscaler.kubernetes.io/scale-down-disabled=true
```

扩缩容事件如何反馈给用户？

用户可在弹性伸缩控制台查询伸缩组的伸缩活动，也可查看 k8s 的事件。在以下三种资源上都会有对应的事件：

- kube-system/cluster-autoscaler-status config map
 - **ScaledUpGroup** - CA 触发扩容。
 - **ScaleDownEmpty** - CA 删除了一个没有运行 Pod 的节点。
 - **ScaleDown** - CA 缩容。
- node
 - **ScaleDown** - CA 缩容。
 - **ScaleDownFailed** - CA 缩容失败。
- pod
 - **TriggeredScaleUp** - CA 由于此 Pod 触发扩容。
 - **NotTriggerScaleUp** - CA 无法找到可扩容的伸缩组使得此 Pod 可调度。
 - **ScaleDown** - CA 尝试驱逐此 Pod 来缩容节点。

集群相关

最近更新时间：2023-05-26 10:13:28

创建集群常见问题

创建集群时，云服务器可以不选取公网 IP 么？

云服务器可以不选取公网 IP，无公网IP的云服务器只能拉取镜像仓库下我的镜像，不能拉取 dockerhub 以及第三方镜像。

无公网 IP，但有外网带宽的云服务器可以通过绑定弹性 IP 来访问 Internet。

创建集群时，选择所属网络的作用是什么？

选择的所属网络和子网，是集群内云服务器的所在子网，用户可以通过添加不同的云服务器到不同可用区的子网下进行跨可用区容灾。

创建集群支持什么类型的机型？

支持 CVM 提供的标准型、计算型、高 IO 型、GPU 型、黑石型中等机器，具体以 TKE 控制台展示的机型为主。

当前容器服务宿主机支持什么操作系统？

当前支持操作系统详情见 TKE 支持的 [公共镜像列表](#)。

容器服务如何设置节点的自定义 Kubelet 参数？

- 该功能白名单形式开放，您可以通过 [提交工单](#) 申请使用。
- 在 [新增节点](#) 页面、[添加已有节点](#) 页面、[新增节点池](#) 页面设置节点的自定义 Kubelet 参数。如下图所示：

容器服务创建集群如何设置自定义 Kubernetes 组件参数？

具体操作步骤可参考 [自定义 Kubernetes 组件启动参数](#)。

容器服务支持独立集群部署模式吗？

TKE 为您提供集群完全自主可控的 Master 独立部署模式。该模式下，Kubernetes 集群的 Master 和 Etcd 会部署在您购置的 CVM 上，并且您拥有 Kubernetes 集群的所有管理和操作权限。详情请参考 [Master 独立部署模式](#)。

TKE Serverless 集群创建后支持修改可用区吗？

容器服务 TKE 的 Serverless 集群创建成功后，是不支持更换或者在添加其他可用区的。

TKE 集群支持获取客户端 ipv6 地址吗？

目前 TKE 集群不支持获取客户端 ipv6 地址。

TKE 集群可以导出配置吗？

TKE 容器服务器不支持将您现在的集群进行导出。

TKE 集群创建提示需完成节点异常检测 plus 参数设置

创建集群提示需要完成节点异常检测 plus 参数设置，需要您检查下是否有勾选 NodeProblemDetectorPlus 组件，如已勾选，则需要配置相关参数。组件详情见 [NodeProblemDetectorPlus 说明](#)。

说明：

如暂无法评估是否需要使用此组件，可以取消勾选，直接创建集群。后续如需使用相关功能，在组件管理中再次操作安装即可。

TKE 集群如何启用 IPVS？

具体操作可参考 [集群启用 IPVS](#)。

容器服务如何查看集群访问凭证？

1. 登录 [容器服务控制台](#)，选择左侧导航栏中的**集群**，进入集群管理界面。
2. 单击需要连接的集群 ID/名称，进入集群详情页。
3. 选择左侧导航栏中的**基本信息**，即可在“基本信息”页面中查看“集群 API Server 信息”模块中该集群的访问地址、外网/内网访问状态、Kubeconfig 访问凭证内容等信息。

容器服务怎么才能添加其他账号名下的服务器到集群？

目前不支持将其他账号下的云服务器添加到集群，当前仅支持添加同一 VPC 下的云服务器。

使用 K8S 官方 SDK 连接集群 API Server 报错 certificate verify failed: self signed certificate 如何处理？

K8S 集群的 apiserver 所使用的 tls 证书一般都是自签名证书，在 python 的低版本 request 库和 nodejs 的默认设置中，不信任自签名的证书，因此会报错，可通过如下方案解决：

1. 新增跳过服务端校验的配置。

若使用 nodejs sdk，需设置环境变量：

```
export NODE_TLS_REJECT_UNAUTHORIZED="0" ;
```

若使用 python sdk，需在 kubeconfig 里设置

```
clusters:
- cluster:
  insecure-skip-tls-verify: true
```

2. 将自签名证书颁发机构根证书的公钥放到系统的受信根证书列表中。

集群网络常见问题

自定义 Webhook 开发时，针对集群网络需要注意什么？

在进行自定义 Webhook 开发时，不要拦截 kube-system namespace 下 Pod，否则会导致集群网络使用异常。

扩展云服务器常见问题

扩展云服务器有什么限制？

只能选择当前集群所在的地域，但可以选择不同的可用区，允许集群跨可用区部署。

云服务器的数量有限制么？

有，用户所有按量计费云服务器不能超过用户配额，具体详情请看 [云服务概览页](#)。

销毁云服务器常见问题

销毁云服务器后，该主机下部署的容器怎么办？

销毁云服务器时，该主机下的容器等资源也会随之销毁。若某服务的容器数量小于期望运行的容器数量时，集群将在其他主机上启动容器，直到运行的容器数量等于期望运行容器数量为止。

如何选择 Containerd 和 Docker

最近更新时间：2023-08-10 17:22:55

如何选择 Containerd 和 Docker

如何选择运行时组件？

注意：

Kubernetes 1.24通过 Dockershim 对 Docker 的支持已移除，新建节点的容器运行时请使用 Containerd，通过 Docker 构建的镜像可以继续使用。

Containerd 是更为稳定的运行时组件，支持 OCI 标准，但不支持 Docker API。

容器运行时（Container Runtime）是 Kubernetes（K8S）最重要的组件之一，负责管理镜像和容器的生命周期。Kubelet 通过 Container Runtime Interface（CRI）与容器运行时交互，以管理镜像和容器。

TKE 支持用户选择 Containerd 和 Docker 作为运行时组件：

- Containerd 调用链更短，组件更少，更稳定，占用节点资源更少。建议选择 Containerd。
- 当您遇到以下情况时，请选择 docker 作为运行时组件：
 - 如需使用 docker in docker。
 - 如需在 TKE 节点使用 docker build/push/save/load 等命令。
 - 如需调用 docker API。
 - 如需 docker compose 或 docker swarm。

如何修改运行时组件？

1. 登录 [容器服务控制台](#)，选择左侧导航栏中的集群。
2. 在“集群管理”列表页面，选择目标集群 ID，进入该集群基本信息页面。
3. 在“集群基本信息”中修改运行时组件。如下图所示：

说明：

修改运行时组件及版本，只对集群内无节点池归属的增量节点生效，不会影响存量节点。

Runtime components ⓘ



Containerd 和 Docker 组件常用命令是什么？

Containerd 不支持 docker API 和 docker CLI，但是可以通过 cri-tool 命令实现类似的功能。

镜像相关

镜像相关功能	Docker	Containerd
显示本地镜像列表	docker images	crictl images
下载镜像	docker pull	crictl pull
上传镜像	docker push	无
删除本地镜像	docker rmi	crictl rmi
查看镜像详情	docker inspect IMAGE-ID	crictl inspect IMAGE-ID

容器相关

容器相关功能	Docker	Containerd
显示容器列表	docker ps	crictl ps
创建容器	docker create	crictl create
启动容器	docker start	crictl start
停止容器	docker stop	crictl stop
删除容器	docker rm	crictl rm
查看容器详情	docker inspect	crictl inspect
attach	docker attach	crictl attach
exec	docker exec	crictl exec
logs	docker logs	crictl logs
stats	docker stats	crictl stats

POD 相关

POD 相关功能	Docker	Containerd
----------	--------	------------

POD 相关功能	Docker	Containerd
显示 POD 列表	无	<code>crictl pods</code>
查看 POD 详情	无	<code>crictl inspectp</code>
运行 POD	无	<code>crictl runp</code>
停止 POD	无	<code>crictl stopp</code>

调用链区别有哪些？

- Docker 作为 K8S 容器运行时，调用关系如下：

```
kubelet --> docker shim (在 kubelet 进程中) --> dockerd --> containerd
```

- Containerd 作为 K8S 容器运行时，调用关系如下：

```
kubelet --> cri plugin (在 containerd 进程中) --> containerd
```

其中 dockerd 虽增加了 swarm cluster、docker build、docker API 等功能，但也会引入一些 bug，而与 Containerd 相比，多了一层调用。

包括 exec, preStop, ipv6, 日志 stdout 的格式等。

Stream 服务的区别

```
kubectl exec/logs
```

 等命令需要 kubelet 在 apiserver 跟容器运行时之间建立流转发通道。

Stream 服务的原理是什么？

可以通过了解 `kubectl exec` 命令的原理来了解 CRI 的 Stream Service 是如何工作的：

1. dockershim 或 containerd 在启动后会监听某个端口，用以运行 Stream 服务。
2. 当执行 `kubectl exec` 等命令时，请求经过 kube-apiserver 找到 Pod 对应的节点，转发到 kubelet。
3. 此时 kubelet 会请求 CRI-Service (dockershim 或 containerd-cri) 的 GetExec 接口，CRI-Server 会为本次请求生成一个随机的 Token，记录后和 CRI Stream server 监听端口组合成 URL，返回给 kubelet。
4. kubelet 将 kube-apiserver 发过来的 HTTP 请求升级 websocket，并作为 apiserver 和 CRI Stream 之间的 proxy 转发数据。

如何在 Containerd 中使用并配置 Stream 服务？

Docker API 本身提供 stream 服务，Dockershim 位于 kubelet 内部有默认的配置 "127.0.0.1:0"。

Containerd 的 stream 服务需要单独配置：

```
[plugins.cri]
stream_server_address = "127.0.0.1"
stream_server_port = "0"
enable_tls_streaming = false
```

K8S 1.11 前后版本配置区别是什么？

Containerd 的 stream 服务在 K8S 不同版本运行时场景下配置不同。

- 在 K8S 1.11 之前：
Kubelet 不会做 stream proxy，只会做重定向。即 Kubelet 会将 containerd 暴露的 stream server 地址发送给 apiserver，并让 apiserver 直接访问 containerd 的 stream 服务。此时，您需要给 stream 服务转发器认证，用于安全防护。
- 在 K8S 1.11 之后：
K8S 1.11 引入了 [kubelet stream proxy](#)，使 containerd stream 服务只需要监听本地地址即可。

容器 Exec 的区别

Docker 和 Containerd 在 Exec 的实现上略有区别，区别主要在 Execsync 的实现上，也即执行单条命令的情况。

因此 `kubectl exec` 命令在不指定参数 `-it` 时和 pod lifecycle 中的 ExecProbe 在 Runtime 类型不同的节点上表现可能稍有不同。

kubectl exec 的区别

- docker exec 时会以 **当前exec首进程结束** 为 exec 结束的标志。
- CRI exec 会以 **当前exec中所有进程结束** 为本次 exec 结束。

如 `kubectl exec <pod-id> -- bash -c "nohup sleep 10 &"` 命令，在 Runtime 是 docker 的节点上会在两秒左右结束；而在 containerd 的节点上需要等到 sleep 进程退出后才能结束。

pod exec probe 的区别

kubelet 在实现 exec probe 时使用了 CRI Runtime 的 ExecSync 接口，因此 exec probe 和 `kubectl exec # no -t -i` 的表现一致，也即：

- 在 Docker 节点上，exec probe 如果残留子进程仍会正常退出。
- 在 Containerd 节点上，exec probe 会等到 probe 中所有的进程退出再结束。

区别导致的影响主要出现在 pod lifecycle 的 postStartHook 和 preStopHook 中，如果在 hook 中使用 exec probe 并且出现残留子进程的情况，在 containerd 的节点上可能会遇到 Pod 长期卡在 containerCreating 状态。原因是 kubelet

在 syncPod 时会逐个容器拉起，并执行 probe，如果某个 probe 因上述原因阻塞住，会导致后续容器无法启动。

在 ExecProbe 中拉起子进程并退出父进程属于 K8S 中未定义的行为，具体表现可能会和运行时版本、种类相关，因此建议尽量不要在 probe 执行过于复杂的操作。

容器网络的区别

在正常情况下，Pod 中的容器会共享同一个 Network Namespace，因此 Pod 需要在创建 Sandbox 容器时将网络准备好。为了更好的说明区别，我们简单介绍下 Pod 网络初始化的流程：

1. kubelet 调用 CRI Runtime（dockershim 或 containerd）创建 Pod 的 Sandbox 容器。
2. CRI Runtime 调用底层 docker 或 containerd 创建 pause 容器（此时 pause 进程还没启动，但已经初始化 Network Namespace）。
3. CRI Runtime 调用 CNI 执行在 Network Namespace 中创建 veth 并加入 cbr0 网桥等网络初始化操作。
4. CRI Runtime 启动 pause 容器，pause 进程被拉起。
5. kubelet 继续调用 CRI Runtime 执行创建容器等后续操作。

Docker 和 containerd 在**创建 pause 容器并初始化 Network Namespace**和**调用 CNI 初始化 veth**这两步有区别。

创建 pause 容器

Containerd 是为了 kubernetes 设计的 CRI Runtime，没有独立的网络模块；但 Docker 在设计时带有自己的网络功能的，因此 docker 在创建 Pause 容器时，会进行 Docker 特有的网络设置。该设置导致和 Containerd 最大的区别是在不使用 IPv6 的情况下，Docker 会将容器 Network Namespace 中内核参数 `net.ipv6.conf.all.disable_ipv6` 设置为1，也即关闭容器内的 ipv6 选项。

同样的 Pod 在 Docker 的节点上只开启了 IPv4，而在 Containerd 的节点上会同时开启 IPv4 和 IPv6。

同时开启 IPv4 和 IPv6 的情况中，DNS 解析可能会同时发出v4、v6两个版本的包。在某些情况，业务如果需要频繁进行 DNS 解析，可能会触发 DNS 解析库的 Bug（取决于 Pod 业务的实现时的依赖）。在 Containerd 节点上，可以通过给 Pod 添加 init container 来针对 Pod 关闭 IPv6 设置。代码如下：

```
apiVersion: v1
kind: Pod
...
spec:
  initContainers:
  - image: busybox
    name: sysctl
    command: ["sysctl", "-w", "net.ipv6.conf.all.disable_ipv6=1"]
    securityContext:
      privileged: true
  ...
```

调用 CNI

两者在调用 CNI 上没有实质区别。

对比项	Docker	Containerd
谁负责调用 CNI	Kubelet 内部的 docker-shim	Containerd 内置的 cri-plugin
如何配置 CNI	Kubelet 参数 <code>--cni-bin-dir</code> 和 <code>--cni-conf-dir</code>	Containerd 配置文件 (toml) : <pre>[plugins.cri.cni] bin_dir = "/opt/cni/bin" conf_dir = "/etc/cni/net.d"</pre>

容器日志的区别

对比项	Docker	Containerd
存储路径	如果 Docker 作为 K8S 容器运行时，容器日志的落盘将由 docker 来完成，保存在类似 <code>/var/lib/docker/containers/\$CONTAINERID</code> 目录下。Kubelet 会在 <code>/var/log/pods</code> 和 <code>/var/log/containers</code> 下面建立软链接，指向 <code>/var/lib/docker/containers/\$CONTAINERID</code> 该目录下的容器日志文件。	如果 Containerd 作为 K8S 容器运行时，容器日志的落盘由 Kubelet 来完成，保存至 <code>/var/log/pods/\$CONTAINER_NAME</code> 目录下，同时在 <code>/var/log/containers</code> 目录下创建软链接，指向日志文件。
存储大小	Pod 中的每个容器，docker 默认会保留 $100\text{MB} \times 10 = 1\text{G}$ 日志	Pod 中的每个容器，containerd 默认会保留 $10\text{MB} \times 5 = 50\text{MB}$ 日志
配置参数	在 docker 配置文件中指定： <pre>"log-driver": "json-file", "log-opts": {"max-size": "100m", "max-file": "5"}</pre>	- 方法一：在 kubelet 参数中指定： <pre>--container-log-max-files= --container-log-max-size="100Mi"</pre> - 方法二：在 KubeletConfiguration 中指定： <pre>"containerLogMaxSize": "100Mi", "containerLogMaxFiles": 5,</pre>
把容器	把数据盘挂载到 “data-root” (缺省是	创建一个软链接 <code>/var/log/pods</code> 指

日志保存到数据盘	<code>/var/lib/docker</code>) 即可。	数据盘挂载点下的某个目录。 在 TKE 中选择“将容器和镜像存储在数据盘”，会自动创建软链接 <code>/var/log/pods</code> 。
----------	------------------------------------	---

自建 DNS 导致节点初始化异常

最近更新时间：2022-05-09 11:40:29

背景信息

在使用容器服务 TKE 的自定义镜像时，为了能够解析到业务内部的相关服务，用户在自定义镜像中修改了 DNS 的解析顺序或将腾讯官方的 DNS 解析地址完全替换为自建 DNS。

操作影响

上述情况可能会导致节点在注册进集群的过程中，无法解析到腾讯云的官方资源库，进而大概率出现节点初始化失败、网络、存储等相关组件功能异常等情况。

- 节点初始化：在节点初始化流程中可能会报错误信息 “Failed to resolve address, dns may be changed”。
- 网络组件：网络组件如 IPAMD 功能依赖腾讯云内网 DNS 解析，解析不到腾讯云的官方资源库可能会导致网络组件功能不可用。
- 存储组件：存储组件如 CBS-CSI 的 mount / unmount 失败。

说明：

集群内的相关组件安装需要依赖腾讯云官方资源库，解析地址为：

```
nameserver 183.60.83.19
```

```
nameserver 183.60.82.98
```

解决措施

将腾讯云 DNS Nameserver 配置为自建 DNS 的上游

建议将 `/etc/resolv.conf` 配置中的 `nameserver` 添加到自建 DNS 服务器的上游，因为部分服务依赖腾讯云内部 DNS 解析，如果未将其设为自建 DNS 的上游，可能导致部分服务无法正常工作。本文以 [BIND 9](#) 为例修改配置文件，将上游 DNS 地址写入 `forwarders` 中，示例如下：

```
options {  
forwarders {  
183.60.83.19;  
183.60.82.98;
```

```
};  
...
```

完成上述操作后的 24 小时内，节点初始化流程的自动重试策略会不断尝试执行节点初始化工作直到解析到腾讯云官方资源库；超过 24 小时后，需要用户删除节点重新创建。若上述解决措施不生效，请 [提交工单](#) 来寻求帮助。

注意：

自建 DNS Server 和请求源不在同个 Region，可能会导致部分不支持跨域访问的腾讯域名失效。

服务类

最近更新时间：2022-06-10 16:48:46

创建服务常见问题

服务的名称为什么不能重复？

服务名称是当前集群下的服务的唯一标识，服务之间可以通过服务名称+访问端口的形式互相访问。

创建服务能否使用非腾讯云或 dockerhub 镜像的第三方镜像？

您可以通过登录到主机执行 `docker login` 命令登录到第三方镜像仓库拉取。

使用外网服务的有什么前置条件？

确保集群内的云服务器拥有外网带宽，否则外网服务将创建失败。

内存限制，CPU 限制如何填写？

详情参考 [容器服务资源限制说明](#)。

创建服务时的特权级是什么意思？

开启该选项会使得容器内程序具有真正的 `root` 权限。在容器内程序需要进行高级系统操作时建议开启，如搭建 `nfs` 服务器。

负载均衡可以在创建时就指定安全组吗？

可以，目前支持以下两个方案，实现服务使用负载均衡时指定安全组：

- 使用已有负载均衡。先创建负载均衡并配置安全组，再挂载给服务。详情请参见 [Service 使用已有 CLB](#)。
- 可在服务中通过 `TkeServiceConfig` 配置安全组，负载均衡创建时会根据配置使用对应安全组。如需使用此功能，请 [提交工单](#) 进行申请。

说明：

集群内进行服务访问时，建议不要通过负载均衡 IP 进行访问，以避免出现访问不通的情况。

一般情况下，4层负载均衡会绑定多台 Node 作为 `real server` (rs)，使用时需要限制 `client` 和 `rs` 不能存在于同一台云服务器上，否则会有有一定概率导致报文回环失败。

当 Pod 去访问负载均衡时，Pod 为源 IP，当其传输到内网时负载均衡也不会做 `snat` 处理将源 IP 转化为 Node IP，则负载均衡收到报文时无法判断是哪个 Node 发送的，避免回环策略就无法生效，所有的 `rs` 都可能被转发。当转发到 `client` 所在的 Node 上时，负载均衡就无法收到回包，从而导致访问不通。

更新服务容器数量常见问题

更新容器数量需注意哪些问题？

需确认 CPU、内存资源充足，否则容器将创建失败。

能否将容器数量设为0？

可以，通过将容器数量设置为0，保存服务配置并且释放资源占用。

更新服务配置常见问题

更新服务是否支持滚动更新？

支持滚动更新和快速更新两种方式。

公网负载均衡可以切换为内网负载均衡吗？

可以，目前支持公网切换至 VPC 内网、VPC 内网切换至公网及 VPC 不同子网间切换。详情请参见 [Service 生命周期管理](#)。

注意：

- 若为服务负责负载均衡资源的生命周期管理，则负载均衡及其公网 IP 将会被释放。
- 公网切换内网的过程并不是瞬间的，公网负载均衡服务下线到内网负载均衡并提供服务，此过程需要一定的时间。建议您先在集群中配置一个内网服务资源，并进行测试。等到流量切换完成之后，再删除原有公网服务资源。

删除服务常见问题

删除服务后服务创建的负载均衡会自动销毁么？

删除服务时，将会同时删除创建该服务时自动创建的负载均衡。若在创建服务时选用的是已有负载均衡，则该负载均衡将不会受到任何影响。

删除服务是否会影响业务数据？

删除服务不会删除业务容器，数据不会受到影响，该操作无需提前备份数据。

服务运行常见问题

如何设置容器系统时间为北京时间？

容器默认使用 UTC 时间，使用容器时经常碰到容器系统时间和北京时间差8小时的问题，解决方法是在 `dockerfile` 中创建时区文件。详情请参见 [解决容器内时区不一致问题](#)。

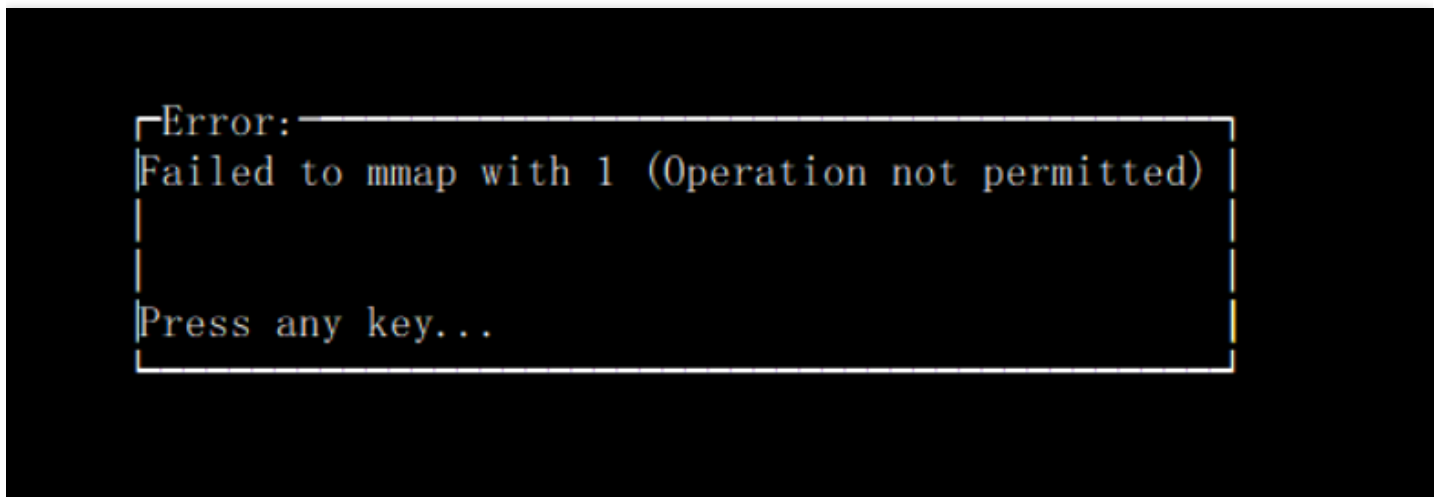
Dockerhub 部分镜像如 ubuntu、php 和 busybox 等在容器服务里运行异常怎么办？

运行异常是因为没有设置启动命令或者默认的启动命令为 `bash`，导致容器执行完启动程序就退出了。要使容器一直运行，容器里 PID 为1的进程必须是常驻进程，否则 PID 为1的进程一结束容器就退出了。对于部分镜像如 `centos` 等，可以使用 `/bin/bash` 作为运行命令，`-c sleep 800000` 作为运行参数来创建服务，在控制台填运行参数时 `-c` 和 `sleep 800000` 必须放在两行。

目前已知的使用默认参数启动不了服务的镜像包括这些：`clearlinux`、`ros`、`mageia`、`amazonlinux`、`ubuntu`、`clojure`、`crux`、`gcc`、`photon`、`java`、`debian`、`oraclelinux`、`mono`、`bash`、`buildpack-deps`、`golang`、`sourcemage`、`swift`、`openjdk`、`centos`、`busybox`、`docker`、`alpine`、`ibmjava`、`php`和`python`。

容器执行 `perf top -p` 查看进程 CPU 情况提示 “Operation not permitted”

容器执行 `perf top -p` 查看进程 CPU 情况提示 “Operation not permitted”，如下图所示：



Docker 默认配置文件阻止了重要的系统调用，`perf_event_open` 可能会泄漏主机上的大量信息所以被禁止使用。如果您需要调用请配置特权容器或者直接修改 `Pod yamI` 字段 `privileged` 为 `true`，您需自行评估安全风险。

TKE Serverless 集群

TKE Serverless 集群相关

最近更新时间：2023-03-27 11:08:16

本文汇总了 TKE Serverless 的集群常见问题，介绍集群相关常见问题的出现原因及解决办法。

为什么 Pod 规格与填写的 Request/Limit 不一致？

在分配 Pod 资源量时，TKE Serverless 需要对工作负载设置的 Request 及 Limit 进行计算，自动判断 Pod 运行所需的资源量，并非按照设置的 Request 及 Limit 值进行资源分配。详情请参见 [CPU Pod 规格计算方法](#) 及 [GPU Pod 规格计算方法](#)，进一步了解如何通过 Request、Limit 自动计算指定资源规格。

如何新增或修改 TKE Serverless 集群的容器网络？

在创建集群时，需要选择一个 VPC 网络作为集群网络，同时指定一个子网作为容器网络，详情请参见 [容器网络说明](#)。Serverless 集群的 Pod 会直接占用容器网络子网 IP。在使用集群过程中，如需新增或修改容器网络，需要通过新增/移出超级节点的操作来实现，具体操作请参考以下步骤。

步骤1：新建超级节点以新增容器网络

1. 登录容器服务控制台，选择左侧导航栏中的 [集群](#)。
2. 单击需要修改容器网络的集群 ID，进入集群基本信息页。
3. 选择左侧的 **超级节点**，在 **超级节点** 页面，单击 **新建**。
4. 在 **新建超级节点** 页面，选择 IP 充足的容器网络，单击 **确定** 即可完成新建。



步骤2：移出超级节点以删除容器网络

注意

移出超级节点后必须保证 Serverless 集群至少存在一个超级节点，即若此时只存在一个超级节点，则无法执行移出操作。

移出超级节点之前，需要将此超级节点上的 Pod 全部驱逐到其他超级节点上（不包含 DaemonSet 管理的 Pod），完成驱逐后才可执行移出操作，否则会导致移出节点失败。具体操作请参考以下步骤。

1. 登录容器服务控制台，选择左侧导航栏中的 [集群](#)。
2. 单击需要修改容器网络的集群 ID，进入集群基本信息页。
3. 选择左侧的**超级节点**，在**超级节点**页面选择节点名称右侧的**更多 > 驱逐**。如下图所示：



4. 在**驱逐节点**页面，确定节点信息，单击**确定**。驱逐后，该超级节点的状态会变更为“已封锁”，将不再向该节点上调度 Pod。

注意

驱逐会导致 Pod 重建，请谨慎操作。

5. 在**超级节点**页面中，选择节点名称右侧的**移出**。

6. 在**删除节点**页面中，单击**确定**即可完成删除节点操作。

Pod 因子网 IP 耗尽调度失败时，该如何处理？

当 Pod 因容器网络的子网 IP 耗尽而发生调度失败时，会在节点日志中观察到如下图所示两种事件：

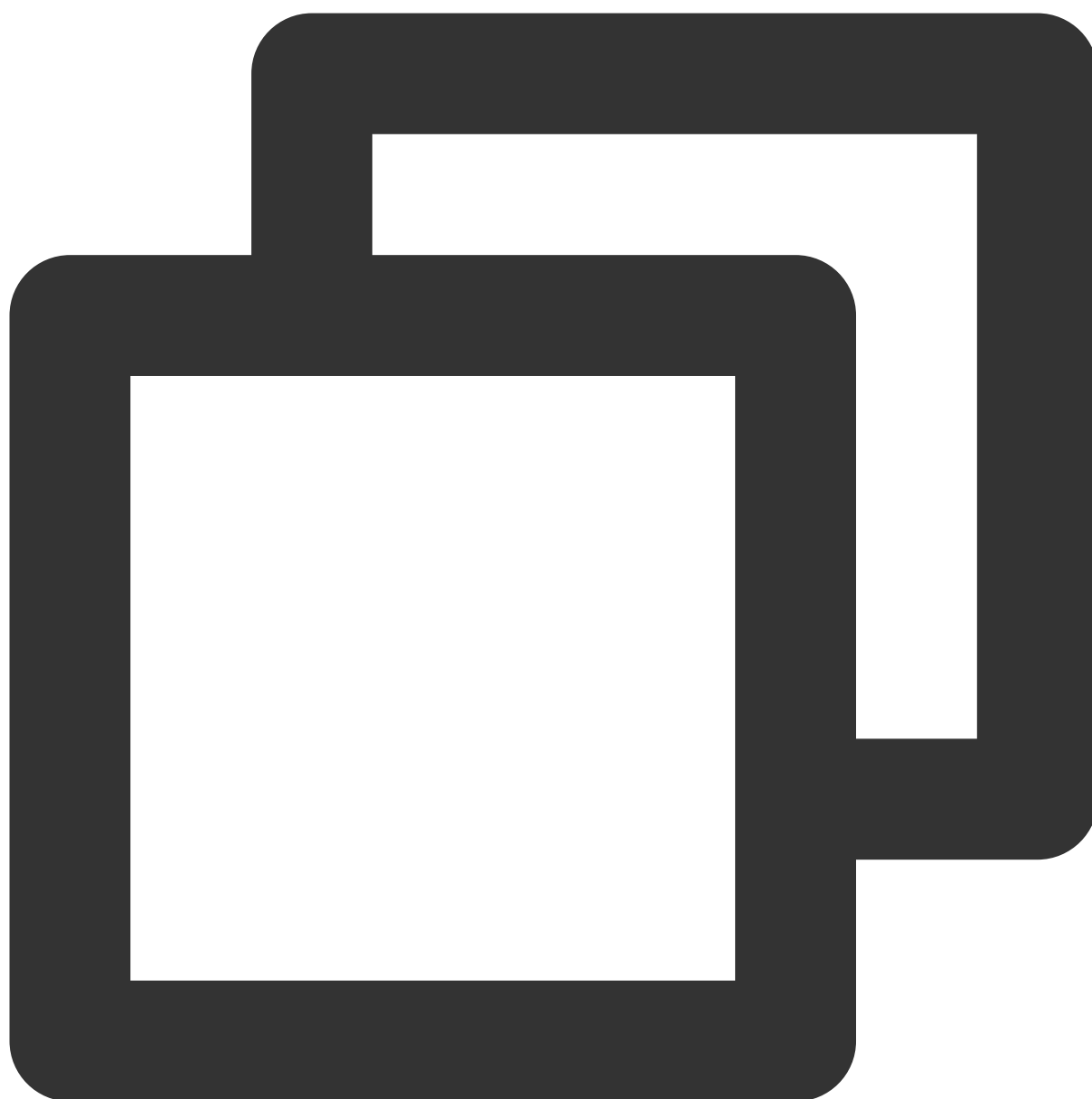
事件一：



事件二：

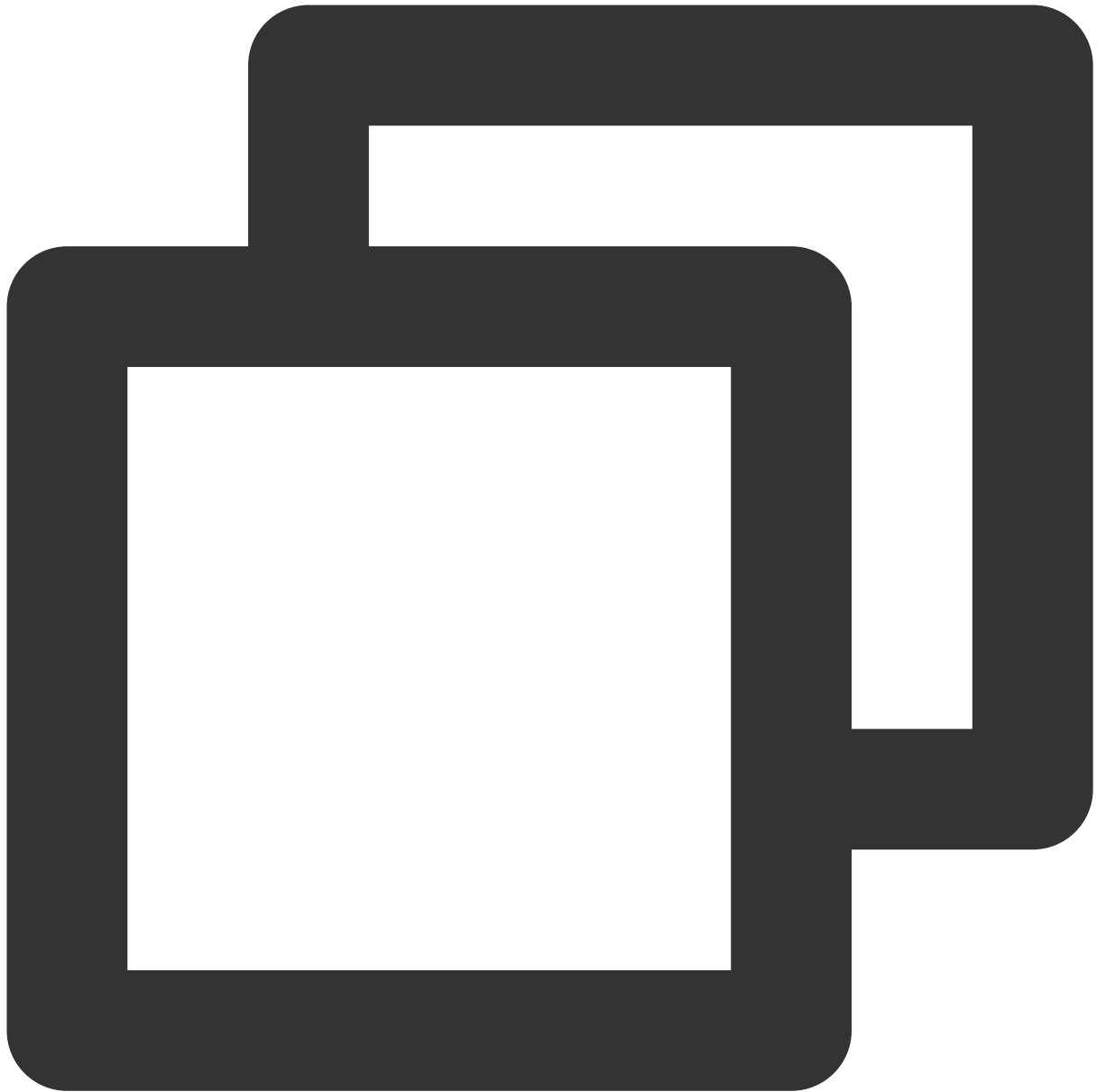


您可通过 [容器服务控制台](#) 或在命令行执行以下命令查询超级节点的 YAML。

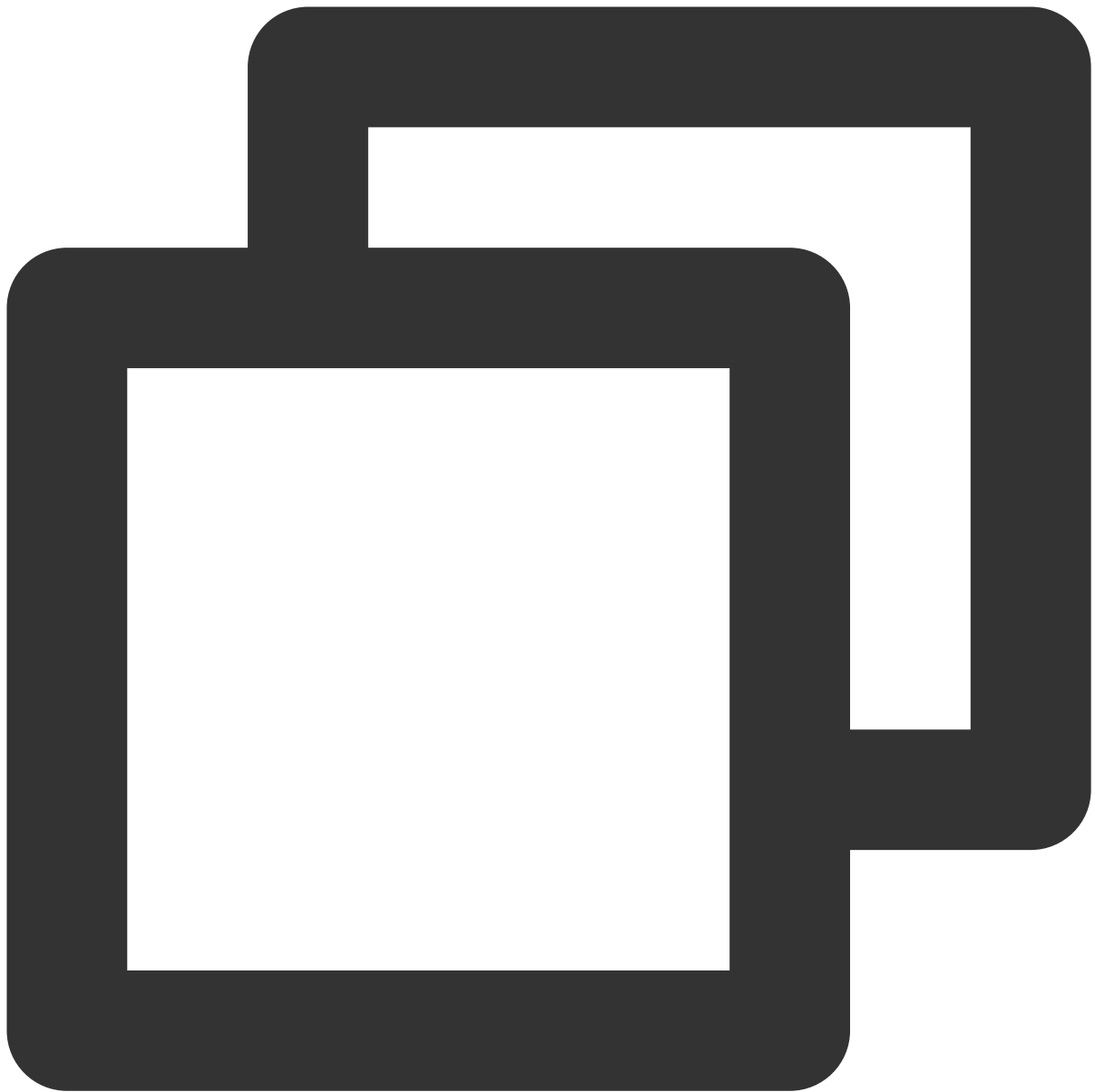


```
kubectl get nodes -oyaml
```

返回结果如下所示：



```
spec:
  taints:
  - effect: NoSchedule
    key: node.kubernetes.io/network-unavailable
    timeAdded: "2021-04-20T07:00:16Z"
```



```
- lastHeartbeatTime: "2021-04-20T07:55:28Z"  
  lastTransitionTime: "2021-04-20T07:00:16Z"  
  message: eklet node has insufficient IP available of subnet subnet-bok73g4c  
  reason: EKLetHasInsufficientSubnetIP  
  status: "True"  
  type: NetworkUnavailable
```

说明此时 Pod 是因容器网络的子网 IP 耗尽而发生调度失败。发生以上情况时，需要新建超级节点去新增子网，以此扩展集群 Pod 可用网段，新建超级节点的操作请参见 [新建超级节点](#)。

TKE Serverless 安全组使用指引和说明有哪些？

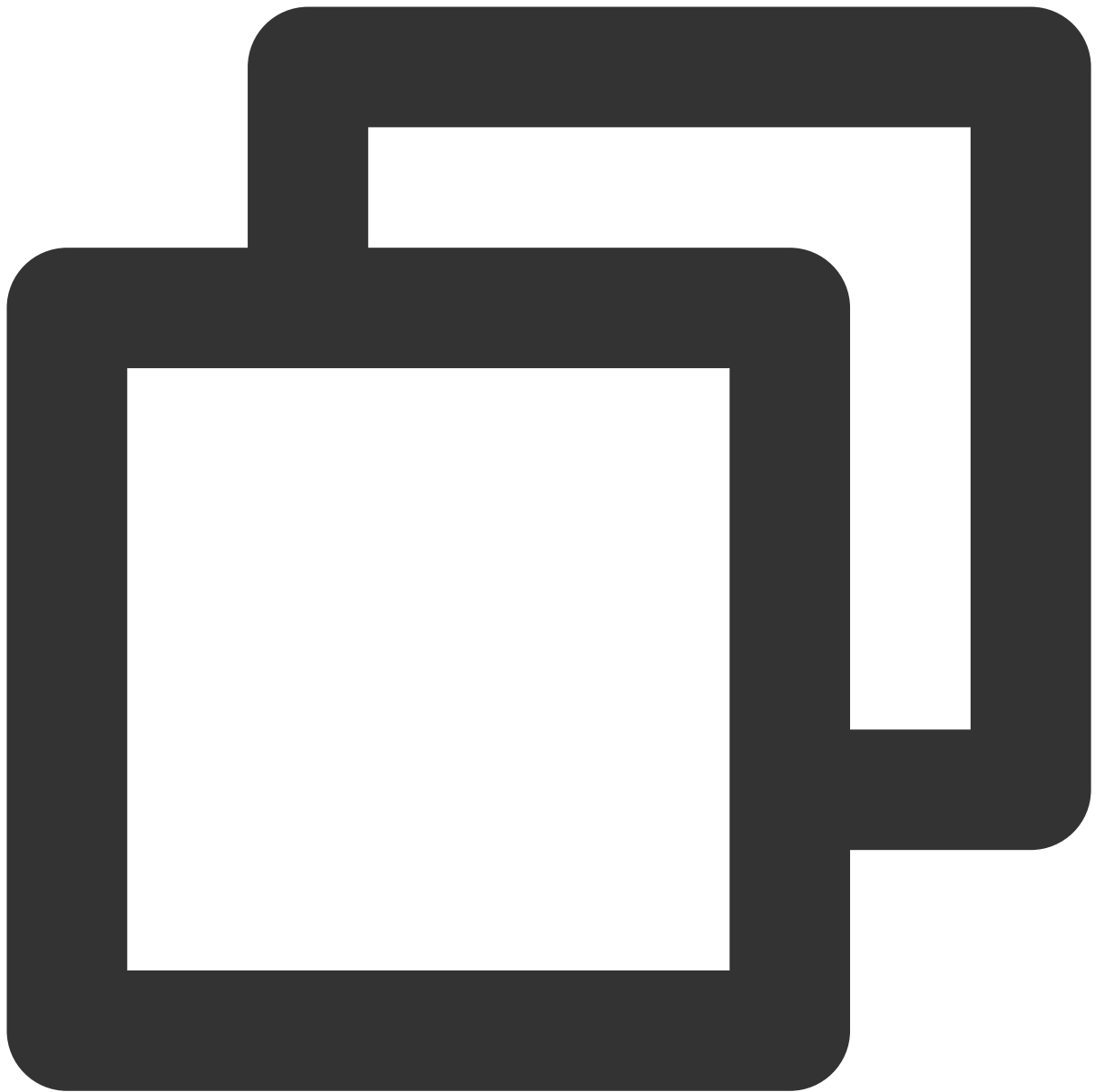
在创建 Serverless 集群 Pod 时，若未指定安全组，则默认使用 default 安全组，您也可以通过 Annotation `eks.tke.cloud.tencent.com/security-group-id` : 安全组 ID 为 Pod 指定安全组，请确保同地域已存在该安全组 ID。关于此 Annotation 的详细说明请参加 [Annotation 说明](#)。

如何设置容器终止消息？

Kubernetes 可以通过 `terminationMessagePath` 设置容器退出的消息来源，即当容器退出时，Kubernetes 将从容器的 `terminationMessagePath` 字段中指定的终止消息文件中检索终止消息，并使用此内容来填充容器的终止消息，消息默认值为：`/dev/termination-log`。

此外，您还可以设置容器的 `terminationMessagePolicy` 字段，进一步自定义容器终止消息。该字段默认值为 `File`，即仅从终止消息文件中检索终止消息。您可以根据需求设置为 `FallbackToLogsOnError`，即在容器因错误退出时，如果终止消息文件为空，则使用容器日志输出的最后一部分内容来作为终止消息。

代码示例如下：



```
apiVersion: apps/v1beta2
kind: Deployment
metadata:
  name: nginx
spec:
  containers:
  - image: nginx
    imagePullPolicy: Always
    name: nginx
  resources:
    limits:
```

```
cpu: 500m
memory: 1Gi
requests:
  cpu: 250m
  memory: 256Mi
terminationMessagePath: /dev/termination-log
terminationMessagePolicy: FallbackToLogsOnError
```

通过以上配置，当容器错误退出且消息文件为空时，Get Pod 会发现 stderr 的输出显示在 containerStatuses 中。

如何使用 Host 参数？

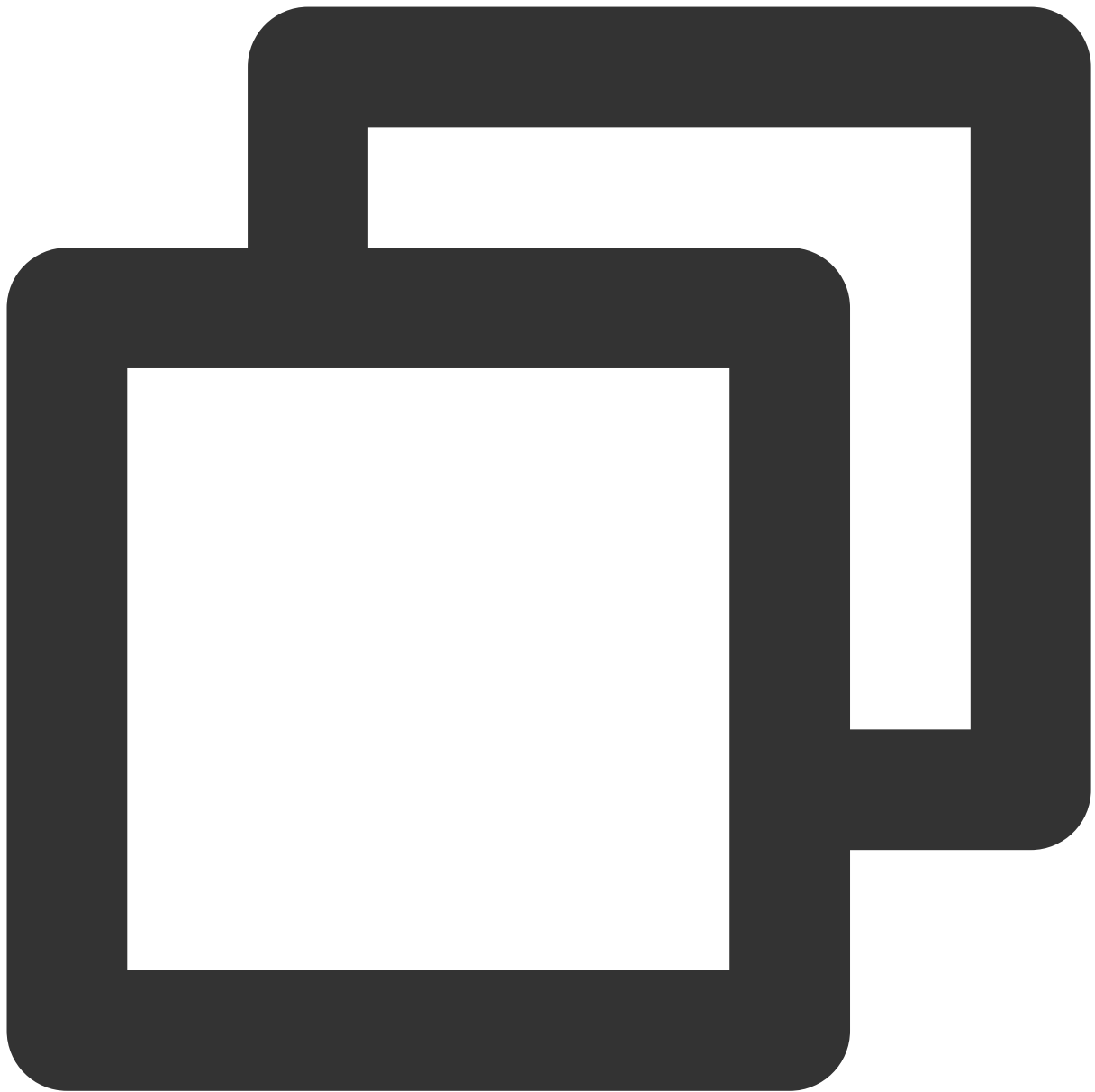
在使用 Serverless 集群时需要注意以下事项：

Serverless 集群虽然没有节点，但兼容 Hostpath、Hostnetwork: true、DnsPolicy: ClusterFirstWithHostNet 等与 Host 相关的参数。您在使用时请注意，因为没有节点，这些参数提供能力并不能完全与标准 k8s 对齐。

例如，期望使用 Hostpath 共享数据，但调度到同一个超级节点上的两个 Pod 查看到的是不同子机的 Hostpath，而且 Pod 重建后，Hostpath 的文件也将同时删除。

如何挂载 CFS/NFS？

在 Serverless 集群中，支持使用腾讯云 [文件存储 CFS](#)，也支持使用自建的文件存储 NFS 以 Volume 的形式挂载到 Pod 上，以实现持久化的数据存储。Pod 挂载 CFS/NFS 的 YAML 示例如下：



```
apiVersion: v1
kind: Pod
metadata:
  name: test-pd
spec:
  containers:
  - image: k8s.gcr.io/test-webserver
    name: test-container
    volumeMounts:
    - mountPath: /cache
      name: cache-volume
```

```
volumes:
- name: nfs
  nfs:
    path: /dir
    server: 127.0.0.1
---
```

`spec.volumes`：设置数据卷名称、类型、数据卷的参数。

`spec.volumes.nfs`：设置 NFS/CFS 盘。

`spec.containers.volumeMounts`：设置数据卷在 Pod 的挂载点。

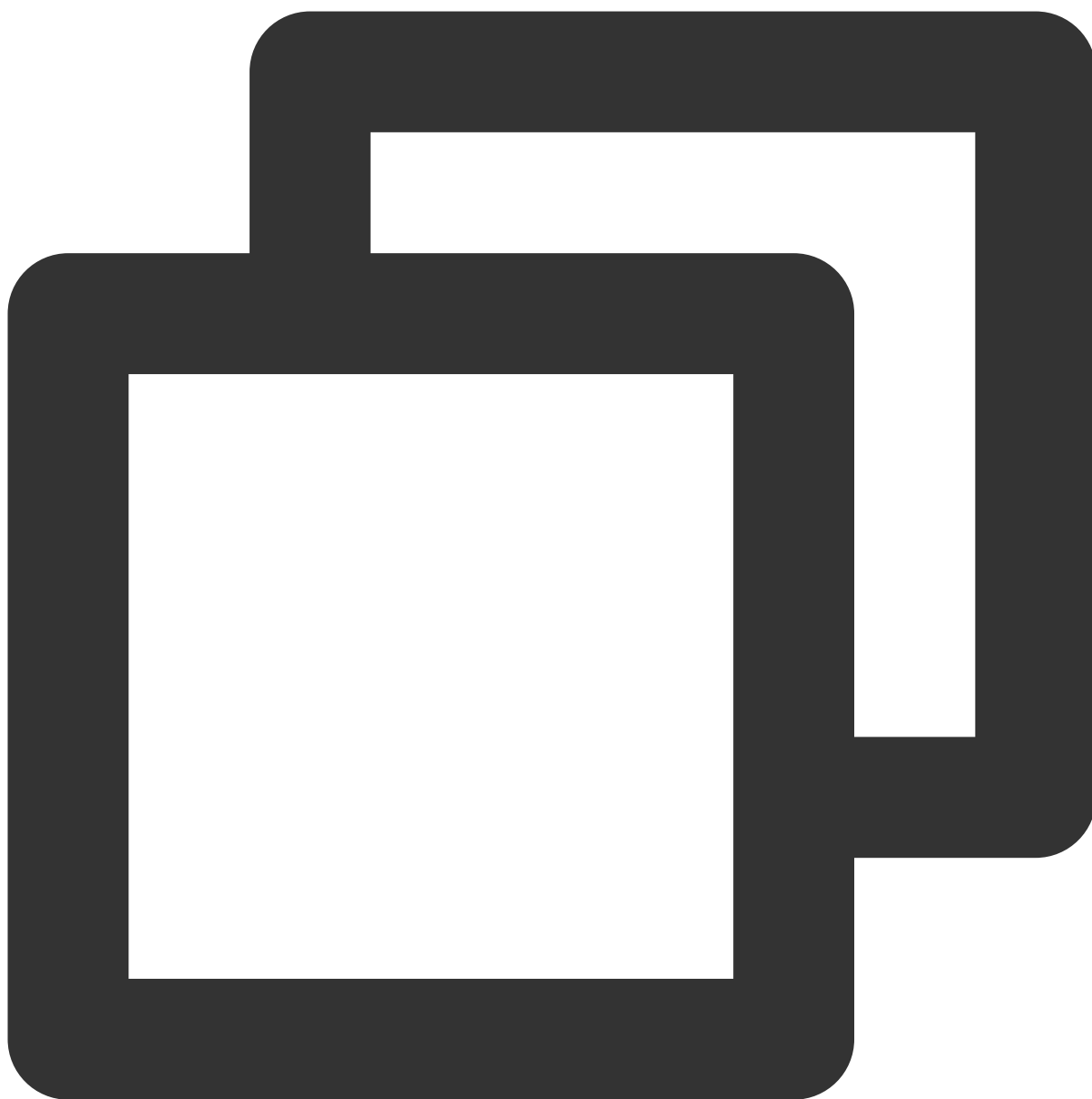
关于 Pod 挂载 Volume 的详细操作，请参见 [存储卷的使用说明](#)。

如何通过镜像复用加快容器启动速度？

TKE Serverless 支持缓存容器镜像以便在下次用同样镜像启动容器时加快启动速度。

复用条件：

1. 对于同一工作负载的 Pod，如果缓存时间内在同一个可用区（Zone）有 Pod 创建且销毁过，新建的 Pod 默认不重复拉取相同的镜像。
2. 如果不同的工作负载（包括 Deployment、Statefulset、Job）的 Pod 想复用镜像，可以使用如下 annotation：



```
eks.tke.cloud.tencent.com/cbs-reuse-key
```

同一个用户账号下，有同一个 annotation value 的 Pod，缓存时间内会尽量复用启动镜像，建议 annotation value 填写镜像名：`eks.tke.cloud.tencent.com/cbs-reuse-key: "image-name"`。

缓存时间：2小时。

如何解决复用镜像异常问题？

当启用复用镜像功能时，当创建 Pod 时，`$kubectl describe pod` 可能见到如下错误：

```
no space left on device: unknown
```

```
Warning FreeDiskSpaceFailed 26m eklet, eklet-subnet-xxx failed to garbage collect  
required amount of images. Wanted to free 4220828057 bytes, but freed 3889267064  
bytes
```

恢复方法：

无需任何操作，等待若干分钟后 Pod 会自动 running。

原因：

```
no space left on device: unknown
```

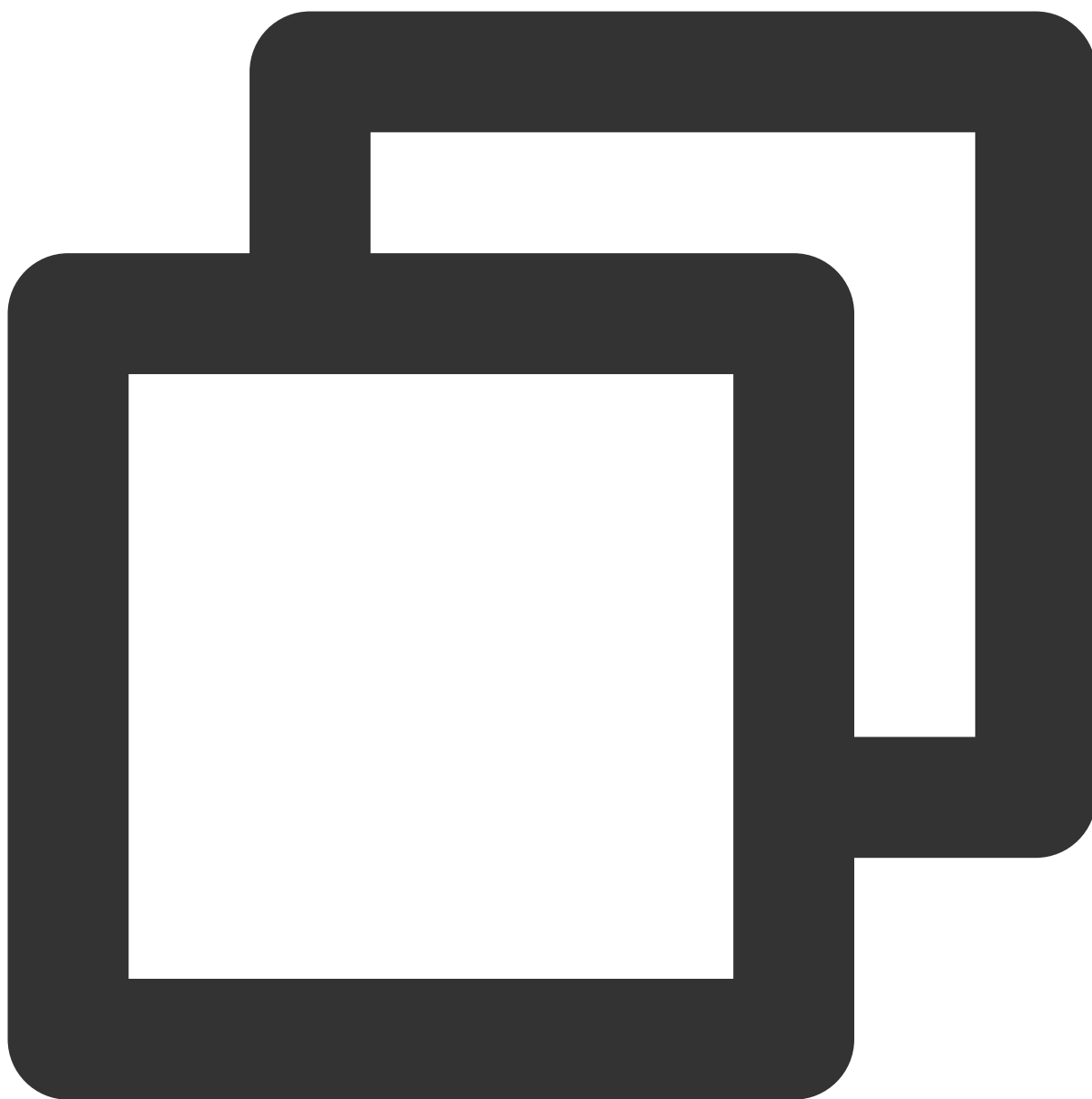
Pod 默认复用系统盘时，系统盘内的原有镜像占满磁盘空间，导致磁盘当前没有足够的空间下载新镜像，所以报错“no space left on device: unknown”。TKE Serverless 支持定时镜像回收机制，遇到磁盘空间占满时，会自动将系统盘中现有的多余镜像删除，给当前磁盘腾出可用空间。（耗时若干分钟）

```
Warning FreeDiskSpaceFailed 26m eklet, eklet-subnet-xxx failed to garbage collect  
required amount of images. Wanted to free 4220828057 bytes, but freed 3889267064  
bytes
```

该条日志说明当前 Pod 下载镜像需要4220828057空间，但目前只清空了3889267064空间的数据。产生该条 event 的原因为磁盘上有多个镜像，目前只清理了部分镜像，TKE Serverless 的定时镜像回收机制会继续清理镜像，直到能成功拉取到新镜像为止。

挂载自建的 nfs 时，事件报 Operation not permitted 如何处理？

如果您使用自建的 nfs 实现持久化存储时，连接时事件报 Operation not permitted。您需要修改自建 nfs 的 /etc/exports 文件，添加 /<path><ip-range>(rw,insecure) 参数，示例如下：



```
/data/ 10.0.0.0/16(rw,insecure)
```

Pod 磁盘满了 (ImageGCFailed) 如何处理？

TKE Serverless 的 Pod 默认免费提供 20G 可用的系统盘空间，若系统盘空间满了，可以通过如下方式进行处理。

1. 清理未使用的容器镜像

如果使用空间达到 80%，TKE Serverless 后台会触发容器镜像的回收流程，尝试回收未使用的容器镜像来释放磁盘空间。如果未能释放任何空间，则会有一条事件提醒：ImageGCFailed: failed to garbage collect required amount of

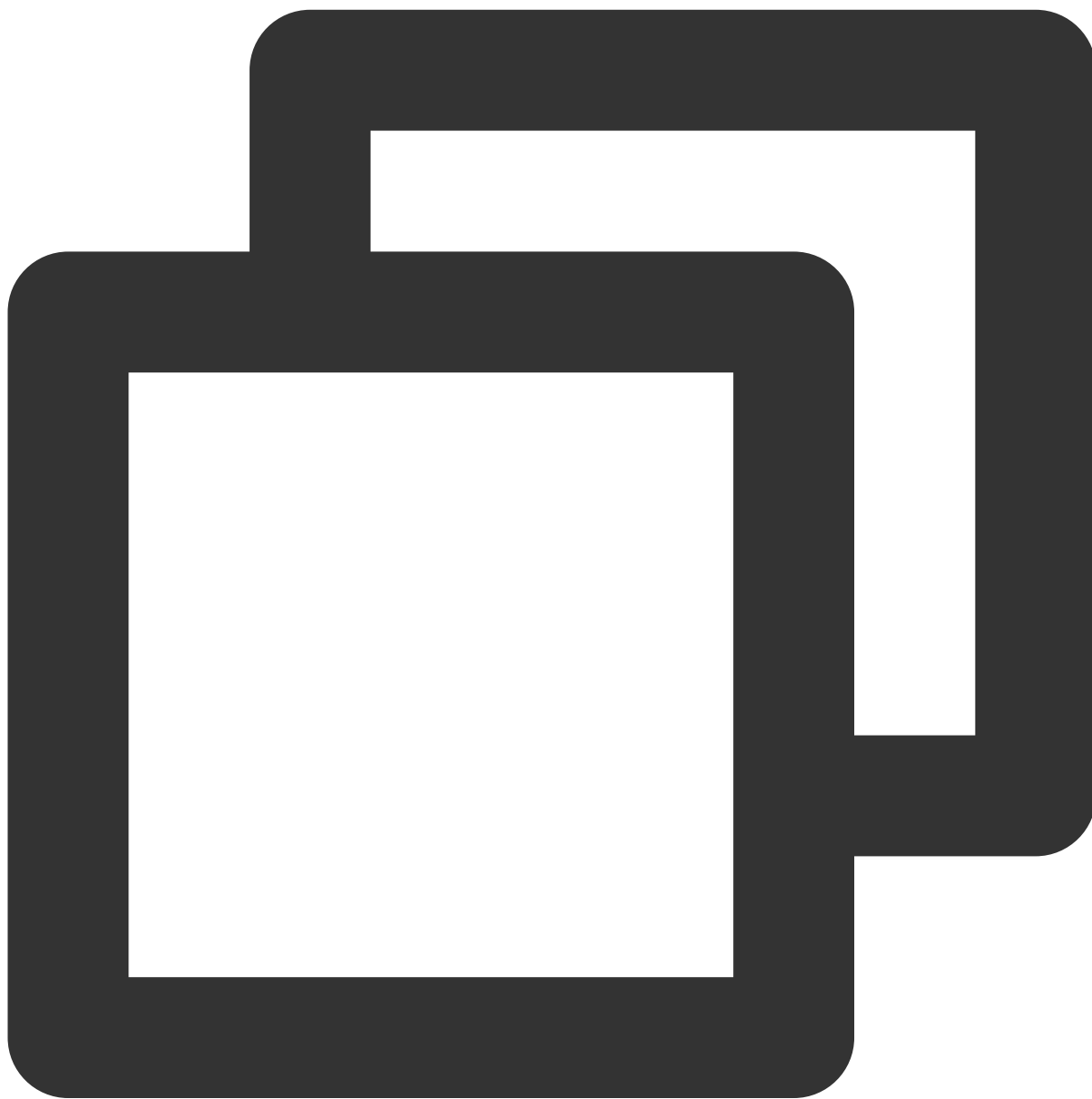
images，提醒用户磁盘空间不足。

常见磁盘空间不足的原因有：

业务有大量临时输出。您可以通过 `du` 命令确认。

业务持有已删除的文件描述符，导致磁盘空间未释放。您可以通过 `lsdf` 命令确认。

如果业务希望调整容器镜像回收的阈值，可以设置如下 `annotation`：



```
eks.tke.cloud.tencent.com/image-gc-high-threshold: "80"
```

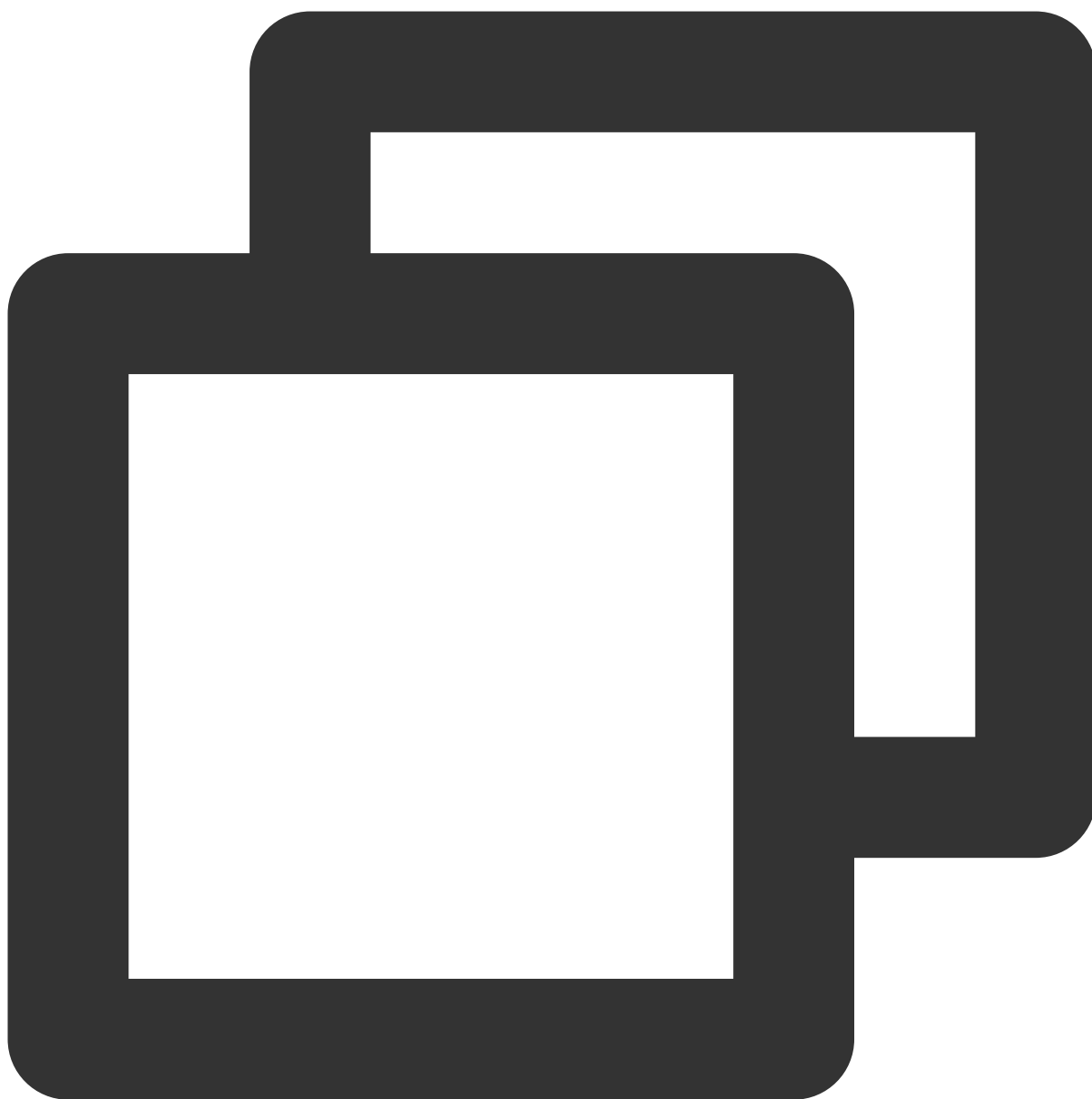
2. 清理已退出的容器

如果业务原地升级过，或者容器异常退出过，已退出的容器仍会保留，直到磁盘空间达到 85% 时才会清理已退出的容器。清理阈值可以使用如下 Annotation 调整：



```
eks.tke.cloud.tencent.com/container-gc-threshold: "85"
```

如果已退出的容器不想被自动清理（例如需要退出的信息进一步排障的），可以通过如下 Annotation 关闭容器的自动清理，但副作用是磁盘空间无法自动释放：



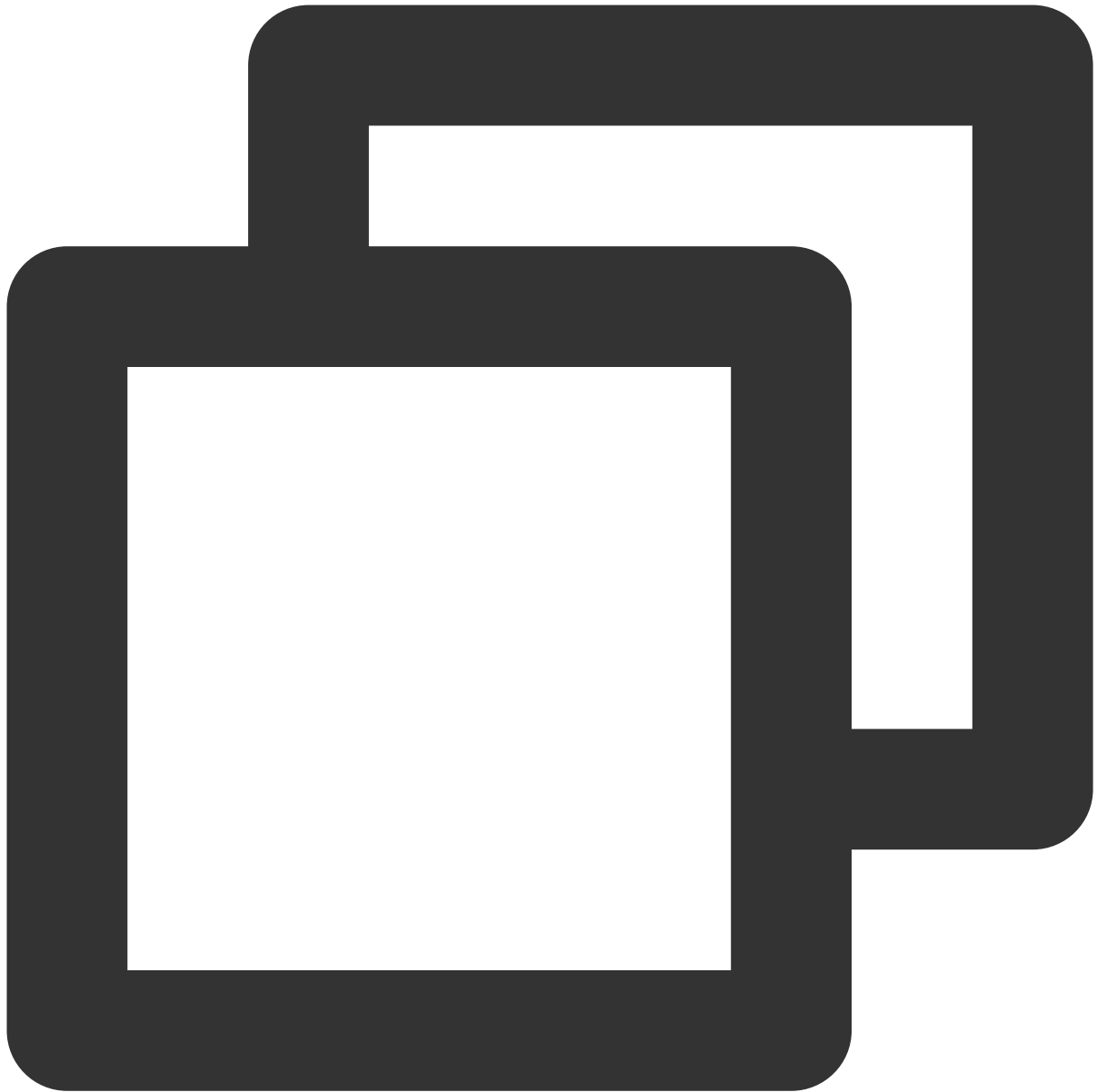
```
eks.tke.cloud.tencent.com/must-keep-last-container: "true"
```

说明

此特性上线时间为 2021-09-15，故在此时间前创建的 Pod，并未带有此特性。

3. 重启磁盘用量高的 Pod

业务需要在容器的系统盘用量超过某个百分比后直接重启 Pod，可以通过 Annotation 配置：



```
eks.tke.cloud.tencent.com/pod-eviction-threshold: "85"
```

只重启 Pod，不会重建子机，退出和启动都会进行正常的 `gracestop`、`prestop`、健康检查。

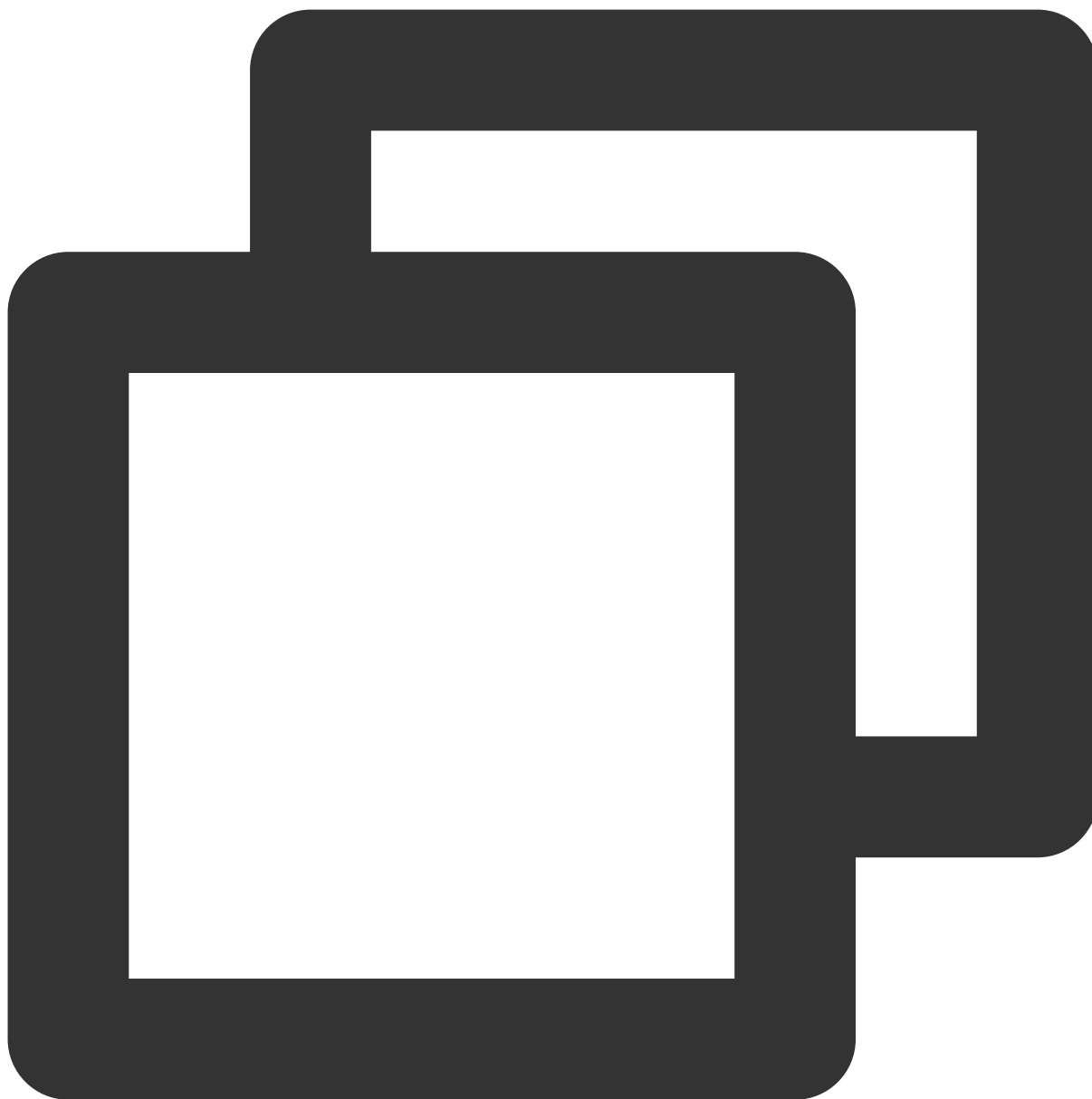
说明

此特性上线时间在 2022-04-27，故在此时间前创建的 Pod，需要重建 Pod 来开启特性。

9100 端口问题

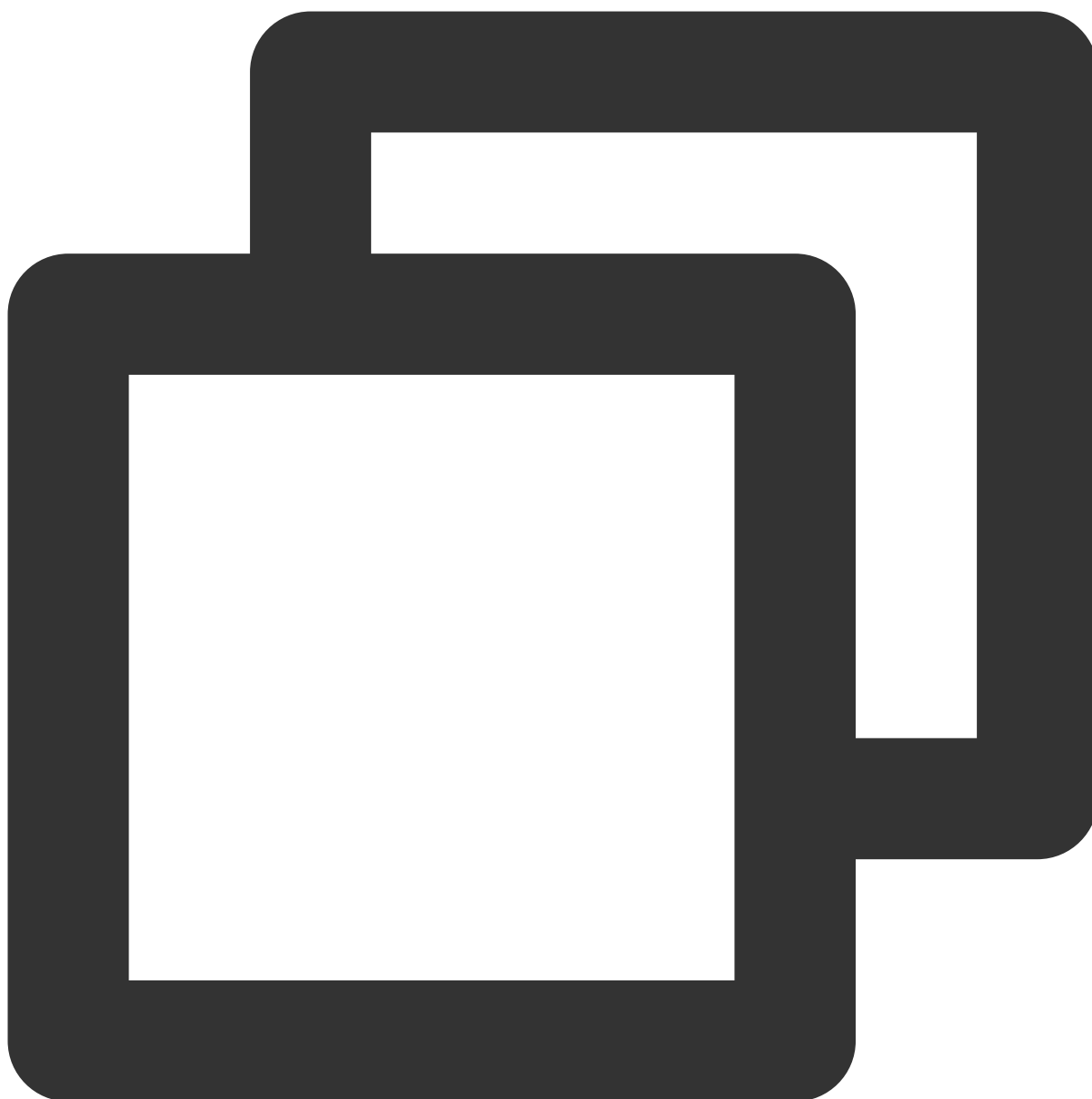
TKE Serverless 的 Pod 默认会通过 9100 端口对外暴露监控数据，用户可以执行以下命令访问 9100/metrics 获取数据：

获取全部指标：



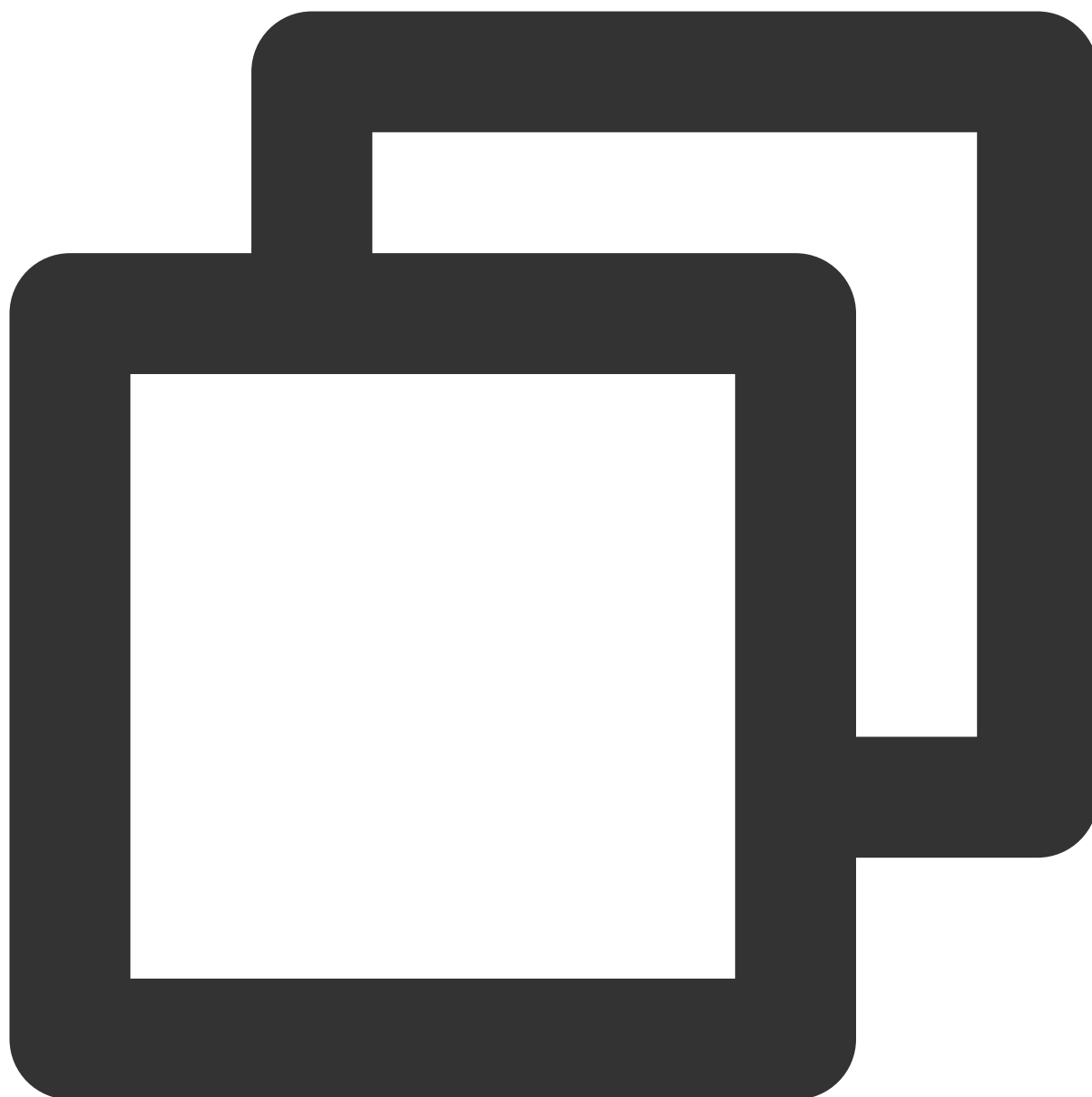
```
curl -g "http://<pod-ip>:9100/metrics"
```

大集群建议去掉 ipvs 指标：



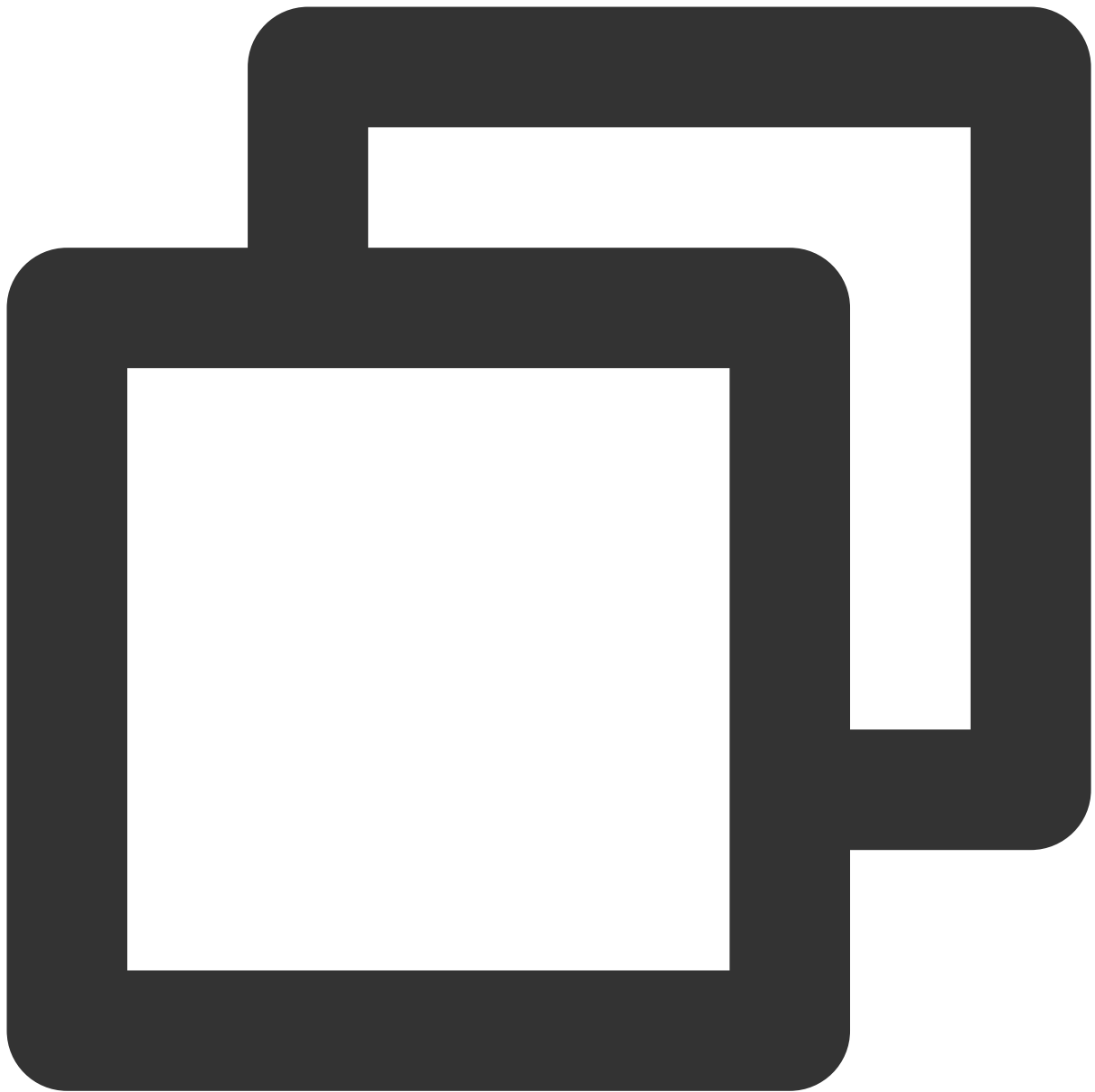
```
curl -g "http://<pod-ip>:9100/metrics?collect[]=ipvs"
```

如果业务本身需要监听 9100 端口，则可以在创建 Pod 时，通过使用 9100 之外的端口来收集监控数据，避免跟业务的 9100 端口冲突。配置方式如下：



```
eks.tke.cloud.tencent.com/metrics-port: "9110"
```

如果没有变更监控暴露的端口，业务直接监听 9100 端口，则在 TKE Serverless 新的网络方案里，将报错提醒用户 9100 端口已经被使用：



```
listen() to 0.0.0.0:9100, backlog 511 failed (1: Operation not permitted)
```

出现报错提醒时，需要为 Pod 添加一个 Annotation：`metrics-port`，以变更监控端口，再重建 Pod。

注意

如果 Pod 带有公网 eip，则需要设置安全组，注意 9100 端口问题，并放通需要的端口。

负载均衡相关

最近更新时间：2023-03-27 11:08:16

本文汇总了负载均衡相关常见问题，介绍与 Service/Ingress CLB 相关的各种常见问题的出现原因及解决办法。

本文档需要您：

熟悉 K8S 的 [基本概念](#)。例如 Pod、工作负载/Workload、Service、Ingress 等。

熟悉 [腾讯云控制台](#) 容器服务之 TKE Serverless 集群的常规操作。

熟悉通过 kubectl 命令行工具操作 K8S 集群中的资源。

注意：

您可通过多种方式操作 K8S 集群中的资源，本文档向您介绍如何通过腾讯云控制台进行操作，及如何通过 kubectl 命令行工具进行操作。

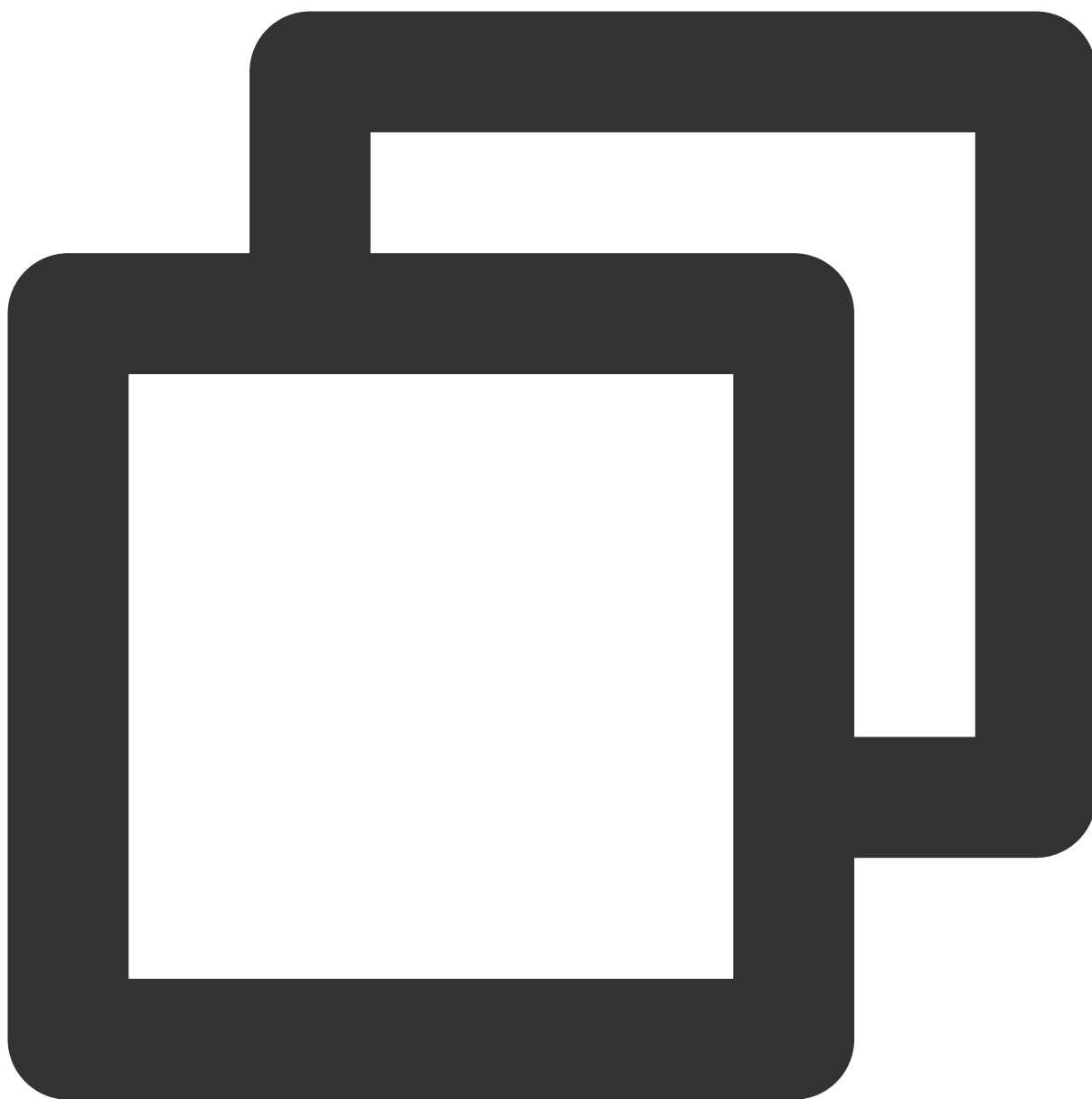
TKE Serverless 会为哪些 Ingress 创建 CLB 实例？

TKE Serverless 会为满足如下条件的 Ingress 创建 CLB 实例：

对 Ingress 资源的要求	其他说明
annotations 中含有如下键值对： kubernetes.io/ingress.class: qcloud	如果不希望 TKE Serverless 为 Ingress 创建 CLB 实例，例如希望使用 Nginx-ingress，则只需保证 annotations 中不包含前述键值对。

如何查看 TKE Serverless 为 Ingress 创建的 CLB 实例？

如果成功为 Ingress 创建 CLB 实例，TKE Serverless 会将 CLB 实例的 VIP 写入 Ingress 资源的 `status.loadBalancer.ingress` 中，并且将如下键值对写入 annotations 中：



```
kubernetes.io/ingress.qcloud-loadbalance-id: CLB 实例 ID
```

如需查看 TKE Serverless 为 Ingress 创建的 CLB 实例，具体步骤如下：

1. 登录容器服务控制台，选择左侧导航栏中的 [集群](#)。
2. 在集群列表页面，选择集群 ID 进入集群管理页面。
3. 在集群管理页面，选择左侧**服务与路由** > **Ingress**。
4. 在 “Ingress” 页面，查看 CLB 实例 ID 及其 VIP。如下图所示：



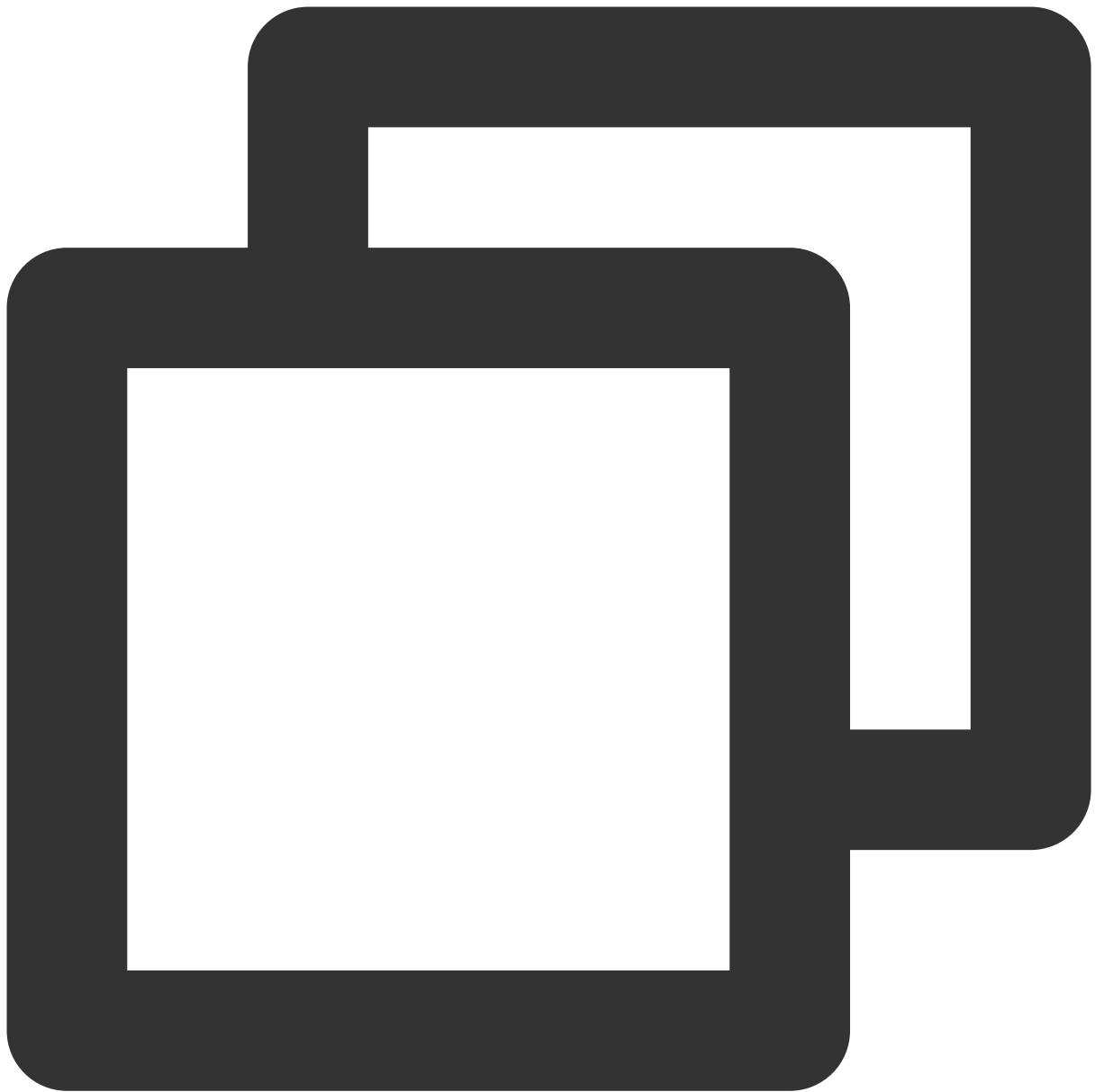
TKE Serverless 会为哪些 Service 创建 CLB 实例？

TKE Serverless 会为满足如下条件的 Service 创建 CLB 实例：

K8S 版本	对 Service 资源的要求
所有 TKE Serverless 支持的 K8S 版本	spec.type 为 LoadBalancer
魔改版 K8S（kubectl version 返回的 Server GitVersion 带有 "eks." 或 "tke." 后缀）	spec.type 为 ClusterIP，并且 spec.clusterIP 的值不是None（即非 Headless 的 ClusterIP 类型的 Service）
非魔改版 K8S（kubectl version 返回的 Server GitVersion 不带 "eks." 或 "tke." 后缀）	spec.type 为 ClusterIP，并且明确指定 spec.clusterIP 为空字符串（""）

注意：

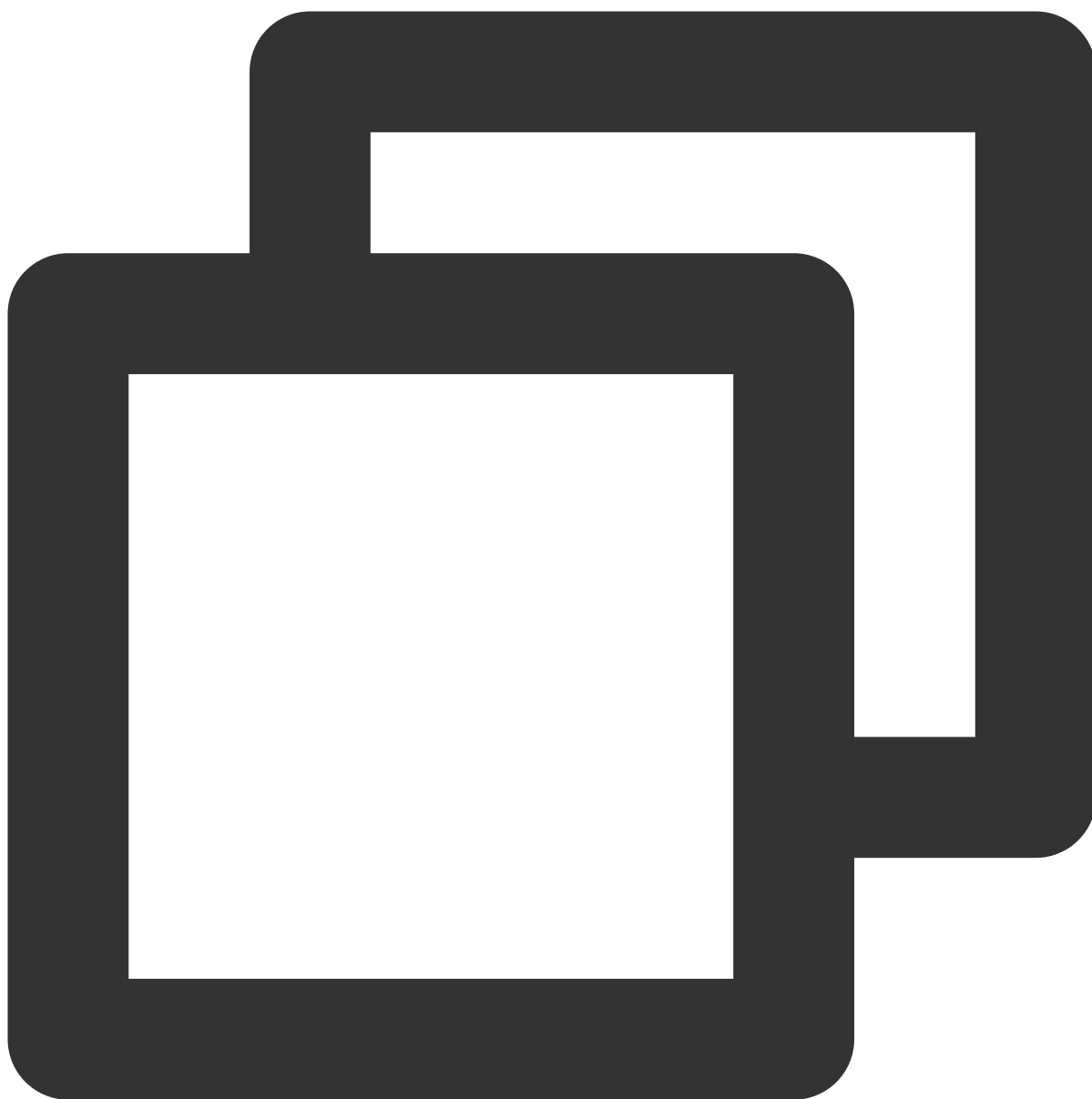
如果成功创建 CLB 实例，TKE Serverless 会将如下键值对写入 Service annotations 中：



```
service.kubernetes.io/loadbalance-id: CLB 实例 ID
```

如何查看 TKE Serverless 为 Service 创建的 CLB 实例？

如果成功为 Service 创建 CLB 实例，TKE Serverless 会将 CLB 实例的 VIP 写入 Service 资源的 `status.loadBalancer.ingress` 中，并且将如下键值对写入 annotations 中：



```
kubernetes.io/ingress.qcloud-loadbalance-id: CLB 实例 ID
```

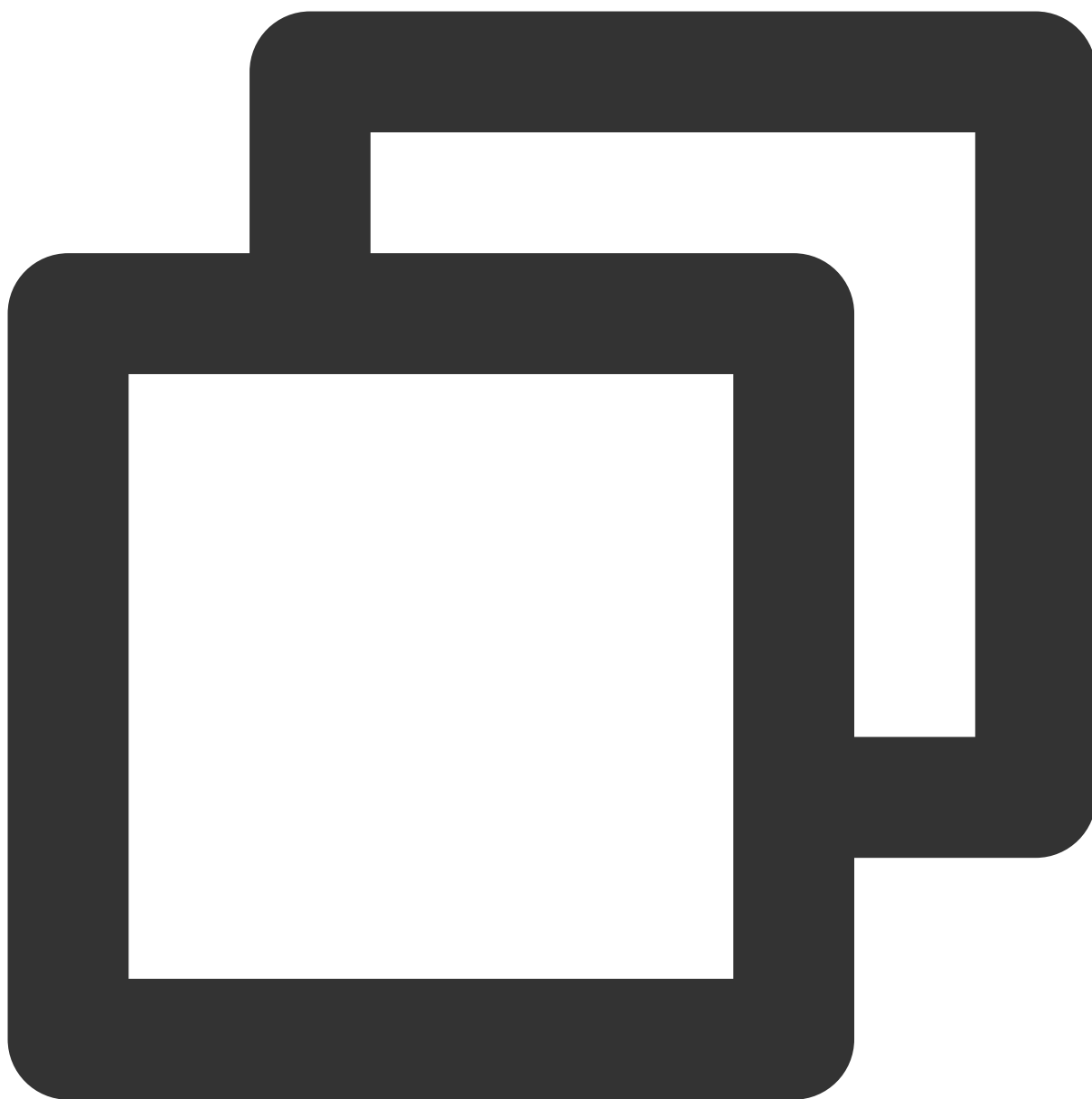
如需查看 TKE Serverless 为 Service 创建的 CLB 实例，具体步骤如下：

1. 登录容器服务控制台，选择左侧导航栏中的 [集群](#)。
2. 在集群列表页面，选择集群 ID 进入集群管理页面。
3. 在集群管理页面，选择左侧**服务与路由** > **Service**。
4. 在 **Service** 页面，查看 CLB 实例 ID 及其 VIP。如下图所示：



为什么 Service 的 ClusterIP 无效（无法正常访问）或没有 ClusterIP ？

对于 spec.type 为 LoadBalancer 的 Service，目前 TKE Serverless 默认不分配 ClusterIP，或者分配的 ClusterIP 无效（无法正常访问）。如果用户同时需要使用 ClusterIP 访问 Service，可以通过在 annotations 中加入如下键值对指示 TKE Serverless 基于内网 CLB 实现 ClusterIP：



```
service.kubernetes.io/qcloud-clusterip-loadbalancer-subnetid: Service CIDR 子网 ID
```

Service CIDR 子网 ID 在创建集群时指定，为 `subnet-*****` 字符串。您可在 CLB 基本信息页面中查看该子网 ID 信息。

注意：

只有使用魔改版 K8S（kubectl version 返回的 Server GitVersion 带有 "eks." 或 "tke." 后缀）的 TKE Serverless 集群才支持该特性。对于早期创建的、使用非魔改版 K8S（kubectl version 返回的 Server GitVersion 不带 "eks." 或 "tke." 后缀）的 TKE Serverless 集群，您需要升级 K8S 版本后使用该特性。

如何指定 CLB 实例类型（公网或内网）？

您可以通过容器服务控制台 或通过 kubectl 命令行工具指定 CLB 实例类型：

通过容器服务控制台操作

通过 kubectl 命令行工具操作

对于 Ingress，通过“网络类型”选择“公网”或“内网”：



对于 Service，通过“服务访问方式”控制，其中“VPC内网访问”对应内网 CLB 实例：



默认创建的 CLB 实例是“公网”类型。

如需创建“内网”类型的 CLB 实例，则需为 Service 或 Ingress 加上相应的 annotation：

资源类型	需要在 annotations 中添加的键值对
Service	service.kubernetes.io/qcloud-loadbalancer-internal-subnetid: 子网 ID
Ingress	kubernetes.io/ingress.subnetId: 子网 ID

注意：

子网 ID 是形如 subnet-***** 的字符串，并且该子网必须在创建集群时为“集群网络”指定的 VPC 中，该 VPC 信息可在腾讯云控制台集群“基本信息”中查询到。

如何指定使用已有 CLB 实例？

您可以通过容器服务控制台 或通过 kubectl 命令行工具指定使用已有 CLB 实例：

通过容器服务控制台 操作

通过 kubectl 命令行工具操作

在创建 Service 或 Ingress 时，可以选择“使用已有” CLB 实例。对于 Service，还可以在 Service 创建之后，通过“更新访问方式”切换到“使用已有” CLB 实例。

在创建 Service/Ingress 或变更 Service 时，为 Service 或 Ingress 加上相应的 annotation 即可：

资源类型	需要在 annotations 中添加的键值对
------	-------------------------

Service	service.kubernetes.io/tke-existed-lbid: CLB 实例 ID
Ingress	kubernetes.io/ingress.existLbid: CLB 实例 ID

注意：

“已有 CLB 实例” 不能是 “TKE Serverless 为 Service 或 Ingress 创建的 CLB 实例”，并且 TKE Serverless 不支持多个 Service/Ingress 共用同一个已有 CLB 实例。

如何查看 CLB 实例的访问日志？

仅7层 CLB 实例支持配置访问日志，但 TKE Serverless 为 Ingress 创建的7层 CLB 实例默认不开启访问日志。如需开启 CLB 实例的访问日志，可在 CLB 实例的详情页面进行操作。具体步骤如下：

1. 登录容器服务控制台，选择左侧导航栏中的 [集群](#)。
2. 在集群列表页面，选择集群 ID 进入集群管理页面。
3. 在集群管理页面，选择左侧**服务与路由 > Ingress**。
4. 在 “Ingress” 页面，选择 CLB 实例 ID 进入 CLB 基本信息页面。如下图所示：



5. 在 CLB 基本信息页面的“访问日志（七层）”中，单击



，在弹出窗口中进行开启。如下图所示：



为什么 TKE Serverless 没有为 Ingress 或 Service 创建 CLB 实例？

请参考 [TKE Serverless 会为哪些 Ingress 创建 CLB 实例](#) 及 [TKE Serverless 会为哪些 Service 创建 CLB 实例](#) 问题的回答，确认对应的资源是否具备对应的条件。如具备条件但未成功创建 CLB 实例，可通过 `kubectl describe` 命令查看“资源”相关事件。

通常 TKE Serverless 会输出相关的 Warning 类型事件。示例图如下，输出事件表明子网中已经无可用的 IP 资源，所以无法成功创建 CLB 实例。



如何在多个 Service 中使用相同的负载均衡？

TKE Serverless 集群默认多个 Service 不可共用同一个 CLB 实例。如果您希望 Service 复用其他 Service 占用的 CLB，请添加此 annotation 并将 value 填写为 "true"。 `service.kubernetes.io/qcloud-share-existed-lb: true`，关于此 annotation 的详细说明请参见 [Annotation 说明](#)。

为什么访问 CLB VIP 时失败？

请您按照以下步骤进行分析：

查看 CLB 实例类型

1. 在“Service”或“Ingress”页面，选择 CLB 实例 ID 进入 CLB 基本信息页面。如下图所示：



2. 您可在 CLB 基本信息页面中查看上述 CLB 实例的“实例类型”。

确认访问 CLB VIP 的环境是否正常

若 CLB 实例“实例类型”为内网，则其 VIP 只能在所属的 VPC 内访问。

由于 TKE Serverless 集群中 Pods 的 IP 是 VPC 内的弹性网卡的 IP，所以可以在 Pods 中访问集群内任何 Service 或 Ingress 的 CLB 实例的 VIP。

注意

通常 LoadBalancer 系统都存在回环问题（例如 [AzureLoad Balancer 问题排查指南](#)），请勿在工作负载所属的 Pods 中通过该工作负载（经 Service 或 Ingress）对外暴露的 VIP 访问该工作负载提供的服务。即 Pods 不要通过 VIP（包含“内网类型”及“外网类型”）访问 Pods 自己提供的服务。否则可能导致延迟增加，或者（在 VIP 对应的规则下只有一个 RS/Pod 时）访问不通。

若 CLB 实例“实例类型”为公网，则其 VIP 可以在有公网访问能力的环境中访问。

若要在集群内访问公网 VIP，请确保已经通过配置 NAT 网关或其他方式为集群开启了公网访问能力。

查看 CLB 下的 RS 包括（且仅包括）预期的 Pods 的 IP + 端口

您可在 CLB 管理页面中，选择[监听器管理](#)页面，查看（7层协议）转发规则、（4层协议）绑定的后端服务。其中的 IP 地址预期即为各个 Pod 的 IP。TKE Serverless 为某个 Ingress 创建的 CLB 示例图如下：



确认对应的 Endpoints 是否正常

如果已正确地为工作负载（Workload）设置了标签（Labels），并且正确地为 Service 资源设置了选择算符（Selectors），则在工作负载的 Pods 成功运行之后，即可通过 `kubectl get endpoints` 命令查看 Pods 被 K8S 列入到 Service 对应的 Endpoints 的就绪地址列表中。示例图如下：

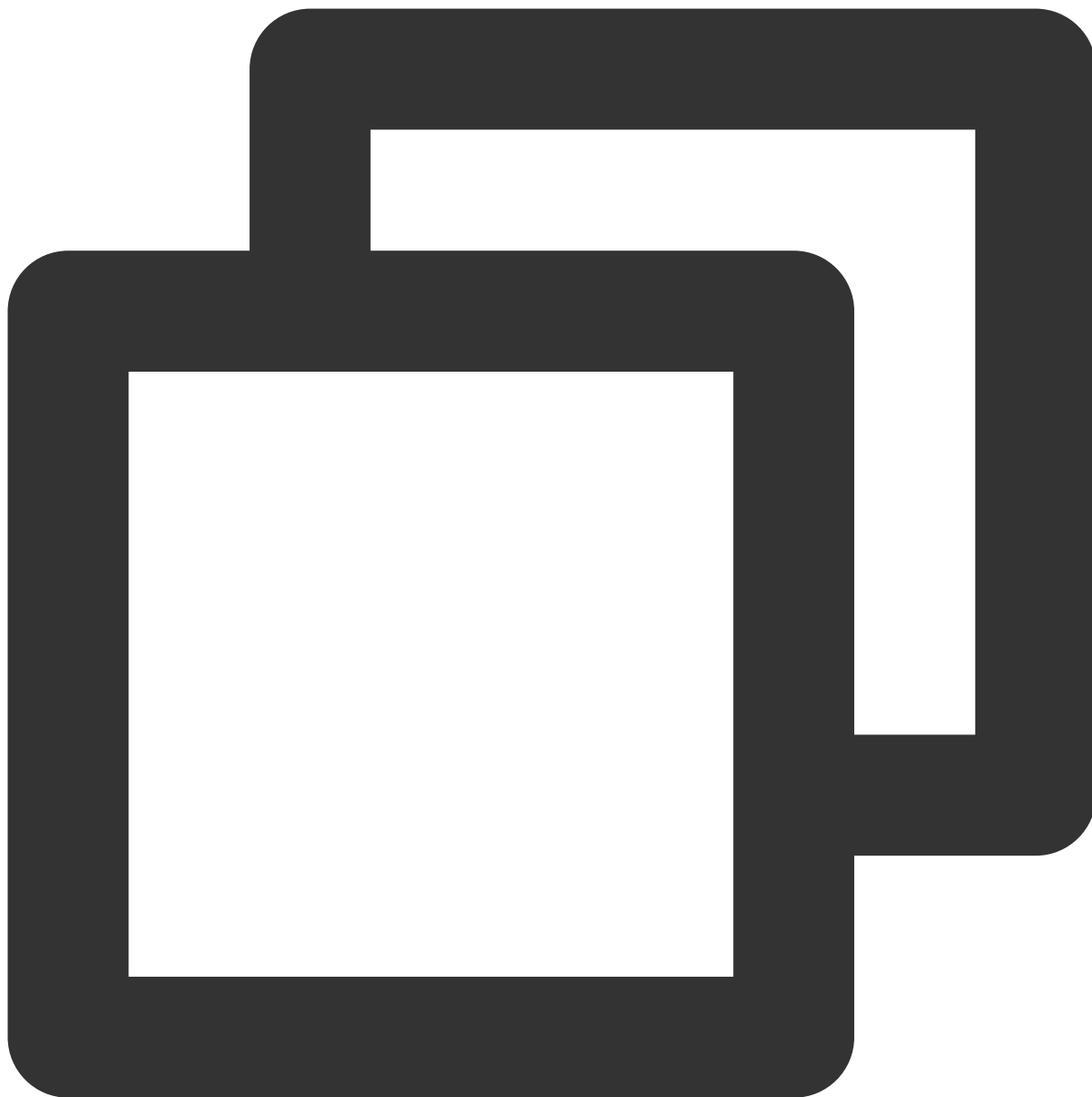


而已创建、但状态异常的 Pods 会被 K8S 列入到 Service 对应的 Endpoints 的未就绪地址列表中。示例图如下：



注意：

对于异常的 Pods，可通过 `kubectl describe` 命令查看异常的原因。示例命令如下：



```
kubectl describe pod nginx-7c7c647ff7-4b8n5 -n demo
```

确认 Pods 是否能够正常提供服务

即使 Running 状态的 Pods 也可能无法正常对外提供服务。例如，未监听指定的协议+端口、Pods 内部逻辑错误、处理过程阻塞等。您可通过 `kubectl exec` 命令登录至 Pod 内，并使用 `telnet/wget/curl` 命令或自定义的客户端工具直接访问 Pod IP+端口。若 Pod 内直接访问失败，则需要进一步分析导致 Pod 无法正常提供服务的原因。

确认 Pod 绑定的安全组是否放通 Pods 提供服务的协议和端口

安全组控制 Pods 的网络访问策略，如同 Linux 服务器中的 IPTables 规则。请结合实际情况进行查看：

通过容器服务控制台创建工作负载

使用 kubectl 命令创建工作负载

交互流程会强制要求指定一个安全组，TKE Serverless 将使用该安全组控制 Pods 的网络访问策略。用户选择的安全组会被存储在工作负载的 `spec.template.metadata.annotations`，最终添加到 Pods 的 `annotations` 中。

示例如下：



若您通过 kubectl 命令创建工作负载，且没有（通过 annotations）为 Pods 指定安全组，则 TKE Serverless 会使用账号下的同地域的默认项目的 `default` 安全组。查看步骤如下：

1. 登录私有网络控制台，选择左侧导航栏中的 [安全组](#)。
2. 在“安全组”页面上方，选择同地域的默认项目。
3. 可在列表中查看 `default` 安全组，可单击 [修改规则](#) 查看详情。如下图所示：



联系我们

若至此仍未找到问题原因，您可通过 [提交工单](#) 联系 TKE 团队解决问题。

超级节点相关

最近更新时间：2023-02-14 14:40:42

如何禁止 Pod 调度到某个按量计费超级节点？

默认情况下，TKE 普通集群添加了按量计费超级节点节点池后，会在 Node 资源不足时，自动向按量计费超级节点调度 Pod。Serverless 集群则会自动在多个按量计费超级节点随机调度 Pod。

此时如果您不希望向某按量计费超级节点（代表某子网/可用区）调度，可通过以下两种方式封锁按量计费超级节点实现禁止调度：

- 通过 [容器服务控制台](#) 对节点进行**封锁**操作，详情可参见 [封锁节点](#)。
- 通过命令行执行如下命令实现禁止调度：

```
$kubectl cordon $按量计费超级节点名称
```

如何禁止 TKE 普通集群在资源不足时自动调度到按量计费超级节点？

需要在 kube-system 命名空间下，新建名为 eks-config 的 configmap。执行命令如下：

```
$kubectl create configmap eks-config --from-literal=AUTO_SCALE_EKS=false
```

将 AUTO_SCALE_EKS 的 value 设置为 false，即可关闭 TKE 普通集群向按量计费超级节点的自动调度机制。

如何手动调度 Pod 到按量计费超级节点？

默认按量计费超级节点会自动添加 Taints 以降低调度优先级，如需手动调度 Pod 到按量计费超级节点或指定按量计费超级节点，通常需要为 Pod 添加对应的 Tolerations。但并非所有的 Pod 均可以调度到按量计费超级节点上，详情请参见 [超级节点调度说明](#)。为方便使用，您可以在 Pod Spec 中指定 nodeSelector。示例如下：

```
spec:
  nodeSelector:
    node.kubernetes.io/instance-type: eklet
```

TKE 的管控组件会判断该 Pod 是否可以调度到按量计费超级节点，若不支持则不会调度到按量计费超级节点。

如何强制调度 Pod 到按量计费超级节点，无论按量计费超级节点是否支持该 Pod？

注意：

强制调度到按量计费超级节点，Pod 可能会创建失败，失败原因可参见 [如何手动调度 Pod 到按量计费超级节点](#)。

如需强制调度 Pod 到按量计费超级节点，除了为 Pod 指定 `nodeSelector` 或 `nodeName`，也必须添加对应的 `Tolerations`。

1. 为 Pod Spec 中指定 `nodeSelector`。示例如下：

```
spec:
  nodeSelector:
    node.kubernetes.io/instance-type: eklet
```

或在 Pod Spec 指定 `nodeName`。示例如下：

```
spec:
  nodeName: $按量计费超级节点名称
```

2. 为 Pod 添加 `Tolerations`。示例如下：

```
spec:
  tolerations:
    - effect: NoSchedule
      key: eks.tke.cloud.tencent.com/eklet
      operator: Exists
```

如何自定义按量计费超级节点 DNS？

TKE Serverless 集群支持按量计费超级节点特性，您可通过在 `yaml` 中定义 `annotation` 的方式，实现自定义 DNS 等能力。详情可参见 [按量计费超级节点 annotation 说明](#)。

镜像仓库相关

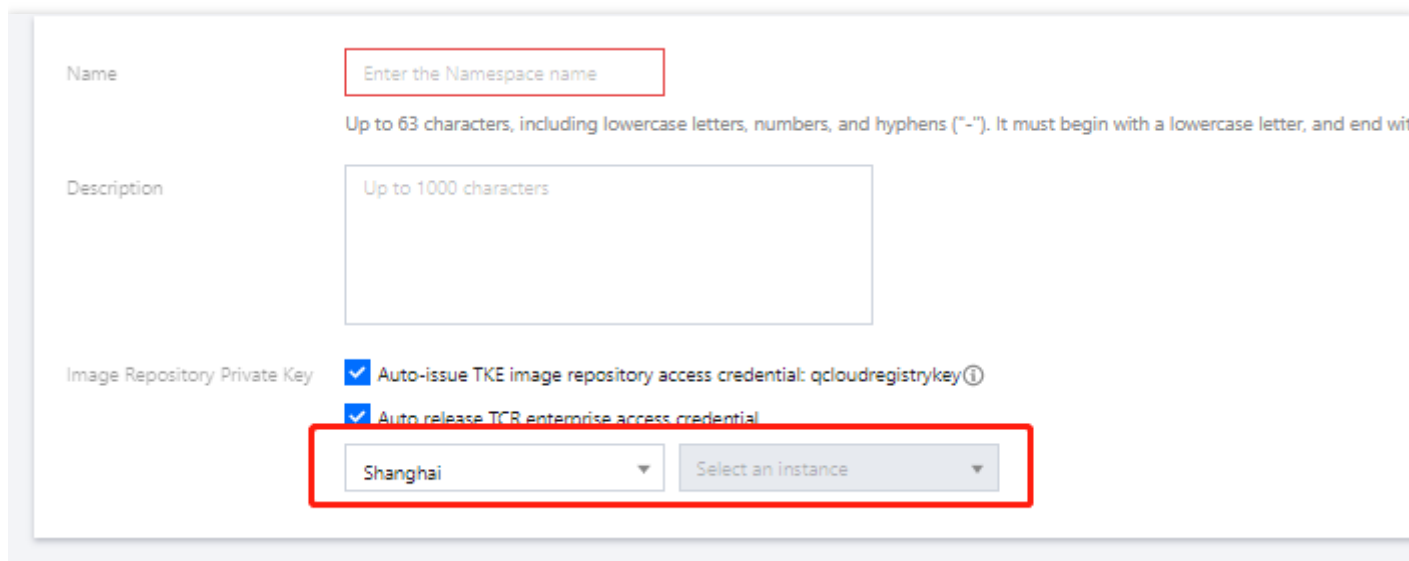
最近更新时间：2023-05-23 10:21:14

Serverless 集群如何使用容器镜像服务 TCR？

Serverless 集群如需使用容器镜像服务 TCR，需要确保 [已选择对应的镜像访问凭证](#) 和 [Serverless 集群到镜像服务网络打通](#)。

确保已选择对应的镜像访问凭证

容器镜像默认私有，因此在创建工作负载时，需选择实例对应的镜像访问凭证。



The screenshot shows the 'Create TCR Access Credential' form in the Tencent Cloud console. The 'Name' field is highlighted with a red box. The 'Image Repository Private Key' section has two checked options: 'Auto-issue TKE image repository access credential: qcloudregistrykey' and 'Auto-release TCR enterprise access credential'. Below these, a red box highlights the 'Shanghai' region dropdown and the 'Select an instance' dropdown.

您可参考以下步骤创建镜像访问凭证：

1. 登录 [容器服务控制台](#)。
2. 在集群列表页面，单击您需要创建访问凭证的Serverless 集群名称，进入Serverless 集群详情页。
3. 选择左侧的“命名空间”，单击**新建**。
4. 在新建“新建Namespace”页中，勾选**自动下发容器镜像服务企业版访问凭证**。
5. 单击**创建Namespace**进行创建。此时在新建的命名空间上，即可选择镜像的访问凭证。

确保 Serverless 集群到镜像服务网络打通

Serverless 集群与容器镜像服务之间的网络默认不通，因此在拉取镜像时会报网络不通的错误：

```
dial tcp x.x.x.x:443: i/o timeout
```

解决步骤

处理方式有以下两种：

处理方式	说明
方式1：内网访问方式（推荐）	容器镜像服务新建内网访问链路，并配置内网域名解析，Serverless 集群通过新建的内网链路访问容器镜像服务。
方式2：外网访问方式	Serverless 集群开启外网访问，通过公网访问容器镜像服务，同时容器镜像服务也需允许公网访问。

方式1：内网访问方式（推荐）

1. 新建内网访问链路

- 登录 [容器镜像服务控制台](#)，选择左侧导航栏中的[访问控制](#)>[内网访问](#)。
- 在“内网访问”页面中选择对应的地域和实例，单击**新建**。
- 在弹出的“新建内网访问链路”窗口中，配置私有网络及子网信息。

2. 配置完成并单击**确定**，该内网访问链路将启动新建。

3. 开启域名内网解析。

容器镜像服务使用的域名为“<tcr-name>.tencentcloudcr.com”，此域名在 VPC 内的解析需额外开通。若不开通，“<tcr-name>.tencentcloudcr.com”的域名依旧解析成公网 IP，无法解析成内网 IP，内网访问将会失败。

注意：

新建链路后，需要等待后台生成内网解析 IP，IP 生成后才可开启以下按钮。

方式2：外网访问方式

1. 确保容器镜像服务开启公网访问

- [容器镜像服务控制台](#)，选择左侧导航栏中的[访问控制](#)>[公网访问](#)。
- 在“公网访问”页面中选择对应的地域和实例。
- 为对应的容器镜像服务实例开启公网访问。

测试体验阶段，可将公网 IP 地址段设置为：0.0.0.0/0。后续运行阶段，可将下文步骤2中涉及的 NAT 网关出口的弹性 IP 添加到公网白名单。

2. 为 Serverless 集群开启公网访问。

Serverless 集群默认不开通外网访问，需通过 NAT 网关进行外网访问，详情可参见 [通过 NAT 网关访问外网](#)。

完成 NAT 网关配置后，将 Serverless 集群所在子网关联至 NAT 网关的路由表，并确保 NAT 网关出口的弹性 IP 在访问容器镜像服务的白名单里（[请参见上文 步骤1](#)），Serverless 集群即可正常访问容器镜像服务，从公网拉取镜像。

常见报错指引

报错：Error:\sfailed\sto\sdo\srequest:Head\s"xxxx/manifests/late-st":\sdial\stcp\sxxx:443:\si/o\stimeout

含有“443: i/o timeout”报错，大部分情况是因为 TKE Serverless 到 TCR 网络不通。

请参见上文 [确保 Serverless 集群到镜像服务网络打通](#)，选择其中一种访问方式，打通 TKE Serverless 到 TCR 的网络。

注意

“<tcrcr-name>.tencentcloudcr.com”的域名默认会解析成公网 IP，报错时请注意识别“dial tcp xxx”里的 IP 地址是公网还是内网，按实际情况处理。

报错：

Error:\scode\s=\sUnknown,\spull\saccess\sdenied,\srepository\sdoes\snot\sexist\sor\smy\srequire\sauthorization:\sserver\smessage:\sinsufficient_scope:\sauthorization\sfailed

含有“insufficient_scope: authorization failed”报错，说明 TKE Serverless 到 TCR 网络已经开通，但权限不足。可能原因为命名空间不存在、密钥不正确、密钥不适用正在拉取的镜像等。

报错：Error:\scode\s=\sNotFound,\sfailed\sto\sresolve\sreference\s"xxx:xxx":\snot\sfound

含有“not found”报错，说明镜像不存在。

其他常见报错指引请参见 [容器镜像服务常见问题指引](#)。

Serverless 集群如何使用自建的自签名镜像仓库或 HTTP 协议镜像仓库？

问题描述

在 Serverless 集群使用自建镜像仓库的镜像创建工作负载，可能会遇到“ErrImagePull”报错，拉取镜像失败。如下图所示：

2021-04-02 16:55:20	2021-04-02 16:55:34	Warning	Pod	busybox-764db787c8-4w894.1671fa3df257b8b	Failed	Error: ErrImagePull	2
2021-04-02 16:55:20	2021-04-02 16:55:34	Warning	Pod	busybox-764db787c8-4w894.1671fa3df252cad	Failed	Failed to pull image "10.16.100.174:5000/busybox": error: dial tcp 10.16.100.174:5000: connect: connection refused	2

问题原因

通常在保证网络连通性的前提下，可能有以下两个原因导致该问题：

- 自建镜像仓库采用 HTTPS 协议，但是 HTTPS 协议证书是自签名。
- 自建镜像仓库采用 HTTP 协议。

上述两个问题都可以通过在工作负载 Yaml 配置里的 PodTemplate 中添加 Annotation 来解决。

解决步骤

HTTPS 自签名镜像仓库

如果自建的镜像仓库是 HTTPS 协议且证书是自签名。则需要在 PodTemplate 中，添加如下 Annotation，使其跳过证书检验：

eks.tke.cloud.tencent.com/registry-insecure-skip-verify: 镜像仓库地址（多个用“,”隔开，或者填写 all）

如下图所示：

```
spec:
  replicas: 1
  selector:
    matchLabels:
      k8s-app: sampleapp
      qcloud-app: sampleapp
  template:
    metadata:
      annotations:
        eks.tke.cloud.tencent.com/registry-insecure-skip-verify: your-registry-url
      creationTimestamp: null
      labels:
        k8s-app: sampleapp
        qcloud-app: sampleapp
```

说明：

如果 Pod 内多个容器的镜像从不同仓库拉取，可填写多个镜像仓库地址，中间用 “,” 隔开。或填写 “all”，表示 Pod 内所有容器镜像的相关仓库均跳过证书检验。

HTTP 协议镜像仓库

说明：

默认 Serverless 集群运行时会使用 HTTPS 协议拉取镜像，因此如果镜像仓库支持的协议为 HTTP，则也需要通过 Annotation 进行说明。

通过命令 “\$kubectl describe pod \$podname”，发现输出 “**http: server gave HTTP response to HTTPS client**” 报错，则说明所访问的镜像仓库使用 HTTP 协议。如下图所示：

Events:	Type	Reason	Age	From	Message
	Normal	Scheduled	11m	default-scheduler	Successfully assigned default/busybox-764db787c8-4w894 to eklet-subnet-es398vv5
	Normal	Starting	11m	eklet, eklet	Starting pod sandbox eks-71cpsr8b
	Normal	Starting	10m	eklet, eklet-subnet-es398vv5	Sync endpoints
	Normal	Pulling	9m12s (x4 over 10m)	eklet, eklet-subnet-es398vv5	Pulling image "10.16.100.174:5000/busybox"
	Warning	Failed	9m12s (x4 over 10m)	eklet, eklet-subnet-es398vv5	Failed to pull image "10.16.100.174:5000/busybox": rpc error: code = Unknown desc = failed to pull and unpack image "10.16.100.174:5000/busybox:latest": failed to resolve reference "10.16.100.174:5000/busybox:latest": failed to do request: Head "https://10.16.100.174:5000/v2/busybox/manifests/latest": http: server gave HTTP response to HTTPS client
	Warning	Failed	9m12s (x4 over 10m)	eklet, eklet-subnet-es398vv5	Error: ErrImagePull
	Normal	BackOff	8m58s (x6 over 10m)	eklet, eklet-subnet-es398vv5	Back-off pulling image "10.16.100.174:5000/busybox"
	Warning	Failed	38s (x42 over 10m)	eklet, eklet-subnet-es398vv5	Error: ImagePullBackOff

解决此问题需要在 PodTemplate 中，添加如下 Annotation，使其使用 HTTP 协议访问镜像仓库：

eks.tke.cloud.tencent.com/registry-http-endpoint: 镜像仓库地址(多个用“,” 隔开，或者填写all)

如下图所示：

```
spec:
  replicas: 1
  selector:
    matchLabels:
      k8s-app: sampleapp
      qcloud-app: sampleapp
  template:
    metadata:
      annotations:
        eks.tke.cloud.tencent.com/registry-http-endpoint: your-registry-url
      creationTimestamp: null
      labels:
        k8s-app: sampleapp
        qcloud-app: sampleapp
```

说明：

如果 Pod 内多个容器的镜像从不同仓库拉取，可填写多个镜像仓库地址，中间用“,” 隔开。或填写“all”，表示 Pod 内所有容器镜像均采用 HTTP 协议拉取。

镜像仓库地址填写说明

以上两个 Annotation 均涉及镜像仓库地址的填写，多个仓库地址可用“,” 隔开。

注意：

如果镜像仓库有端口号，则需要带上端口号。

例如，镜像地址为 10.16.100.174:5000/busybox:latest，则 Annotation 的 value 应填为 “10.16.100.174:5000”，即 “eks.tke.cloud.tencent.com/registry-insecure-skip-verify: 10.16.100.174:5000”。

镜像仓库类

最近更新时间：2020-04-21 12:36:35

开通镜像仓库常见问题

命名空间有什么作用？

命名空间是标识用户私人镜像的地址前缀。

镜像仓库的账户是什么？

默认是用户的腾讯云账号（QQ 号）。

开通时创建的密码忘记了怎么办？

可以通过控制台重置密码。

创建镜像常见问题

创建镜像有配额限制么？

默认配额是单个地域可创建500个镜像仓库，单个镜像仓库可创建100个镜像版本。

如果您需要更多的配额项数量，可通过 [配额申请工单](#) 提出配额申请。

创建的镜像可以分享给其他用户么？

暂不提供该功能。

创建的镜像如何使用？

需要先上传可用的镜像版本，通过具体的镜像版本来创建服务。

删除镜像常见问题

如何删除镜像某一个版本？

直接在控制台指定删除某一个具体的版本。

在镜像列表删除镜像会删除该镜像的所有版本么？

是的，删除镜像时会删除该镜像的所有镜像版本，请提前备份好数据。

镜像构建常见问题

使用源码构建镜像时 Dockerfile 路径与构建目录该怎么填？

- 若不填写路径，系统使用以下默认值：
 - Dockerfile 路径默认值：代码仓库根目录下的 Dockerfile（`Dockerfile`）。
 - 构建目录默认值：代码仓库根目录（`./`）。
- 若需填写路径，则填写以项目为根路径的相对路径。如下图所示：

Build Config

Image Address

ccr.ccs.tencentyun.com/test-yunx/helloworld

Code source

 Github

 Gitlab

Image Tag Naming Rules

Please enter the naming rule

- ☐ Branch/label - ☐ Update Time - ☐ Commit No.

Custom prefix, supports variables in the format of \$(Foo)

Overwrite the image tag

Please enter the tag

The generated image also contains the tag

Dockerfile path

qcloud/Dockerfile

Path of the Dockerfile in the code source

Building Directory

qcloud

The working directory for building, should be a relative path

Building Parameters

[Add a variable](#)

Complete

源码构建功能使用 Dockerfile 路径及构建目录细节是什么？

源码构建功能首先 clone 用户指定的仓库，切换到相应的分支（branch）或标签（tag），然后在代码仓库根目录执行 `docker build -f $DOCKERFILE_PATH $WORKDIR` 命令构建出容器镜像。

Dockerfile 文件中的源路径应该怎么写？

对于 `COPY`，`ADD` 等涉及源路径的命令，源路径应该写成相对于构建目录的相对路径。

远程终端类

最近更新时间：2020-04-26 16:18:05

容器里面没有 bash，怎么办？

如果发现没有 **bash**，您可以在命令行中输入您想执行的命令，屏幕会显示该命令的返回结果。您可以将命令行看做一个缺少自动补全等其他功能的精简版 **bash**。建议您先执行安装 **bash** 的命令，再执行后续操作。

为什么运行 apt-get 安装软件如此之慢？

如果您在使用 `apt-get` 安装软件速度过慢，有可能是因为机器访问国外软件源速度过慢引起。我们根据不同操作系统的容器为您提供对应的提速方案，您可根据实际情况进行操作：

注意事项

选择提速解决方案之前，请您准确识别容器操作系统，避免操作无法正常进行。您可参考以下方式进行识别：

- Ubuntu 系统：执行命令 `cat /etc/lsb-release`，验证是否存在 `/etc/lsb-release` 文件。
- CentOS 系统：执行命令 `cat /etc/redhat-release`，验证是否存在 `/etc/redhat-release` 文件。
- Debian 系统：执行命令 `cat /etc/debian_version`，验证是否存在 `/etc/debian_version` 文件。

">

解决方案

Ubuntu 16.04系统

对于系统为 Ubuntu 16.04的容器，您可以在终端运行以下命令，将 **apt** 的源设置为腾讯云的源服务器：

```
cat << EOF > /etc/apt/sources.list
deb http://mirrors.tencentyun.com/ubuntu/ xenial main restricted universe multiverse
deb http://mirrors.tencentyun.com/ubuntu/ xenial-security main restricted universe multiverse
deb http://mirrors.tencentyun.com/ubuntu/ xenial-updates main restricted universe multiverse
deb-src http://mirrors.tencentyun.com/ubuntu/ xenial main restricted universe multiverse
deb-src http://mirrors.tencentyun.com/ubuntu/ xenial-security main restricted universe multiverse
deb-src http://mirrors.tencentyun.com/ubuntu/ xenial-updates main restricted universe multiverse
EOF
```

CentOS 7系统

对于系统为 CentOS 7 的容器，您可以在终端执行以下操作，直接修改源地址提高安装速度：

1. 将以下代码复制并粘贴至容器内运行：

```
cat << EOF > /etc/yum.repos.d/CentOS-Base.repo
[os]
name=Qcloud centos os - \${basearch}
baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/os/\${basearch}/
enabled=1
gpgcheck=1
gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
[updates]
name=Qcloud centos updates - \${basearch}
baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/updates/\${basearch}/
enabled=1
gpgcheck=1
gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
#[centosplus]
#name=Qcloud centosplus - \${basearch}
#baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/centosplus/\${basearch}/
#enabled=1
#gpgcheck=1
#gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
#[cloud]
#name=Qcloud centos contrib - \${basearch}
#baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/cloud/\${basearch}/openstack-kilo/
#enabled=1
#gpgcheck=1
#gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
#[cr]
#name=Qcloud centos cr - \${basearch}
#baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/cr/\${basearch}/
#enabled=1
#gpgcheck=1
#gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
[extras]
name=Qcloud centos extras - \${basearch}
baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/extras/\${basearch}/
enabled=1
gpgcheck=1
gpgkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
#[fasttrack]
#name=Qcloud centos fasttrack - \${basearch}
#baseurl=http://mirrors.tencentyun.com/centos1/\${releasever}/fasttrack/\${basearch}/
```

```
#enabled=1
#gpcheck=1
#gpkey=file:///etc/pki/rpm-gpg/RPM-GPG-KEY-CentOS-7
ENDOF
```

2. 执行以下命令，清空并重建 YUM 缓存。

```
yum clean all && yum clean metadata && yum clean dbcache && yum makecache
```

Debian 系统

对于系统为 Debian 的容器，您可以在终端运行以下命令，将 apt 的源设置为腾讯云的源服务器：

```
cat << ENDOF > /etc/apt/sources.list
deb http://mirrors.tencentyun.com/debian stretch main contrib non-free
deb http://mirrors.tencentyun.com/debian stretch-updates main contrib non-free
deb http://mirrors.tencentyun.com/debian-security stretch/updates main

deb-src http://mirrors.tencentyun.com/debian stretch main contrib non-free
deb-src http://mirrors.tencentyun.com/debian stretch-updates main contrib non-free
e
deb-src http://mirrors.tencentyun.com/debian-security stretch/updates main
ENDOF
```

问题总结

在容器中直接修改源地址为临时解决方案，当容器被重新调度后，您所作的修改将会失效，建议您在创建镜像时解决该问题。具体的操作方法如下：

修改创建容器镜像的 Dockerfile，在 Dockerfile 的 RUN 字段中，添加对应操作系统 [解决方案](#) 中提供的信息修改源地址。例如，在一个基于 Ubuntu 系统的镜像的 Dockerfile 中，加入以下内容：

```
RUN cat << ENDOF > /etc/apt/sources.list
deb http://mirrors.tencentyun.com/ubuntu/ xenial main restricted universe multiverse
deb http://mirrors.tencentyun.com/ubuntu/ xenial-security main restricted universe multiverse
deb http://mirrors.tencentyun.com/ubuntu/ xenial-updates main restricted universe multiverse
#deb http://mirrors.tencentyun.com/ubuntu/ xenial-proposed main restricted universe multiverse
#deb http://mirrors.tencentyun.com/ubuntu/ xenial-backports main restricted universe multiverse
deb-src http://mirrors.tencentyun.com/ubuntu/ xenial main restricted universe multiverse
deb-src http://mirrors.tencentyun.com/ubuntu/ xenial-security main restricted universe multiverse
```

```
verse multiverse
deb-src http://mirrors.tencentyun.com/ubuntu/ xenial-updates main restricted universe multiverse
#deb-src http://mirrors.tencentyun.com/ubuntu/ xenial-proposed main restricted universe multiverse
#deb-src http://mirrors.tencentyun.com/ubuntu/ xenial-backports main restricted universe multiverse
ENDOF
```

当登录容器后，发现没有 vim，netstat 等工具，怎么办？

您可以通过 `apt-get install vim`，`net-tools` 等命令下载您所需要的工具（CentOS 下执行 `yum install vim`）。

为什么运行 apt-get install 命令，提示找不到工具？

您可以通过以下操作，安装软件。

1. 执行以下命令，升级软件列表。

```
apt-get update
```

2. 执行以下命令，安装软件（CentOS 下执行 `yum updateinfo`）。

```
apt-get install
```

如果想在容器中使用自制的工具，怎么办？

您可以在进入远程终端页面后，单击右下方的文件助手，即可进行上传和下载操作。

如何拷贝现场文件，例如 dump 或者日志到本地分析？

您可以在进入远程终端页面后，单击右下方的文件助手，即可进行上传和下载操作。

为什么用不了文件上传到容器或者下载到本地功能？

可能因为您的容器镜像里没有安装 `tar` 程序，您可以通过 `apt-get install vim`，`net-tools` 等命令（CentOS 下执行 `yum install vim`）先安装 `tar` 程序再重试。

为什么之前安装的工具不见了？

可能因为 `kubernetes` 重新调度您的容器，调度过程中会拉取镜像生成新的容器，如果镜像里面没有您之前安装的工具，新的容器是不会包含这些工具的。建议您在制作镜像时，安装一些常用的排错工具。

怎么复制控制台里的文字？

只要选中您想复制的内容，即可复制被选中的文字。

怎么粘贴复制好的文字？

同时按下 Shift + Insert 即可。

为什么会断开链接？

可能因为您在腾讯云其他页面对容器、云服务器进行操作更改了容器的状况，也有可能是长时间（3分钟）不进行任何操作，服务器断开了这个链接。

运行 top 命令等出现 TERM environment variable not set 的错误，怎么办？

执行命令 `export TERM linux` 即可。

为何进入绝对路径较长的目录后，bash 提示符只显示 “<” 和部分路径？

因为默认的 bash 提示符被设置为显示 “用户名@主机名 当前目录”。如果当前路径长于一定长度，bash 默认会显示 “<” 与路径的最后一部分。

事件类

最近更新时间：2022-04-18 15:45:23

Back-off restarting failed docker container

说明：正在重启异常的 Docker 容器。

解决方法：检查镜像中执行的 Docker 进程是否异常退出，若镜像内并无持续运行的进程，可在创建服务的页面中添加执行脚本。

fit failure on node: Insufficient cpu

说明：集群 CPU 不足。

解决方法：原因是节点无法提供足够的计算核心，请在服务页面修改 CPU 限制或者对集群进行扩容。

no nodes available to schedule pods

说明：集群资源不足。

解决方法：原因是没有足够的节点用于承载实例，请在服务页面修改服务的实例数量，修改实例数量或者 CPU 限制。

pod failed to fit in any node

说明：没有合适的节点可供实例使用。

解决方法：原因是服务配置了不合适的资源限制，导致没有合适的节点用于承载实例，请在服务页面修改服务的实例数量或者 CPU 限制。

Liveness probe failed

说明：容器健康检查失败

解决方法：检查镜像内容容器进程是否正常，检查检测端口是否配置正确。

Error syncing pod, skipping

Error syncing pod, skipping failed to "StartContainer" for with CrashLoopBackOff: "Back-off 5m0s restarting failed container

说明：容器进程崩溃或退出。

解决方法：检查容器内是否有持续运行的前台进程，若有检查其是否有异常行为。详情请参考 [构建 docker 镜像指南](#)。

如果以上解决方法未能解决您的问题，您可以 [提交工单](#) 来寻求帮助