

腾讯云数据库 TCHouse-C

实践教程

产品文档



【版权声明】

©2013–2026 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】

及其他腾讯云服务相关的商标均为腾讯集团下的相关公司主体所有。另外，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

实践教程

腾讯云数据仓库 TCHouse-C 数据实时更新使用指引

使用 MySQL 关联 HBase 维表数据到 ClickHouse

腾讯云数据仓库 TCHouse-C 内核使用建议与规范

实践教程

腾讯云数据仓库 TCHouse-C 数据实时更新使用指引

最近更新时间：2026-06-18 17:33:16

功能介绍

腾讯云数据仓库 TCHouse-C 支持数据实时更新功能，即支持完整 UPSERT 语义。当使用 UPSERT 功能时，数据更新和删除时同步生效。

注意：

此功能为腾讯云数据仓库 TCHouse-C 新功能，如有需要，请联系 [在线客服](#) 咨询使用。

适用场景

腾讯云数据仓库 TCHouse-C 的数据实时更新是同步生效的，具体适用场景为：

- 对于某些满足条件的行，修改 1 个或者多个字段的值（排序键，分区键以及唯一键不可修改）。
- 点更新或删除，通过 WHERE 条件将更新或者删除的数据限制在较少的行。
- UPSERT 功能只在具有 UNIQUE KEY 定义的 MergeTree/ReplicatedMerge 表引擎中支持。其他类似于 replacingMergeTree 或者 VersionedCollapsingMergeTree 则不支持备份。

性能

UPSERT 的性能与待更新数据的行数量、待更新数据跨分区分布情况以及数据过滤条件相关。具体而言：

- 单次更新动作中，待更新数据行数越小，更新代价越小；
- 单次更新动作中，待更新数据所属分区数量越少（最好是都属于同一个分区），更新代价越小；
- 单次更新中，数据过滤（WHERE 条件）执行代价越小，更新代价越小。（通常建议 WHERE 条件包含排序键字段）。

使用注意事项

相比开源 ClickHouse，在使用实时更新功能时，需要考虑其特殊性。

避免并发更新相同的 UNIQUE KEY

腾讯云数据仓库 TCHouse-C 使用的是 Merge-On-Write 实现数据按行更新。在一次更新动作中，包含了数据读取和数据写入两个步骤。如果同时对相同的 UNIQUE KEY 进行并发更新操作，会存在结果不确定性问题。

针对该问题，在后续版本中，会引入用户自定义版本字段来在业务层面解决数据冲突问题。

需要使用更多内存资源

腾讯云数据仓库 TCHouse-C 为了实现数据唯一性，使用了 HASH 索引，需要占用一定量的内存。对于 20 亿行的数据，以 UInt64 字段作为 key，占用内存在 70GB 左右，尽量使用高内存机型，不然如果因为数据太多导致内存耗尽，可能存在无法启动的风险。

避免大批量更新

腾讯云数据仓库 TCHouse-C 通过 Merge-On-Write 实现数据更新，大批量更新会放大数据写入节点的延迟。业务层面尽可能采取小批量更新，例如点更新。

UNIQUE KEY 数据唯一性约束只在节点内有效

腾讯云数据仓库 TCHouse-C 在节点内部维护了索引，确保 UNIQUE KEY 唯一性。若相同数据分布在不同节点，集群无法确保其唯一性。业务层面需要将相同的 KEY 写到同一个 SHARD。谨慎进行水平扩容，水平扩容后会改变节点 key 映射关系。

避免单 SQL 执行多分区更新

为了避免系统中存在过的小文件，以及降低更新代价，建议单个 SQL 只执行一个分区中的数据更新。

操作步骤

创建表

如下示例：

```
CREATE TABLE IF NOT EXISTS realtime_insert (
  id UInt32,
  value1 UInt32,
  value2 UInt32
) ENGINE = MergeTree()
UNIQUE KEY id
ORDER BY id;
```

UNIQUE KEY 定义表示该表将会被创建为 upsert MergeTree 类型表，id 则表示唯一键，用于表数据生成唯一索引，以及数据去重。

如果不填写 UNIQUE KEY 定义，则表示该表没有 upsert 特性，和普通 MergeTree 表没有区别。

注意：

- UNIQUE KEY 选定唯一键推荐仅使用1个字段，选取多个字段作为唯一键内存消耗更多。
- UNIQUE KEY 选定的唯一键必须是 ORDER BY 或者 PRIMARY KEY 选取 column 集合的第一个 column，或者前缀集合。

- UNIQUE KEY 支持的 column 类型仅为 Int8/16/32/64/128, UInt8/16/32/64/128, UUID, IPV4, String 其余类型的 column 不能作为 UNIQUE KEY, 并且 column 类型不能是 Nullable, 也不能是表达式。
- 指定 UNIQUE KEY 的表不支持 projection 特性。
- 指定 UNIQUE KEY 的表, Engine 仅支持 MergeTree 和 ReplicatedMergeTree(beta), 其余 Engine 类型不支持。
- 指定了 UNIQUE KEY 的 ReplicatedMergeTree 表不支持 ttl 特性, 也不支持 drop part 和 drop partition。
- 指定 UNIQUE KEY 的表不能对表进行 column 的操作, 例如增加新列, 删除旧列, 更改列类型, 删除列。

INSERT INTO 实时更新

指定了 UNIQUE KEY 的表, 支持写入时实时去重:

```
insert into realtime_insert values (1,2,2);
insert into realtime_insert values (1,3,3);
```

```
9.0.16.14 :) select * from realtime_insert;
SELECT *
FROM realtime_insert
```

Query id: e19c1105-9b13-4d4f-810d-596bc9cfb62b

id	value1	value2
1	3	3

如果两次写入的数据对应的 UNIQUE KEY 字段数值一致, 那么后写入数据会覆盖前一次的数据, 保证数据的最终去重。

UPSERT 实时更新

指定了 UNIQUE KEY 的表, 支持 update 语句用于更新表中的数据。

语法:

```
update [db.]table set column1 = expr1 [, ...] WHERE filter_expr
```

如下是示例:

```
insert into realtime_insert values (1,3,3);
update realtime_insert set value1 = 4 where id = 1;
9.0.16.14 :) select * from realtime_insert;
```

```
SELECT *
FROM realtime_insert
```

Query id: 3972121e-6474-44ee-bd48-5168d2e54c74

id	value1	value2
1	4	3

⚠ 注意:

- update 语法不能更新 UNIQUE KEY 或者 PRIMARY KEY 选定的字段。
- 更新大量数据会消耗较多的性能，执行时间也较长。

DELETE 实时删除

指定了 UNIQUE KEY 的表，支持 DELETE 语句用于删除中的数据。

语法:

```
delete [db.]table WHERE filter_expr
```

以下是示例:

```
insert into realtime_insert values (1,3,3);
insert into realtime_insert values (2,3,3);
delete from realtime_insert where id = 1;
```

9.0.16.14 :) select * from realtime_insert;

```
SELECT *
FROM realtime_insert
```

Query id: 107c7db9-072f-4295-a63d-2a8912cde1e1

id	value1	value2
----	--------	--------

2	3	3

 注意：

删除大量数据会消耗较多的性能，执行时间也较长。

使用 MySQL 关联 HBase 维表数据到 ClickHouse

最近更新时间：2026-06-18 17:33:54

本文介绍了结合 MySQL 数据库、流计算 Oceanus、HBase 以及云数据仓库 TCHouse-C 来构建实时数仓，并通过流计算 Oceanus 读取 MySQL 数据、关联 HBase 中的维表，最终将数据存入云数据仓库 TCHouse-C 进行指标分析，实现完整实时数仓的全流程操作指导。

环境搭建

创建 Oceanus 集群

在 [流计算 Oceanus](#) 控制台的[计算资源 > 新建](#)，创建集群，选择地域、可用区、VPC、存储、日志、设置初始密码等。创建完后的集群如下：

ⓘ 说明

- 若未使用过 VPC、日志、存储这些组件，需要先进行创建。
- VPC 及子网需要和下面的 MySQL、EMR 集群使用同一个，否则需要手动打通（如对等连接）。

←

集群名称

运行中

资源续费

调整配置

关联空间

解绑空间

更多操作

集群信息

集群作业

集群服务

基本信息

集群名称

✎

集群 ID

集群状态

运行中

部署模式 ①

单可用区

地域 / 可用区

新加坡 / 新加坡二区

计费模式

包年包月 / 2026-05-23 10:35:07 到期

Flink 版本

Flink-1.13(默认), Flink-1.16 ①

创建时间

2026-04-23 10:35:07

集群描述

用于文档维护 ✎

标签

暂无标签 ✎

Flink UI 访问策略

所有公网IP均可访问 ✎

集群底座

公有云

资源信息

CU 规格

1核4G

对象存储 COS ①

计算资源(CU)

空闲12 / 总12(CU)

日志服务 CLS ①

尚未绑定日志主题 ⚙

默认日志采集方式

对象存储 COS ⚙

DNS

默认 ✎

VPC

✎ 共 4093 个子网IP, 剩余可用 4090 个

关联空间

已关联工作空间

暂无关联空间

创建 VPC 私有网络

私有网络是一块您在腾讯云上自定义的逻辑隔离网络空间，在构建 MySQL、EMR、TCHouse-C 集群等服务时选择的网络必须保持一致，网络才能互通，否则需要使用对等连接、VPN 等方式打通网络。

登录 [私有网络](#) 控制台，选择私有网络 > +新建，新建私有网络。

←

新建私有网络 VPC

创建网络帮助文档 ②

私有网络信息

所属地域

② 新加坡

VPC 名称

不超过60个字符

IPv4 CIDR

10.0.0.0/16

网络创建后不可更改，请您提前做好网络规划 ②

标签 ①

标签键

标签值

✎

+ 添加标签

① 键值对模板

② 历史记录

子网信息

子网

子网名称	CIDR ①	可用区 ①	可用 IP 地址数量	操作
请输入子网名称 0/60	10.0.0.0/24	请选择	253	✎

+ 添加 VPC

标签 ①

标签键

标签值

✎

+ 添加标签

① 键值对模板

② 历史记录

选择标签后，为以上子网全部添加标签配置

确定

取消

创建云数据库 MySQL 服务

云数据库 MySQL（TencentDB for MySQL）是腾讯云基于开源数据库 MySQL 专业打造的高性能分布式数据存储服务，让用户能够在云中更轻松地设置、操作和扩展关系数据库。

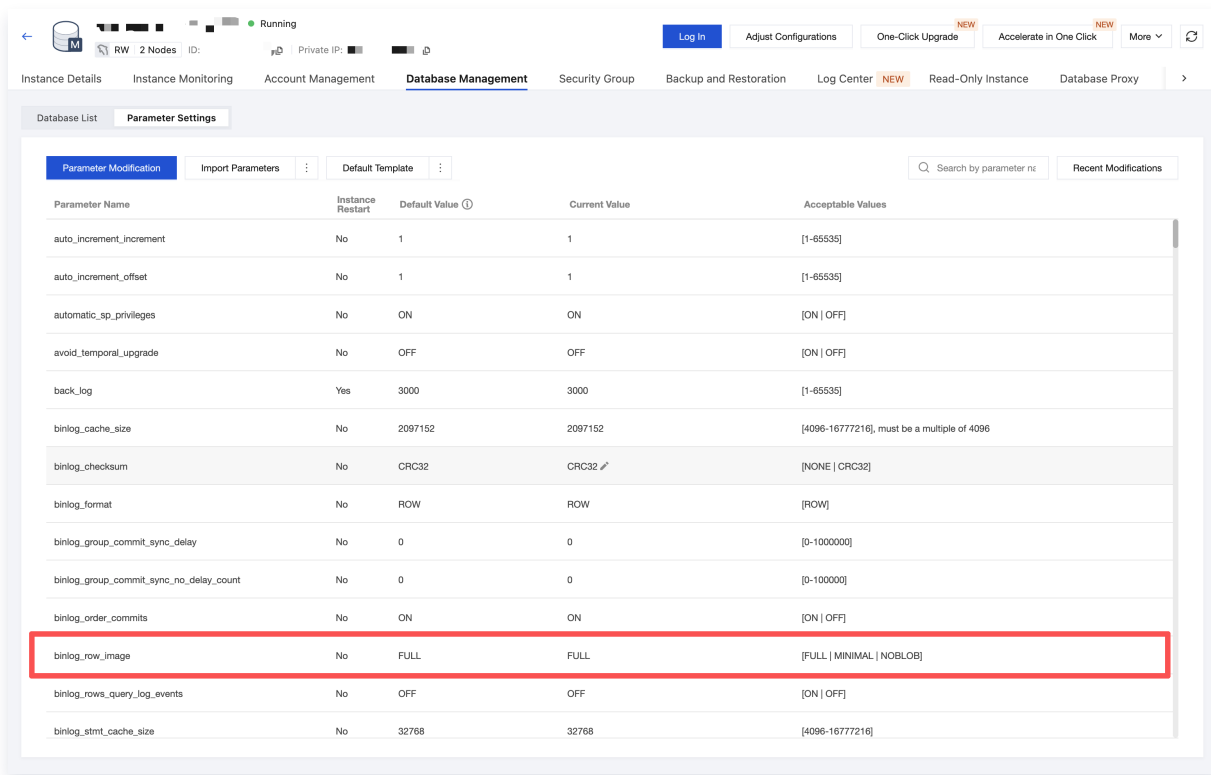
1. 登录 [云数据库 TencentDB](#) 控制台，选择**实例列表** > **新建**，新建实例。



2. 新建 MySQL 服务时，网络需要选择之前创建的。



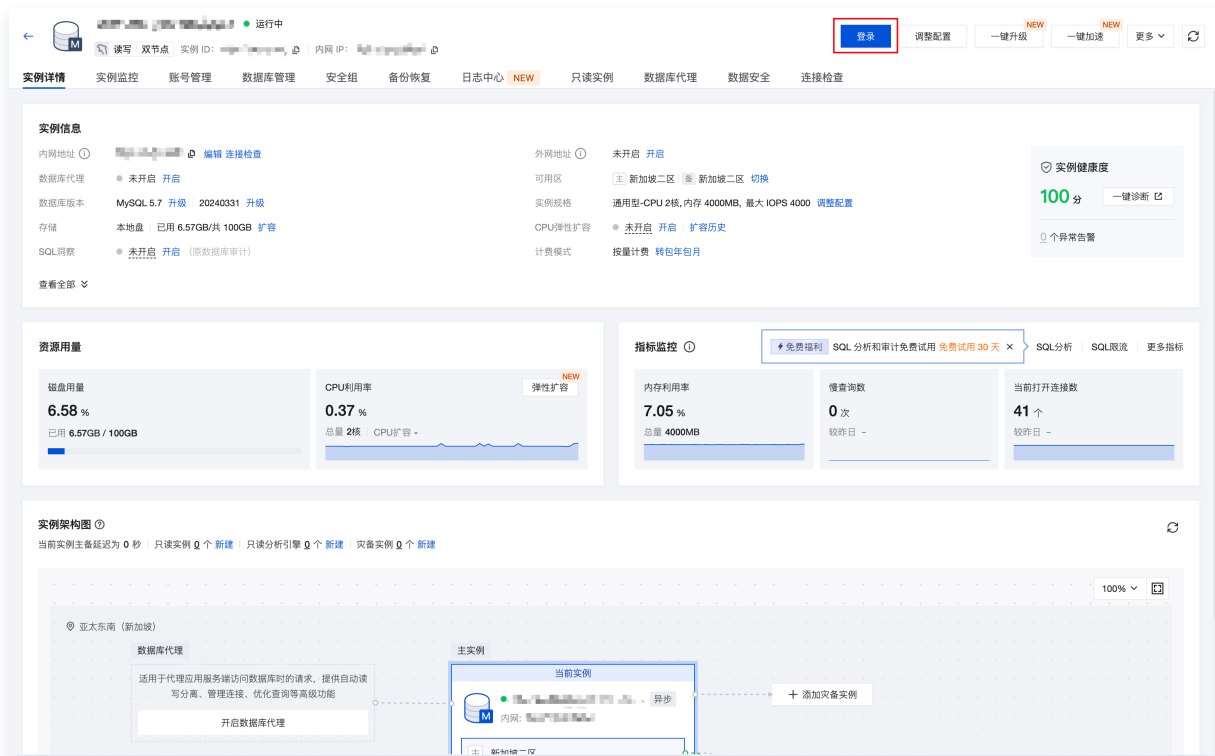
3. 创建完 MySQL 服务后，需要修改 binlog 参数，如图修改为 FULL（默认值为 MINIMAL）。



4. 修改完参数后，登录 MySQL 创建示例所需要的数据库和数据库表。

创建数据库 mysqltestdb

1. 登录 MySQL 创建示例所需要的数据库。



2. 打开 SQL 窗口或者单击可视化页面创建数据库及表。

新建数据库

```
create database mysqltestdb;
```

在新建的数据库上新建表 student:

```
create table `student` (  
  `id` int(11) not null auto_increment comment '主键id',  
  `name` varchar(10) collate utf8mb4_bin default '' comment '名字',  
  `age` int(11) default null comment '年龄',  
  `create_time` timestamp null default current_timestamp comment '数据创建  
时间',  
  primary key (`id`)  
) engine=innodb auto_increment=4 default charset=utf8mb4  
collate=utf8mb4_bin row_format=compact comment='学生表'
```

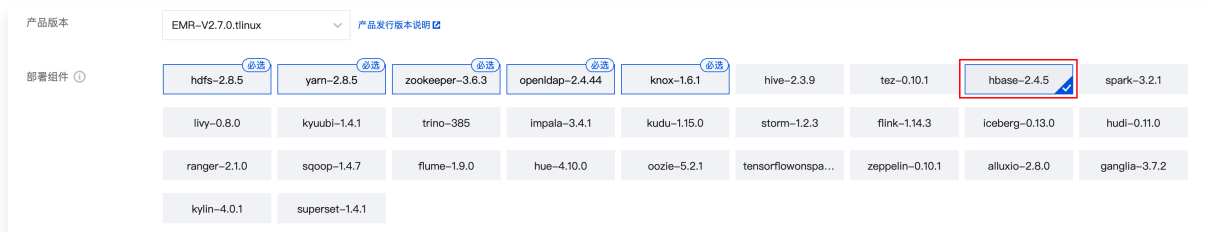
Student 表中插入数据

```
insert into mysqltestdb.student(id,name,age) values(1,"xiaomin",20);
```

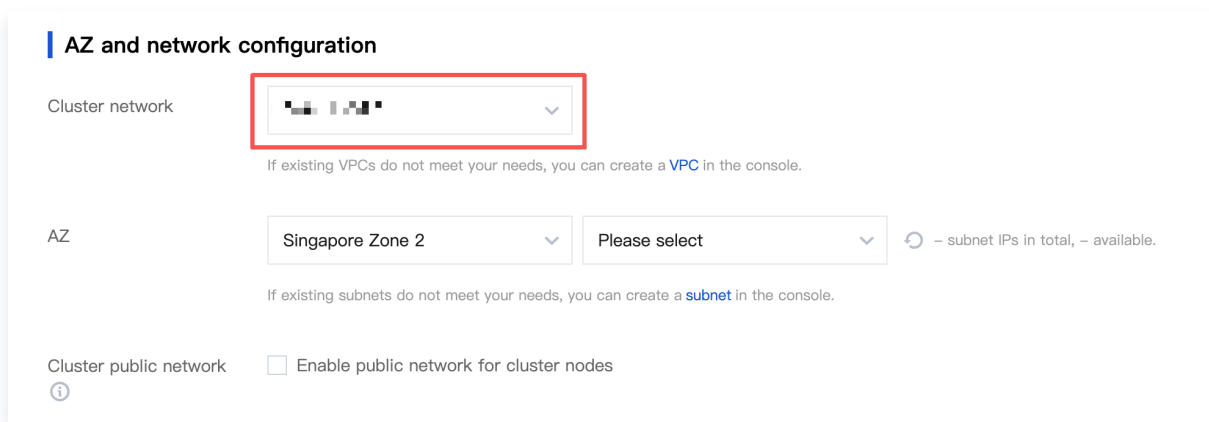
创建 EMR 集群

弹性 MapReduce 是云端托管的弹性开源泛 Hadoop 服务，支持 Spark、HBase、Presto、Flink、Druid 等大数据框架，本次示例主要需要使用 HBase 组件。

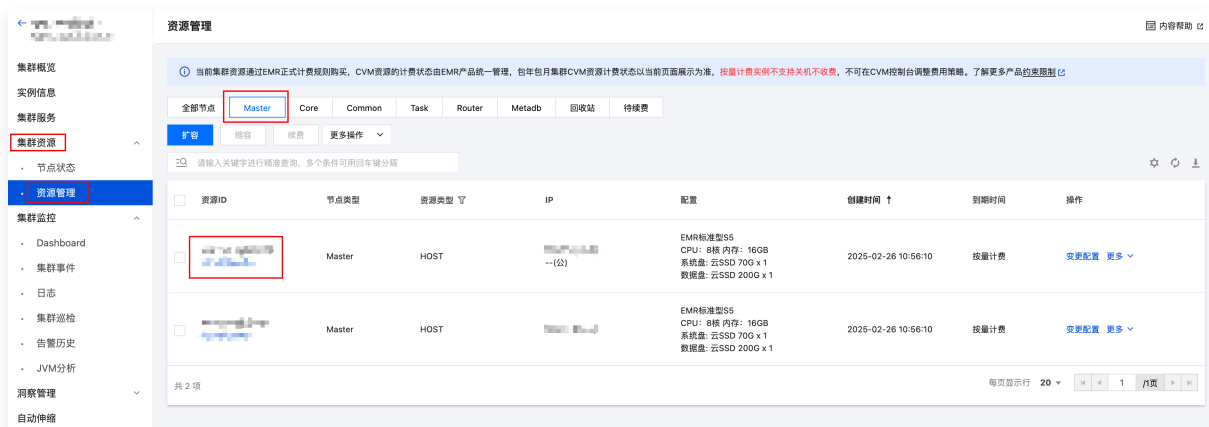
1. 登录 [弹性 MapReduce 控制台](#)，选择**集群列表 > 新建集群**，开始新建集群，具体可参考 [创建 EMR 集群](#)。新建集群时，需选择安装 HBase 组件。



如果是生产环境，服务器配置可根据实际情况选择。网络需要选择之前创建好的 VPC 网络，始终保持服务组件在同一 VPC 下。



2. 在集群列表中，单击新建的集群 ID/名称，进入集群详情页。选择**集群资源 > 资源管理**，即进入 HBase 的 Master 节点。



3. 进入 [云服务器控制台](#)，搜索 EMR 实例 ID，然后单击登录进入服务器。



4. 创建 Hbase 表。

```
# 进入HBase命令
root@172~# hbase shell
```

进入 hbase shell，并新建表：

```
# 建表语句
create 'dim_hbase', 'cf'

# 插入数据
put 'dim_hbase','1','cf:name','MingDeSchool'
```

创建云数据仓库 TCHouse-C

新建集群

1. 登录 [云数据仓库 TCHouse-C](#) 控制台，选择集群列表 > 新建集群，新建集群。



2. 选择网络选择之前新建的 VPC 网络（依然保证各服务在同一网络）。

地域

中国

重庆 广州 香港 中国台北

亚太地区

曼谷 雅加达 首尔 新加坡 东京

欧美

法兰克福 弗吉尼亚 硅谷 圣保罗

处于不同地域的云产品内网不通，购买后不能更换，请您谨慎选择；建议选择最靠近您客户的地域，可降低访问时延。

部署方案

单可用区 多可用区

处于同一私有网络下不同可用区的云产品内网互通；计算节点部署在主可用区，备可用区1；ZK节点部署在主可用区、备可用区1、备可用区2

副本设置

单副本 双副本

单副本不支持集群高可用，单节点（分片）只提供1个副本，不建议作为生产环境
双副本支持集群高可用，单节点（分片）提供2个副本，默认3个Zookeeper节点

网络

如现有的网络不合适，您可以去控制台新建私有网络 或新建子网。

可用区

新加坡二区（主力）

共 253 个子网IP，剩 241 个可用。

集群配置

集群名称

请输入集群名称

长度限制为6-36个字符，只允许包含中文、字母、数字、-、_、.

登录 TCHouse-C

在之前新建的 EMR 下选择一台云服务器单击**登录**，最好选择带有外网 IP 的节点。



安装 ClickHouse 客户端

在此机器上安装 ClickHouse 客户端，clickhouse-client 安装可参见 [快速入门](#)。

登录客户端

登录客户端示例如下：

```
clickhouse-client -h用户自己的ClickHouse服务IP --port 9000
```

新建数据库：

```
create database testdb on cluster default_cluster;
```

新建表：

```
CREATE TABLE testdb.student_school on cluster default_cluster (
```

```
`id` Int32,
`name` Nullable(String),
`school_name` Nullable(String),
`Sign` Int8
) ENGINE = ReplicatedCollapsingMergeTree('/clickhouse/tables/{layer}-{shard}/testdb/ student_school, '{replica}', Sign) ORDER BY id;
```

数据清洗和运算加工

数据准备

按照上面的操作创建表，并向 MySQL 和 HBase 表中插入数据。

创建 Flink SQL 作业

在流计算 Oceanus 控制台创建 SQL 作业，选择响应的内置 Connector。

Source 端

MySQL-CDC Source:

```
--学生信息作为cdc源表
CREATE TABLE `student` (
  `id` INT NOT NULL,
  `name` varchar,
  `age` INT,
  proc_time AS PROCTIME(),
  PRIMARY KEY (`ID`) NOT ENFORCED
) WITH (
  'connector' = 'mysql-cdc',
  -- 必须为 'mysql-cdc'
  'hostname' = 'YoursIp',
  -- 数据库的 IP
  'port' = '3306',
  -- 数据库的访问端口
  'username' = '用户名',
  -- 数据库访问的用户名（需要提供 SHOW DATABASES, REPLICATION SLAVE,
  REPLICATION CLIENT, SELECT, RELOAD 权限）
  'password' = 'YoursPassword',
  -- 数据库访问的密码
  'database-name' = 'mysqltestdb',
  -- 需要同步的数据库
```

```
'table-name' = 'student' -- 需要同步的数据表名
);
```

HBase 维表:

```
--示例使用school学校信息作为维表
CREATE TABLE dim_hbase (
  rowkey STRING,
  cf ROW <school_name STRING>, -- 如果有多个列簇, 写法 cf Row<age INT,name
String>
  PRIMARY KEY (rowkey) NOT ENFORCED
) WITH (
  'connector' = 'hbase-1.4',
  'table-name' = 'dim_hbase',
  'zookeeper.quorum' = '用户自己的hbase服务器zookeeper地址, 多个用逗号隔开'
);
```

Sink 端

创建到 TCHouse-C 的创建表语句。

```
--关联后存入clickhouse表
CREATE TABLE `student_school` (
  stu_id INT,
  stu_name STRING,
  school_name STRING,
  PRIMARY KEY (`id`) NOT ENFORCED
) WITH (
  -- 指定数据库连接参数
  'connector' = 'clickhouse',
  'url' = 'clickhouse://yourIP:8123',
  -- 如果TCHouse-C集群未配置账号密码可以不指定
  --'username' = 'root',
  --'password' = 'root',
  'database-name' = 'testdb',
  'table-name' = ' student_school ',
  'table.collapsing.field' = 'Sign'
);
```

进行逻辑运算

此示例中，只进行了简单的 Join 没有进行复杂的运算。详细运算逻辑可参见 [Oceanus 运算符和内置函数](#)。

```
INSERT INTO
    student_school
SELECT
    student.id as stu_id,
    student.name as stu_name,
    dim_hbase.cf.school_name
FROM
    student
JOIN dim_hbase for SYSTEM_TIME as of student.proc_time
    ON CAST(student.id AS STRING) = dim_hbase.rowkey;
```

结果验证

在数据库中查询数据是否正确。

```
select * from testdb.student_school;
```

腾讯云数据仓库 TCHouse-C 内核使用建议与规范

最近更新时间：2026-05-27 10:17:41

使用 腾讯云数据仓库 TCHouse-C 时需注意以下规范：

写入规范

- 攒批写入：**腾讯云数据仓库 TCHouse-C 必须攒批写入，至少 1000 条/批，建议 5k – 10w 一批写入腾讯云数据仓库 TCHouse-C，每一次写入都会在底层生成 1 个或者多个 part 存储目录，后台任务自动合并小 part 到一个大 part，如果写入频次过高会出现 part 过多，merge 速度跟不上导致写入失败报错：`Too many parts(301). Merges are processing significantly slower than inserts`。
- 减少分布式表直接写入：**为了提高写入和查询性能，应尽可能直接写入本地表，而不是分布式表。写分布式表最终也会转发给本地表，但是分布式表存在写放大以及异步落盘消耗 IO 的问题，写入性能较差。
- 约束数据一致性：**腾讯云数据仓库 TCHouse-C 不支持数据写入的事务保证，因此需要通过外部导入数据模块来控制数据的幂等性。例如，如果某个批次的数据导入异常，可以删除对应的分区数据或清理导入的数据，然后重新导入该分区或批次的数据。也可以使用去重引擎（replacingMergeTree）来保证最终一致性。
- 大规模数据写入：**如果需要进行大规模数据写入，建议提前拆分数据，并按节点均匀地写入腾讯云数据仓库 TCHouse-C 的各个节点。如果存在特定的分布规则，可以在业务侧进行 hash 计算。
- 一次只写入一个分区数据：**为了避免写入性能下降和目录数量过多的问题，应该一次只写入一个分区的数据。如果一批写入数据跨多个分区，会导致底层产生多个 part 文件，消耗更多的 merge 性能，并且不利于幂等控制。

查询规范

单表查询

- 高频过滤和点查询字段使用索引加速。
- 避免使用 `select *` 语句，应该明确需要查询的字段，只查询必要的字段。腾讯云数据仓库 TCHouse-C 底层是列式存储，查询的耗时与查询的字段大小和数量成线性关系。
- 当查询千万以上的数据集时，建议使用 `where` 条件和 `limit` 语句来配合 `order by` 查询，以提高查询效率。
- `select {tablename} final` 能够实现查询（read on merge），但是会减慢查询速度，需要有针对性使用。
- 尽量按分区过滤裁剪，通过指定分区字段可以减少底层数据库扫描的文件数量，提高查询性能。
- 谨慎使用 delete 和 update 的 mutation 操作。腾讯云数据仓库 TCHouse-C 的 update 和 delete 是异步进行的，并且会重写 `where` 条件过滤出的数据 part，是非常重的操作，可能会消耗较多系统资源。此外，update 和 delete 是按照 part 逐个执行，不会保证整体执行的原子性。
- 如果对唯一性要求不高，可以采用近似去重 uniqCombined 来优化去重逻辑，从而提高十倍的查询性能。如果查询允许有误差，可以使用 uniqCombined 替代，否则应该继续使用 distinct 语法。使用 distinct 会对查询性能有

一定影响。

多表关联

1. 为了避免 Join 操作和 shuffle，应尽量使用 Flat 大宽表结构代替多表 Join。
2. 控制 Join 的表数量，尽量保持在3个及以下。
3. 如果查询字段出自单表，可考虑将 Join 改为 in 查询，CK 中的 in 查询支持单字段和 tuple，例如 `SELECT name FROM tab_a WHERE id IN ( SELECT id FROM tab_b WHERE name = 'x' )`。
4. 多表 Join 最好改为两表 Join 和子查询形式。
5. 两表 Join 需大表 Join 小表（数据量控制在百万 – 千万行级别），小表 Join 小表，不允许大表 Join 大表；Join 时大表需在左边，小表在右边。
6. 建议使用列裁剪、分区裁剪，Join 前进行条件过滤，尽可能降低 Join 数据量。
7. 若 Join 时右表为子查询或者分布式表且数据不大，可以采用 GLOBAL JOIN 避免读放大，需要注意的是，GLOBAL JOIN 会触发数据在节点之间传播，占用部分网络流量。如果数据量较大，同样会带来性能损失。

建表规范

1. 高可用集群不可创建非 Replicated 表，非高可用集群不可创建 Replicated 表。
2. 如果对数据最终一致性有强要求，需要使用 ReplacingMergeTree 或者 CollapsingMergeTree 引擎，并定期进行 optimize 或使用 `select {tablename} final` 实现最终去重。
3. 在规划分区时，应该合理规划分区个数，并尽可能利用分区。一张表分区数不建议超过 1000 个，可以在查询时有效帮助进行数据过滤，使用得当可以提升数倍查询性能，通常按天分区是比较普遍的做法。分区也不建议过多，因为腾讯云数据仓库 TCHouse-C 不同分区的数据不会合并，容易导致 part 过多，从而导致查询和重启变得很慢。
4. 建表时尽可能提前规划好表字段，并尽量避免删改字段。删改字段会重写整个表的全量数据，对于大表会消耗大量资源，执行时间可能很长。此外，删改字段期间也容易阻塞其他 DDL 语句，影响表的 merge 操作。如果中途出错，有概率会导致不可预知的数据一致性问题。
5. 禁止修改索引列，对索引列的修改会导致现有索引失效，触发重建索引，期间查询数据不准确。
6. 约束 COS 上存储数据的量，尽可能避免对冷分区进行写入和 mutation 操作。COS 单个桶大约只有1GB的带宽，远低于多节点的本地盘和云盘性能，且网络延迟比较高。如果 COS 上存储过多数据，会严重影响查询效率。针对 COS 分区的写入时，会触发 COS 分区进行 merge，merge 效率也会降低甚至会影响本地盘的数据操作。