

Table Store

Best Practices

Best Practices

This section details recommendations for optimizing your use of Table Store.

A well-designed Primary Key

Table Store dynamically divides table data into partitions according to Partition Keys, and each partition is served by one server node. A Partition Key is the smallest partition unit. The data under a Partition Key cannot be split. In this case, applications should balance data distribution and access distribution across partitions in order to leverage Table Store's capability.

Table Store will sort the rows in a table by the Primary Key. A well-designed Primary Key can better balance data distribution across partitions, making full use of the Table Store's high scalability.

When selecting a Partition Key, the following should be noted:

Data of all rows in a single Partition Key should not exceed 1 GB. While 1GB is not the actual limitation, this is recommended in order to avoid a hot spot.

Data in different Partition Keys of the same table are logically independent.

The access pressure should not be concentrated on a small range of consecutive Partition Keys.

Example

Assume you have a table that stores students' transaction records conducted using student ID cards. In this scenario:

Each student card corresponds to one CardID.

Each seller corresponds to one SellerID.

Each purchase terminal corresponds to a DeviceID, which is globally unique.

For each purchase generated by a purchase terminal, one OrderNumber is recorded. An

OrderNumber generated by a purchase terminal is unique to the terminal, but is not globally unique.

For example, different purchase terminals may generate two separate purchase records using the same OrderNumber. Each OrderNumber generated by the same purchase terminal has a different time stamp. New purchase records will have larger sequential OrderNumbers than the previous purchase records. Every purchase record will be written into the table in real time.

To optimize the use of Table Store, CardID or DeviceID are recommended as the table's Primary Key:

Using CardID is strongly recommended as, generally, the number of purchase records for each card, each day, is similar, thereby the access pressure for each Partition Key will be balanced. This allows for a good utilization of the reserved read/write throughput.

Using DeviceID is recommended as, even though the number of purchase records for each seller per day varies, the number of purchase records generated by each purchase terminal per day can be estimated. This estimation is calculated based on a cashier's operating speed, which determines the number of purchase records that can be generated by their purchase terminal per day. Therefore, DeviceID is suitable as the table's Partition Key to ensure a balanced distribution of access pressure.

The SellerID and OrderNumber are not recommended. The SellerID is not recommended as it indicates the limited number of sellers available, and therefore, does not help to balance access pressure for each Partition Key in scenarios in which a small number of sellers generate the majority of purchase records. The OrderNumber is not recommended due to the sequential increase of purchase orders generated at the same time, resulting in grouped orders in the same time period. This restricts the effectiveness of the read/write throughput.

Note: If OrderNumber must be the Partition Key, you can hash it and use the hash value as the OrderNumber prefixes. This will ensure an even distribution of the data and access distribution pressure.

Spliced Partition Keys

For optimized Table Store use, it is recommended that the data volume of a single partition does not exceed 1 GB. If the total data volume for all rows in a single table partition exceeds 1 GB, you can splice multiple original Primary Key columns into a Partition Key when designing the table.

Example

As in the preceding student card purchase record example, assume the Primary Key columns are [DeviceID, SellerID, CardID, OrderNumber]. DeviceID is the Partition Key for this table and the total

data volume from all rows of a single DeviceID may exceed 1 GB. In this case, splice DeviceID, SellerID, and CardID as the table's first Primary Key column (Partition Key).

The original table is shown as follows:

DeviceID	SellerID	CardID	OrderNumber	attrs
16	'a100'	66661	200001	...
54	'a100'	6777	200003	...
54	'a1001'	6777	200004	...
167	'a101'	283408	200002	...

After splicing DeviceID, SellerID, and CardID to create the Partition Key, the new table is shown as follows:

CombineDeviceIDSellerIDCardID	OrderNumber	attrs
'16:a100:66661'	200001	...
'167:a101:283408'	200002	...
'54:a1001:6777'	200004	...
'54:a100:6777'	200003	...

In the original table, the two rows for DeviceID = 54 belong to two purchase records in the same Partition Key 54. In the new table, these two purchase records have different Partition Keys. By splicing multiple Primary Key columns to form a Partition Key, you can reduce the total data volume for each Partition Key in the table.

However, splicing the Primary Key columns to form a table poses some drawbacks. DeviceID is an integer-type Primary Key column. In the original table, the purchase records of DeviceID = 54 are listed before those of DeviceID = 167. After splicing the first three Primary Key columns into a string-type Primary Key column, the purchase records of DeviceID = 54 will be listed after those of DeviceID = 167. If the application needs to read all purchase records from the DeviceID range [15, 100), the preceding table is not optimal.

In response to this situation, you can add 0s in front of the DeviceIDs. The number of 0s to add is determined by the maximum number of digits for DeviceIDs. If the DeviceID range is [0, 999999], you can add 0s so that all DeviceIDs have 6 digits, and then splice. The resulting table will be as follows:

CombineDeviceIDSellerIDCardID	OrderNumber	attrs
'000016:a100:66661'	200001	...
'000054:a1001:6777'	200004	...
'000054:a100:6777'	200003	...
'000167:a101:283408'	200002	...

However, even after padding 0s in front of the IDs, the table is still not fully optimal. This is because of the two rows with DeviceID = 54; the row with SellerID = 'a1001' is listed after SellerID = 'a100'. This discrepancy is caused by the connector :, which influences the lexicographic order, meaning '000054:a1001' is lexicographically less than '000054:a100:', but 'a1001' is greater than 'a100'. To resolve this issue, choose a character that is less than the ASCII code of all other available characters. In this table, the SellerID value uses uppercase and lowercase letters and digits. It is recommended that , is the connector, because the ASCII code for , is less than the ASCII code of all characters available for the SellerID.

Using , and then splicing, produces the following optimized table:

CombineDeviceIDSellerIDCar dID	OrderNumber	attrs
'000016,a100,66661'	200001	...
'000054,a100,6777'	200003	...
'000054,a1001,6777'	200004	...
'000167,a101,283408'	200002	...

Summary

If the total data size for all rows in a single Partition Key exceeds 1 GB, you can splice multiple Primary Key columns to form a Partition Key to minimize the data size of an individual Primary Key. When splicing the Partition Key, the following should be noted:

When choosing the Primary Key columns to splice, be sure that the original rows of the same Partition Key will have different Partition Keys after splicing.

When splicing integer-type Primary Key columns, you can add 0s before the numbers to ensure the rows' order remains the same.

When selecting a connector, you must consider its effect on the lexicographical order of the new Partition Key. The ideal method is to select a connector with an ASCII code that is less than all other available characters.

Add hash prefixes in Partition Keys

Example

In the Primary Key section, it is recommended that OrderNumber is not used as the table's Partition Key. Since OrderNumbers increase sequentially, purchase records will always be written in the newest OrderNumber range, meaning older OrderNumber ranges will not experience any written pressure.

This will cause an imbalance in access pressure resulting in inefficient use of the reserved read/write throughput. If a sequentially increasing key value needs to be used as the Partition Key, splice a hash prefix to the Partition Key. In this way, the OrderNumbers will be randomly distributed throughout the table to better balance the access pressure distribution.

The purchase records table using OrderNumber as the Partition Key is as follows:

OrderNumber	DeviceID	SellerID	CardID	attrs
200001	16	'a100'	66661	...
200002	167	'a101'	283408	...
200003	54	'a100'	6777	...
200004	54	'a1001'	6777	...
200005	66	'b304'	178994	...

As an example, for the OrderNumbers, you can use the md5 algorithm to calculate a prefix (other hashing algorithms are permitted) and splice it to create the HashOrderNumber. As the hash strings calculated by the md5 algorithm may be too long, you can take only the first few digits to achieve a random distribution of records of sequential OrderNumbers. In this example, the first 4 digits are used to produce the following table:

HashOrderNumber	DeviceID	SellerID	CardID	attrs
'2e38200004	54	'a1001'	6777	...
'a5a9200003	54	'a100'	6777	...
'c335200005	66	'b304'	178994	...
'db6e200002	167	'a101'	283408	...
'ddba200001	16	'a100'	66661	...

When subsequently accessing the purchase records, use the same algorithm to calculate the hash prefix of the OrderNumber to get the HashOrderNumber that corresponds to a purchase record. One drawback of adding a hash prefix to the Partition Key is that the originally contiguous records will be dispersed. This means that the GetRange operation cannot be used to get a range of logically consecutive records.

Write data in parallel

When Table Store tables are split into multiple partitions, these partitions are distributed across multiple Table Store servers. If a batch of data is ordered by the Primary Key to be uploaded to Table

Store, and the data is written in the same order, this may concentrate the written pressure on a certain partition, while the other partitions remain idle. This operation does not fully utilize the reserved read/write throughput and may impact the data import speed.

To resolve this issue, use either of the following methods to increase the data import speed:

Disrupt the original data order and then import. Ensure that the written data is evenly distributed across each Partition Key.

Use multiple worker threads for parallel data import. Split a large data set into multiple smaller sets. The worker threads will then randomly select a smaller set to import.

Distinguish cold data and hot data

Mismanaged time sensitive data can become an issue. Using the previous example of student transaction records, some purchase records may have a higher access probability as applications frequently query the latest record, as well as process and compile statistics based on the latest records. However, old purchase records continue to occupy storage space and become cold. Furthermore, if there is a large volume of cold data in a table (such as CardIDs of students no longer enrolled, yet retained in the system), the reserved read/write throughput is ineffectively utilized, and results in unbalanced access pressure across the partitions.

To effectively manage time sensitive data, use different tables to separate cold and hot data, and set a different reserved read/write throughput for them. For example, purchase records may be divided into different tables according to month, with a new table being created for each month. The reserved read/write throughput can then be set for each table as follows:

- A high reserved read/write throughput can be set for the table with the latest purchase records of the current month to satisfy its access needs (that is, new purchase records have a higher chance of being queried than legacy data).
- A low reserved write throughput and a high reserved read throughput can be set for later tables (of the past few months) in which little or no new data is written, but queries are still performed.
- A low reserved read/write throughput can be set for tables that have exceeded their maintenance period (such as historical records of a year or longer). These tables can then be exported to restore in an OSS archive, or simply deleted.

This section provides recommendations for optimizing Table Store data operations. Notably, it details how to effectively manage Attribute columns and application request errors.

Split tables among Attribute columns

If a table's rows have many Attribute columns, but each operation only accesses a portion of these

columns, you can split the table into multiple tables with the Attribute columns of different access frequencies placed in different tables. For example, in a merchandise management system with rows containing the item quantity, item price, and item description:

- Item quantities and prices are integer-type values that occupy little storage space, but are modified frequently.
- Item descriptions are string-type values that occupy more storage space, but are not modified frequently.

Because the majority of operations will only require updating the integer-type values of item quantities and prices, the table can be split into one table containing these two values, and the other table containing the string-type item descriptions.

Compress text-based Attribute columns

If an Attribute column contains a large amount of text, the Attribute columns can be compressed and stored as a binary-type in Table Store. This action saves space and reduces the capacity units consumed by access operations, thereby reduces the Table Store cost.

Store Attribute columns in OSS

Table Store limits the size of a single Attribute column to up to 2 MB. If you need to store a file that exceeds 2 MB, it is recommended to use Alibaba Cloud Object Storage Service (OSS). OSS is an open storage service capable of storing large files at lower storage prices, compared with Table Store.

If OSS cannot be used, individual values greater than 2 MB can be split into multiple, smaller rows, and then stored in Table Store.

Add error retry intervals

If an application's request fails and returns a "try again" error, it is recommended to wait a period of time before trying the request again. As a best practice, randomized or exponentially-increasing backoffs are helpful in order to avoid an avalanche effect. For more error details, refer to [Use Table Store's API](#).