



MDBW 2022

Under the Hood

Supercharging Atlas Data Lake

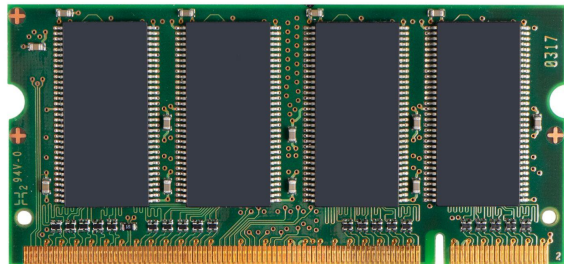


David Golden
Staff Engineer
@xdg





















Welcome!



David Golden, Staff Engineer

david@mongodb.com

twitter: @xdg



Welcome!



David Golden, Staff Engineer

david@mongodb.com

twitter: @xdg



My team...

- Data Federation
- Data Lake Storage



Today...

3

Traps

Principles

Changes

Trap #1: Tiny files





What's in the box?







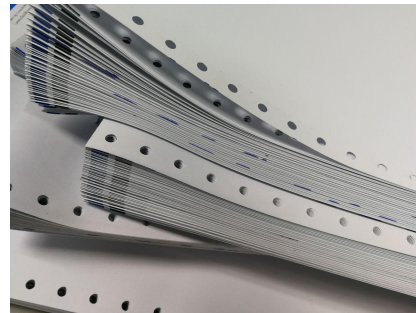


One more problem...

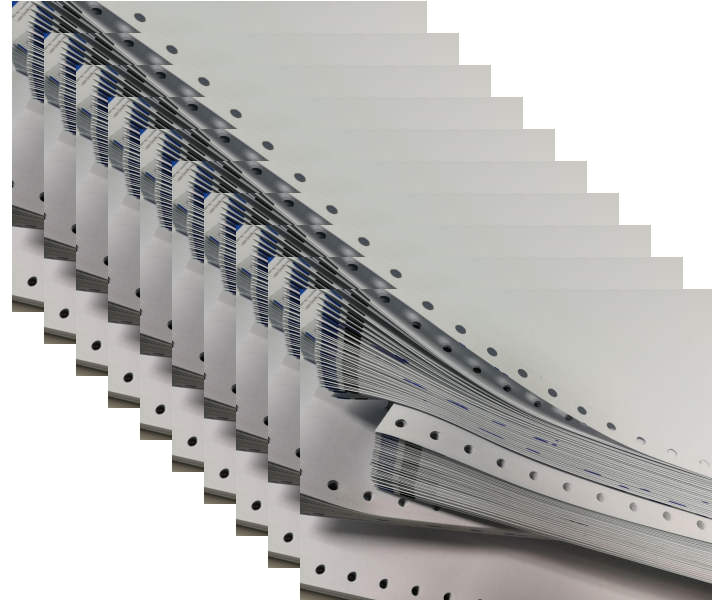








→ 1,000





Why?



Folder $\rightarrow$ bad metaphor



`s3://<bucket>/<key>`



Key $\rightarrow$ anything



7a4f37e9739c4e7a8238a7c67fd9f46e1a48870921144e348c3fdd7fba771



7a4f37e9739c4e7a8238a7c67fd9f46e1a48870921144e348c3fdd7f.json



7a4f37e9739c4e7a8238a7c67fd9f/6e1a48870921144e348c3fdd7f.json



7a4f37e9739c4e7a8238a7c67fd9f/6e1a48870921144e348c3fdd7f.json

Folder?

File?



sales/e9739c4e7a8238a7c67fd9f/6e1a48870921144e348c3fdd7f.json



sales/<uploadDate>/8a7c67fd9f/6e1a48870921144e348c3fdd7f.json



sales/<uploadDate>/<regionID>/<storeID>/21144e348c3fdd7f.json



sales/<uploadDate>/<regionID>/<storeID>/<productID>/dd7f.json



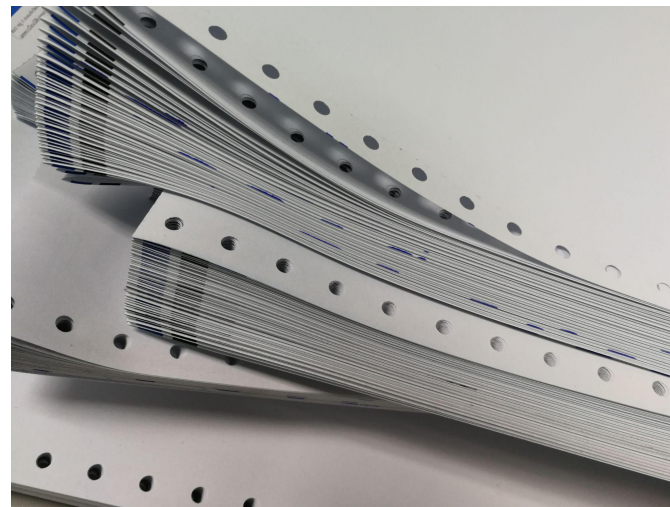
sales/<uploadDate>/<regionID>/<storeID>/<productID>/0001.json



Folders $\rightarrow$ tiny files



sales/<uploadDate>/<regionID>/<storeID>/<productID>/?????.json





Folders $\rightarrow$ pay twice!



Worst case



sales/<uploadDate>/<regionID>/<storeID>/<transactionID>.json

Unique! → single document file





Principle #1

Minimize number of network trips



Big files

(but not too big)

Trap #2:
Wasteful scans



Imagine a **tweet** data lake



Analyze tweets by @xdg



Assume folders...



Assume daily uploads...



File per screen name?

tweets/2022-05-12/xdg.json

How much did I tweet that day?



Roll up by date?

tweets/2022-05-12/0001.json

⋮

tweets/2022-05-12/0847.json

But which file has "xdg"?



Principle #2

Minimize fetching unneeded documents



Use file bounds?

tweets/2022-05-12/abba-cher.json

⋮

tweets/2022-05-12/weezer-zendaya.json

Only the last file has "xdg"



Analyze a whole month?



Daily files with bounds

tweets/2022-05-01/wyclef-ziggymarley.json

⋮

tweets/2022-05-12/weezer-zendaya.json

tweets/2022-05-13/wutangclan-zztop.json

"xdg" is a tiny part of every file



Roll-up



Merging gives a tighter range

tweets/2022-05/xaarian-xylo.json

More of the file is "xdg"

Trap #3: Wide documents





Consider a single tweet...



```
{  
  "id": 54691802283900930,  
  "text": "RT @PostGradProblem: In preparation for the NFL lockout, I will  
be spending twice as much time analyzing my fantasy baseball team during ...",  
  "created_at": "Sun Apr 03 23:48:36 +0000 2011",  
  "truncated": true,  
  "in_reply_to_user_id": null,  
  "in_reply_to_status_id": null,  
  "favorited": false,  
  
  ... [snip] ...  
  
  "user": {  
    "screen_name": "OldGREG85",  
    "notifications": null,  
    "profile_use_background_image": true,  
    "statuses_count": 351,  
  
    ... [snip] ...  
  }  
}
```



Top tweeters?



Count by screen name?



```
{  
  "id": 54691802283900930,  
  "text": "RT @PostGradProblem: In preparation for the NFL lockout, I will  
be spending twice as much time analyzing my fantasy baseball team during ...",  
  "created_at": "Sun Apr 03 23:48:36 +0000 2011",  
  "truncated": true,  
  "in_reply_to_user_id": null,  
  "in_reply_to_status_id": null,  
  "favorited": false,  
  
  ... [snip] ...  
  
  "user": {  
    "screen_name": "OldGREG85",  
    "notifications": null,  
    "profile_use_background_image": true,  
    "statuses_count": 351,  
  
    ... [snip] ...  
  }  
}
```



```
{  
  "$group": {  
    "_id": "$user.screen_name",  
    "cnt": {"$sum": 1}  
  }  
}
```



Only need one field!



```
{
  "id": 54691802283900930,
  "text": "RT @PostGradProblem: In preparation for the NFL lockout, I will
be spending twice as much time analyzing my fantasy baseball team during ...",
  "created_at": "Sun Apr 03 23:48:36 +0000 2011",
  "truncated": true,
  "in_reply_to_user_id": null,
  "in_reply_to_status_id": null,
  "favorited": false,

  ... [snip] ...

  "user": {
    "screen_name": "OldGREG85",
    "notifications": null,
    "profile_use_background_image": true,
    "statuses_count": 351,

    ... [snip] ...

  }
}
```

[illegible]



Download → discard



Principle #3

Minimize fetching unneeded fields



Column storage!



Row → Document
Column → Field



Row storage

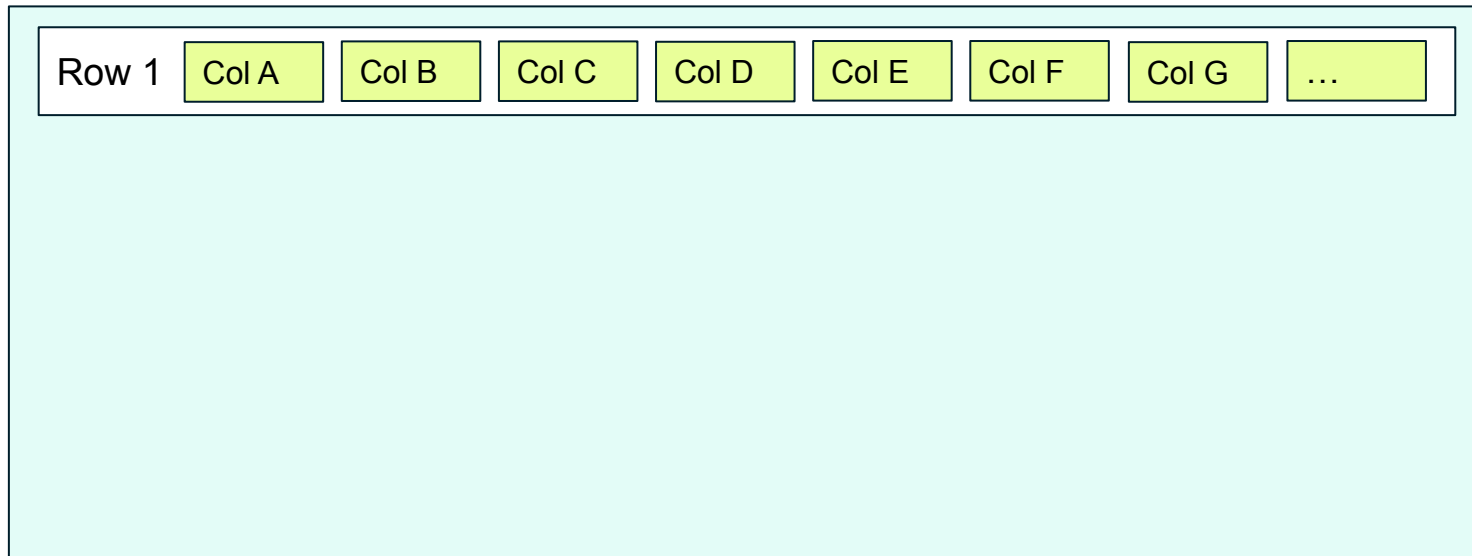
Row file

Row 1



Row storage

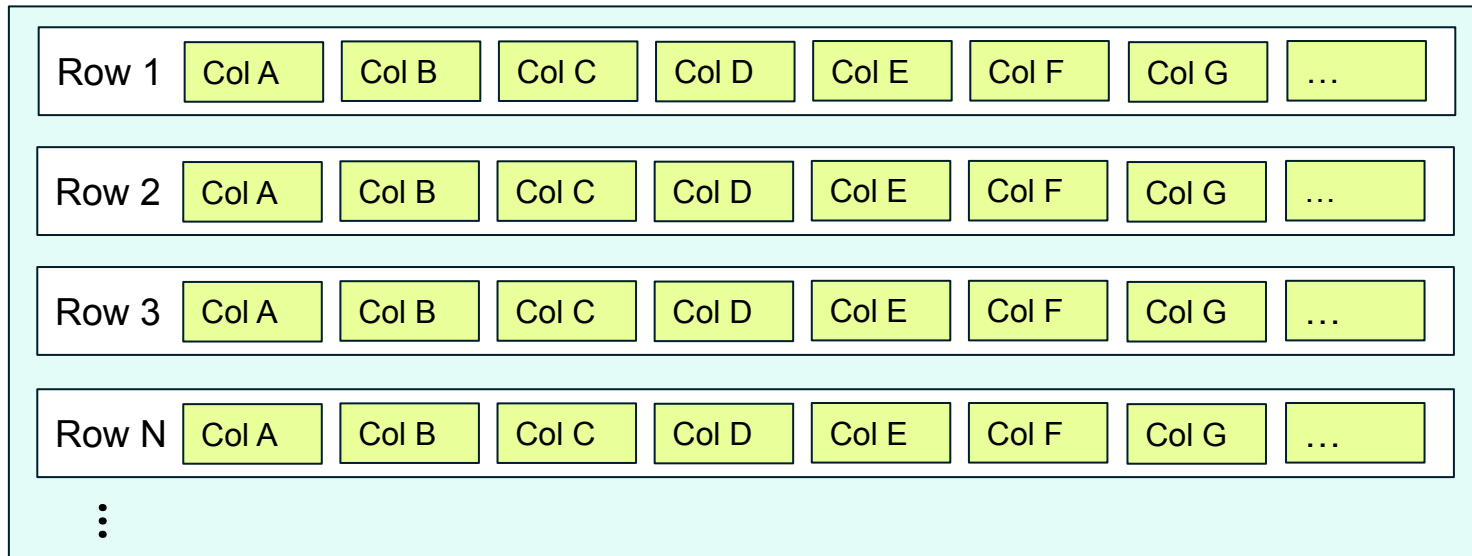
Row file





Row storage

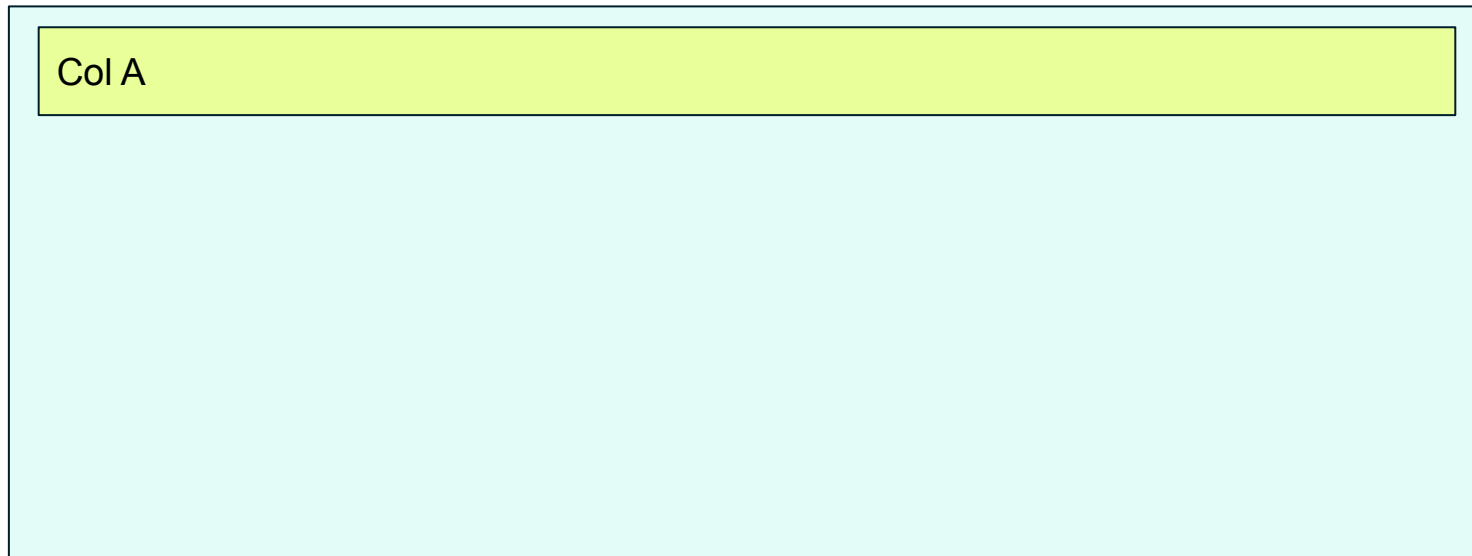
Row file





Column storage

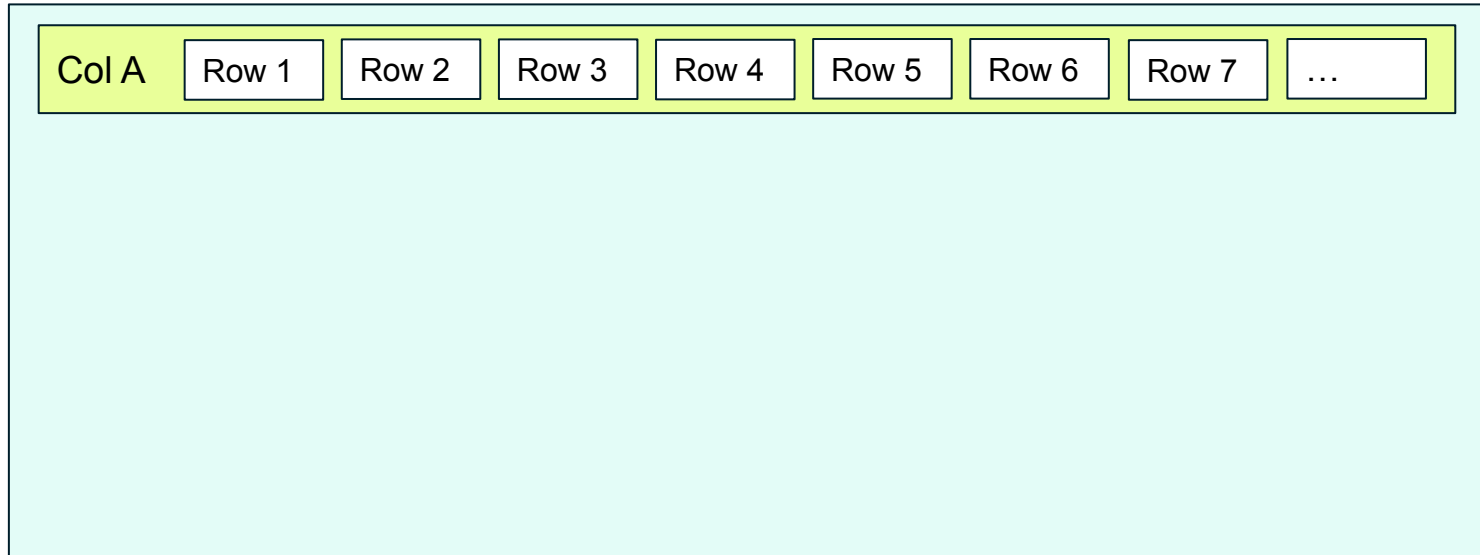
Column file





Column storage

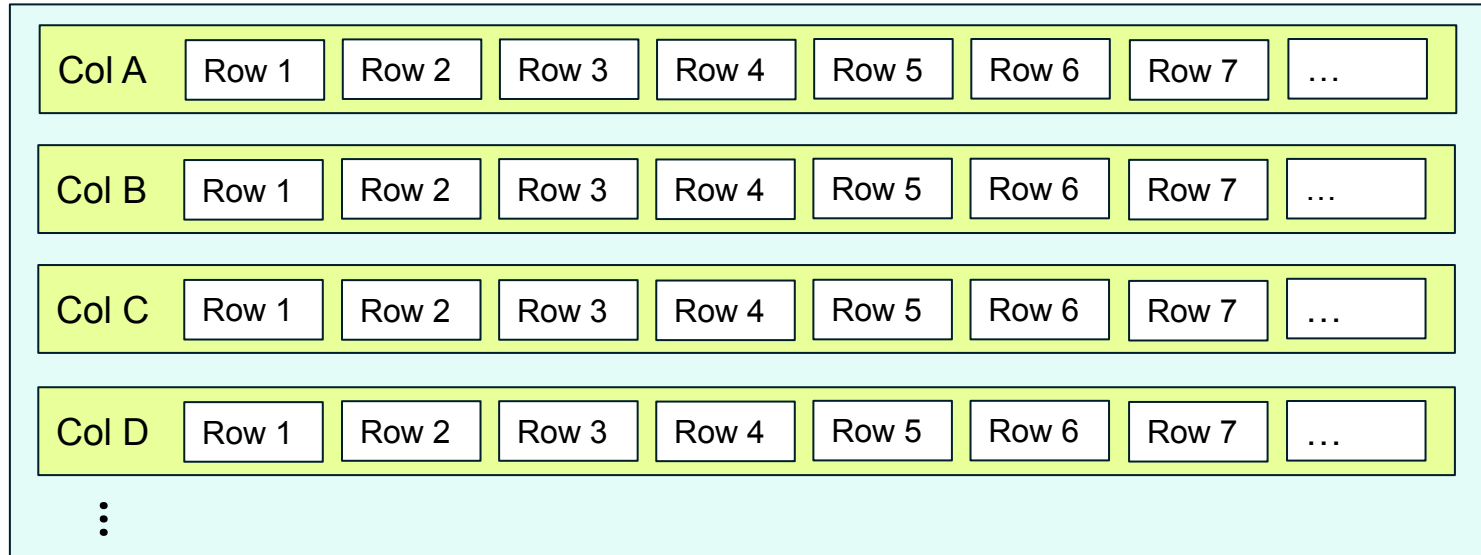
Column file





Column storage

Column file





Column metadata

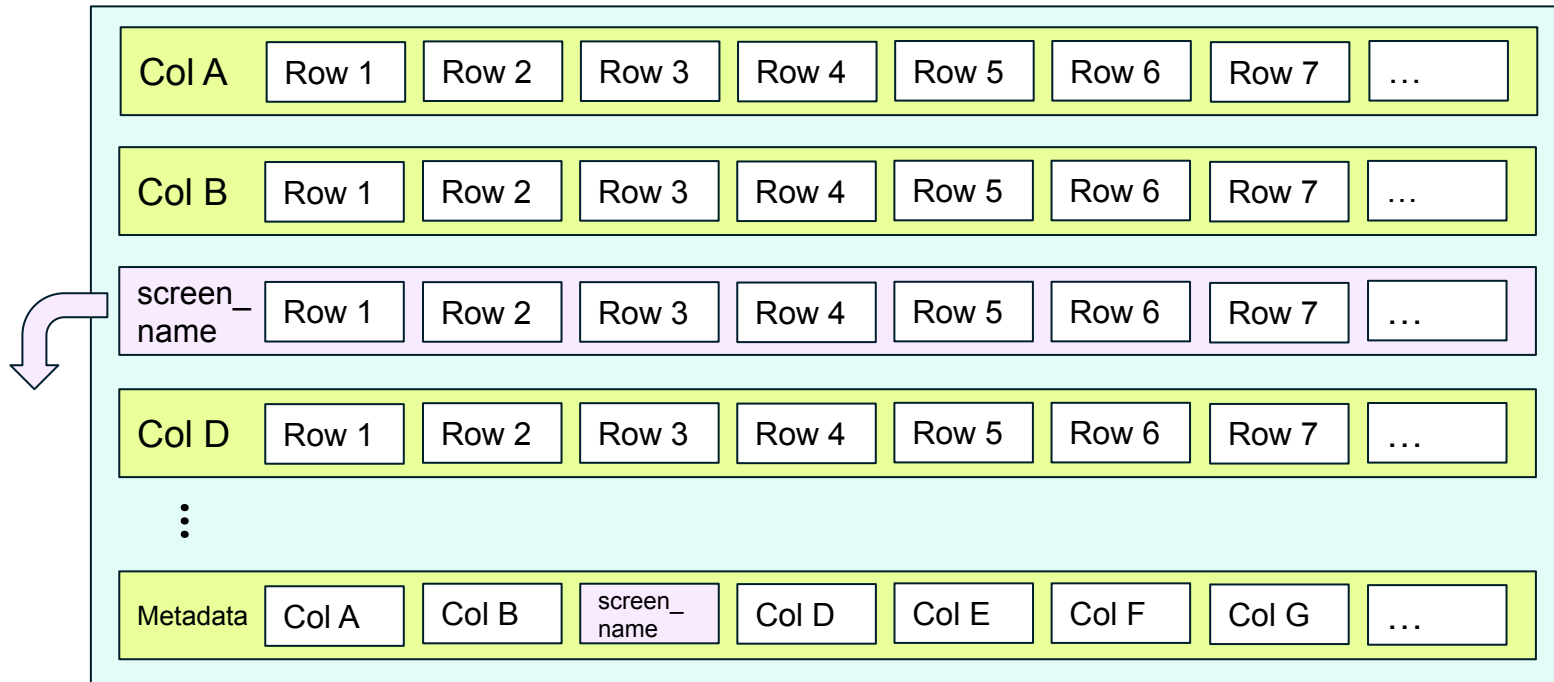
Column file

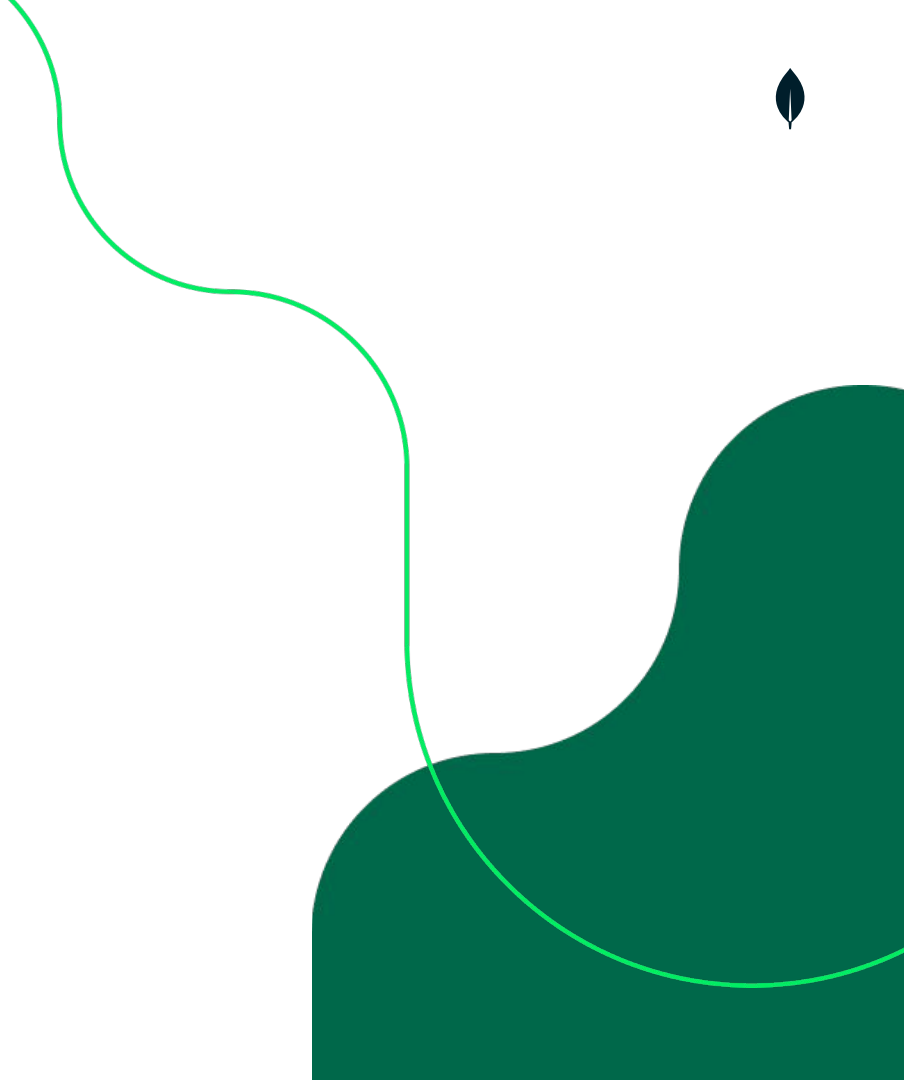
Col A	Row 1	Row 2	Row 3	Row 4	Row 5	Row 6	Row 7	...
Col B	Row 1	Row 2	Row 3	Row 4	Row 5	Row 6	Row 7	...
Col C	Row 1	Row 2	Row 3	Row 4	Row 5	Row 6	Row 7	...
Col D	Row 1	Row 2	Row 3	Row 4	Row 5	Row 6	Row 7	...
⋮								
Metadata	Col A	Col B	Col C	Col D	Col E	Col F	Col G	...



Fetch by column!

Column file





Recap:

Three principles

Tiny files

→ Minimize network trips



Tiny files

→ Minimize network trips



Wasteful scans

→ Minimize unneeded documents

Tiny files → Minimize network trips



Wasteful scans → Minimize unneeded documents

Wide documents → Minimize unneeded fields

Minimize network trips



Minimize unneeded documents

Minimize unneeded fields



Optimized storage

Metadata catalog

Auto-rebalancing



Change #1:

Optimized storage



No folders



Shard key



(A secondary index)



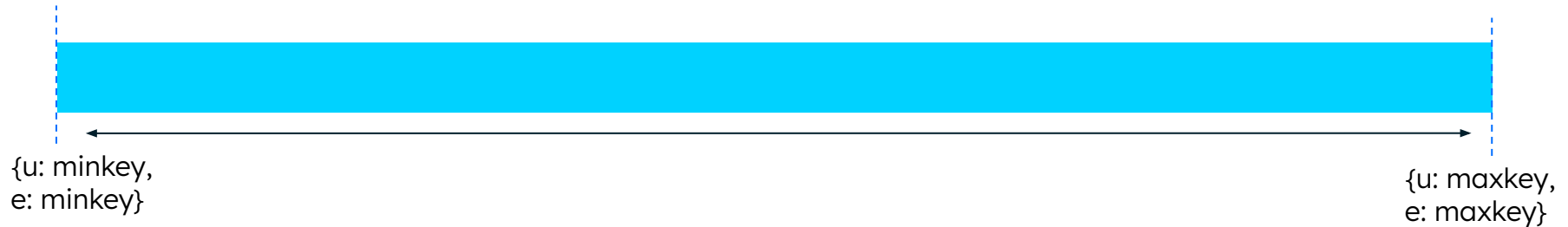
```
{user_id:1, event_time:1}
```



`{u:1, e:1}`

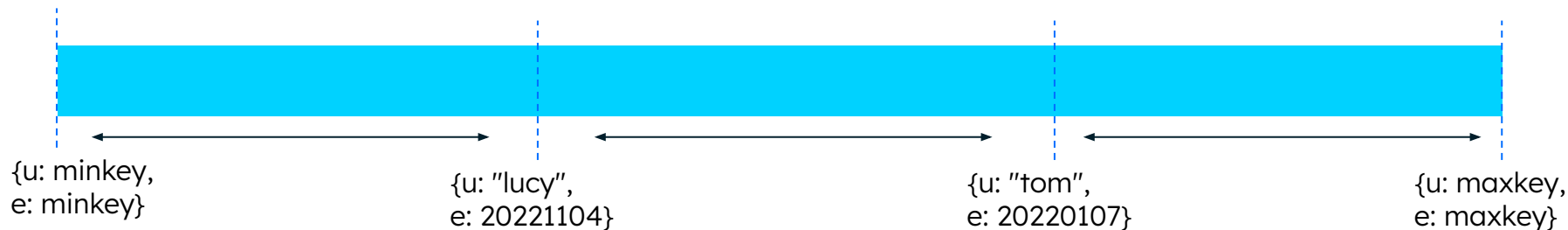


Key space $\rightarrow$ sorted range of all values



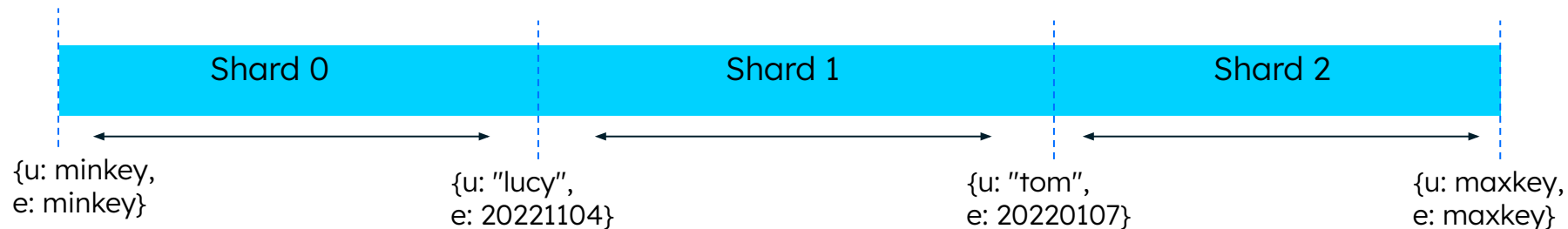


Key space → divide into ranges



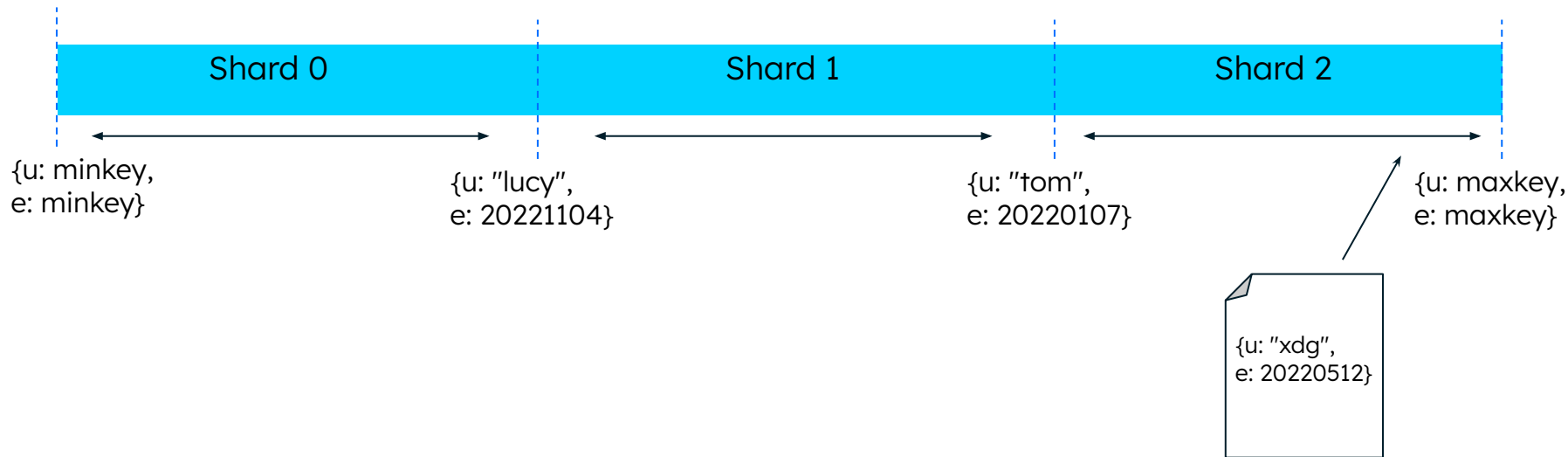


Key space $\rightarrow$ shards





Key space → every doc has a place





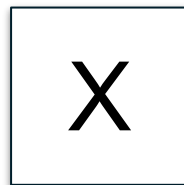
Files?



```
{u: "xdg", e: 20220512}
```

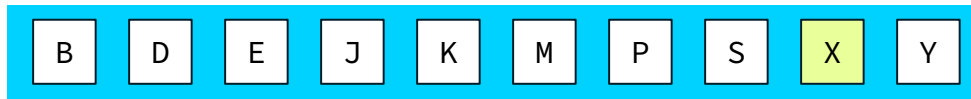


{u: "xdg", e: 20220512}





Files → sort docs by the shard key



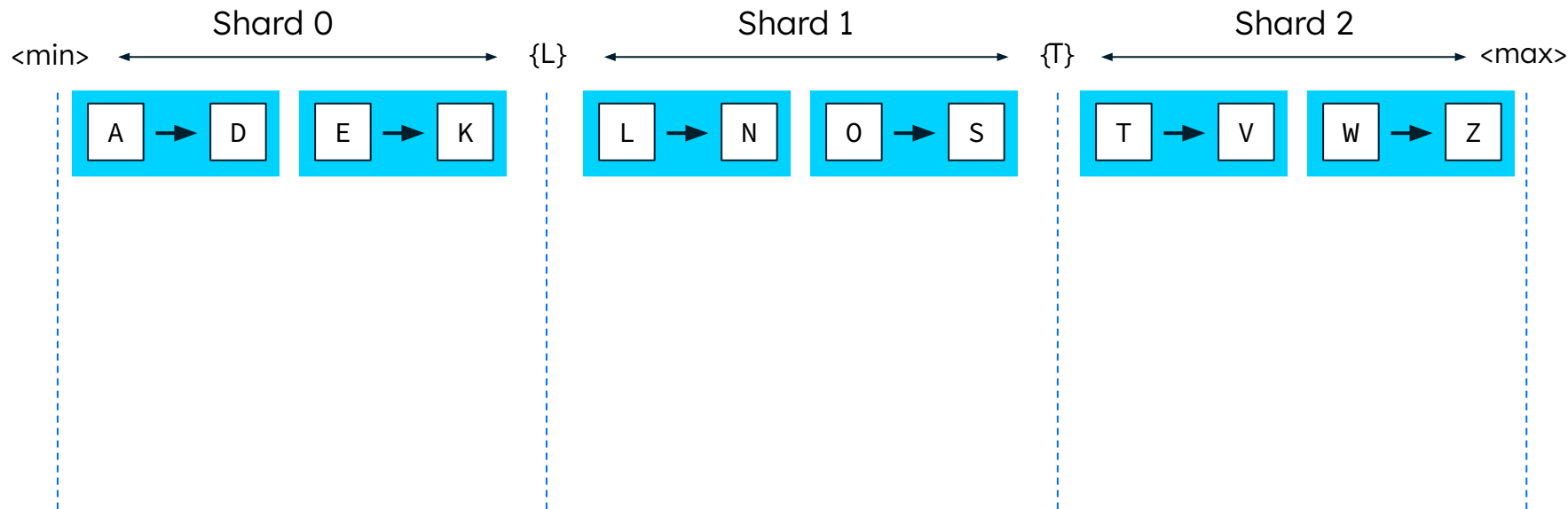


Files → sort docs by the shard key





Key space → every file has a place





Big files

(but not too big)

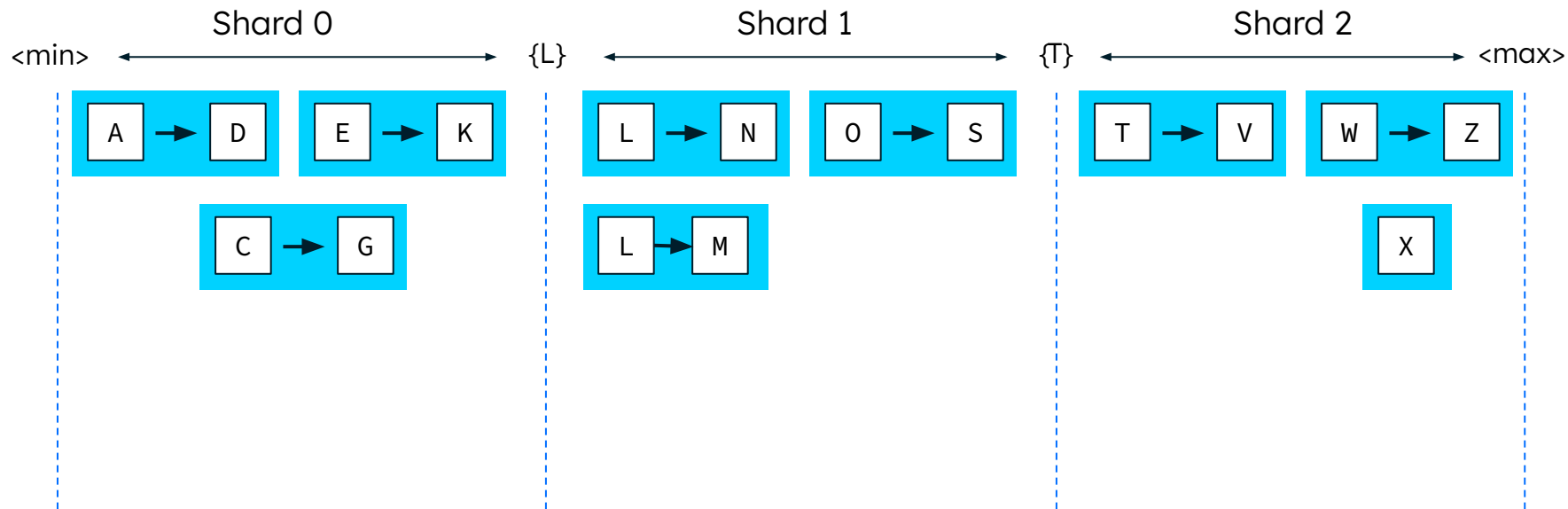


Columns $\rightarrow$ Parquet format

(usually)

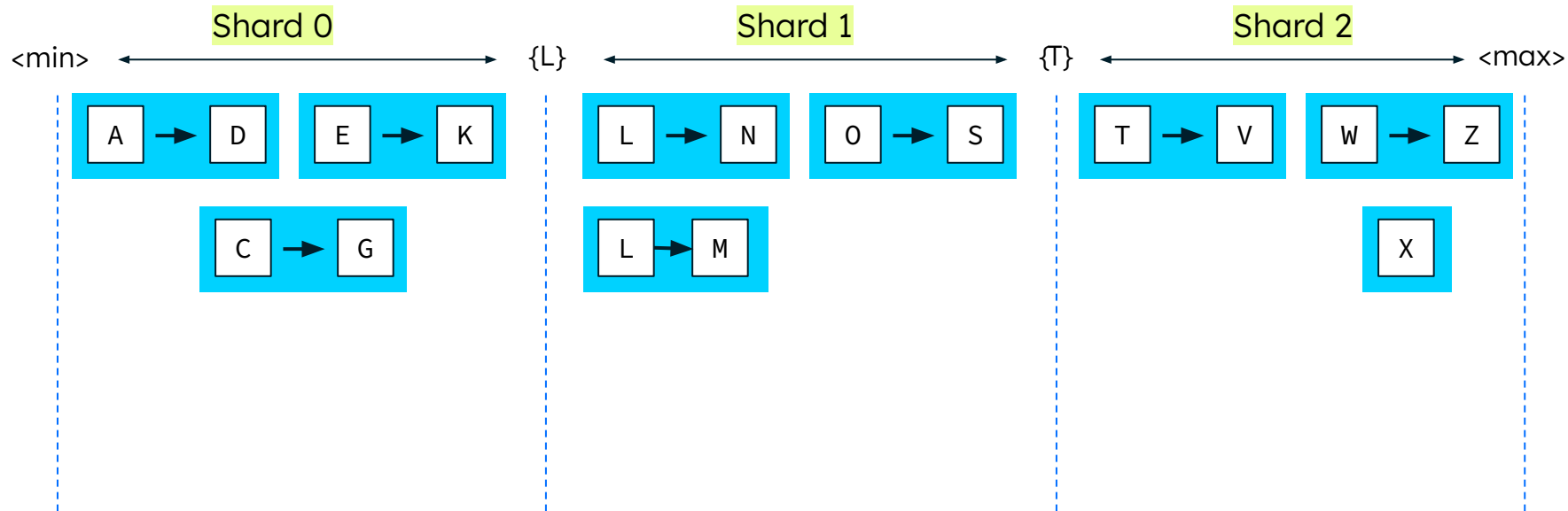


File overlaps are okay!



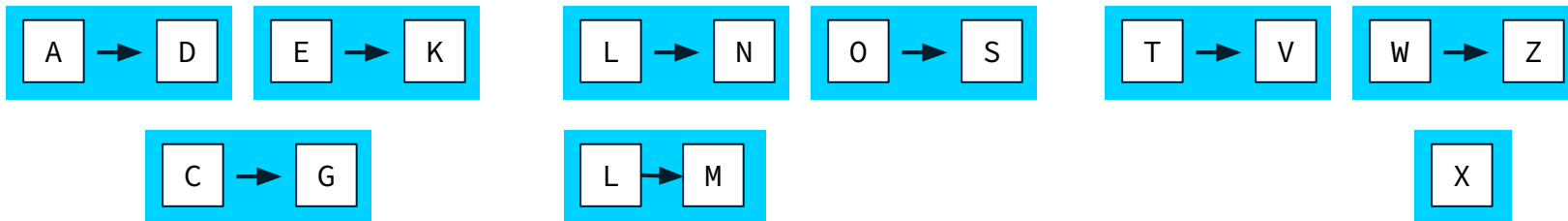


Shard metadata stored outside S3



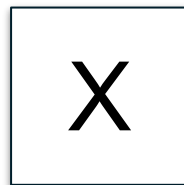


S3 only contains files





{u: "xdg", e: 20220512}

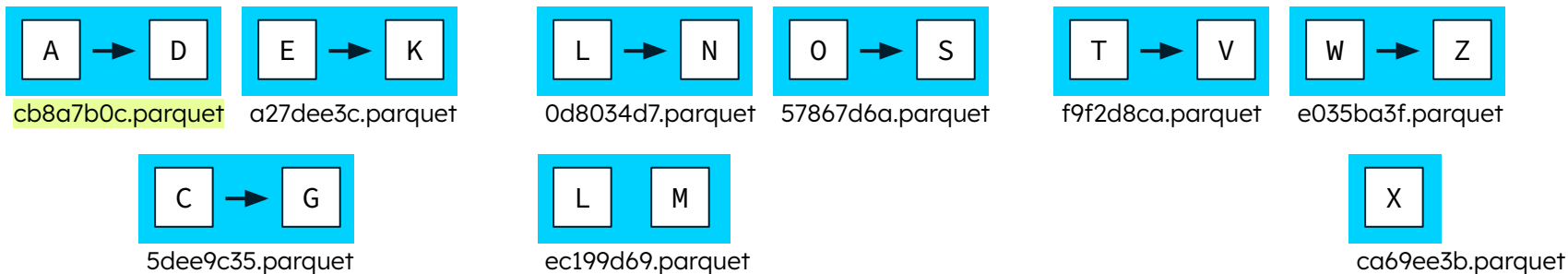




File name $\neq$ values



Unique name from contents





Optimized storage

- ✓ Shard key to group related data
- ✓ Big files, but not too big
- ✓ Column storage (if appropriate)



‘ca69ee3b.parquet’

??



Change #2:

Metadata catalog



Metadata about your data



File metadata

Physical location ("s3://<bucket>/.../cb8a7b0c.parquet")

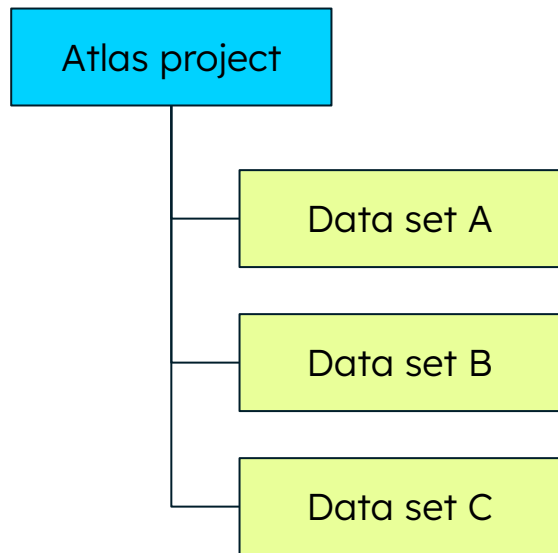
Upper/lower shard key bounds

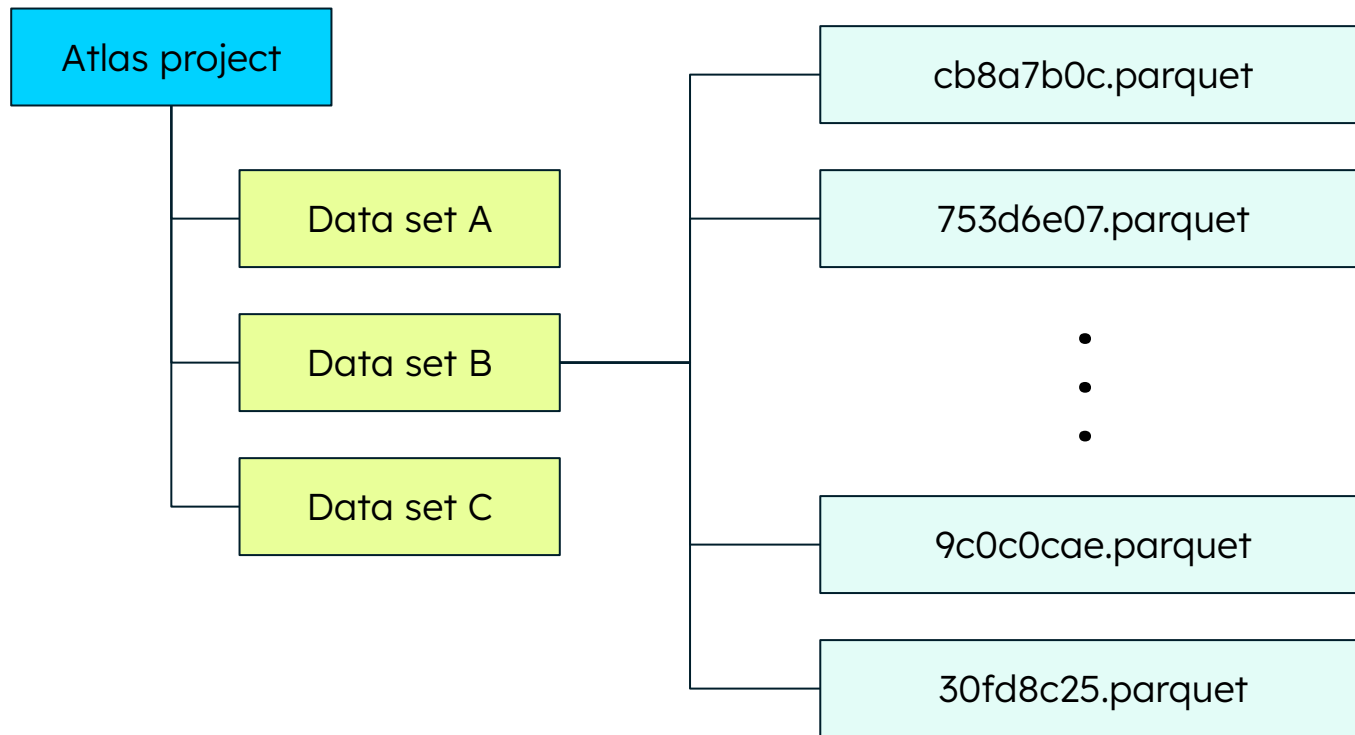
File size

Document count



Project → holds data sets







Data sets → Collection

(Atlas Data Federation)



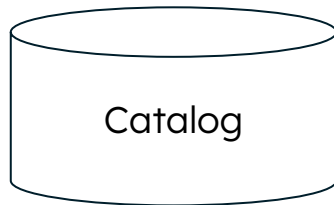
Query plan $\rightarrow$ list of files



```
{ $match: { "user_name": "xdg" } }
```

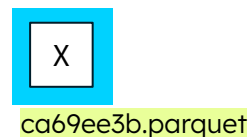
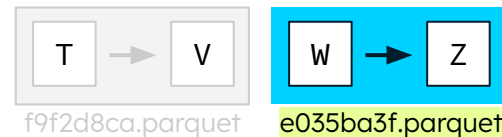
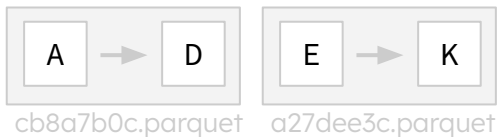
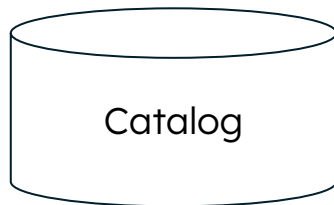


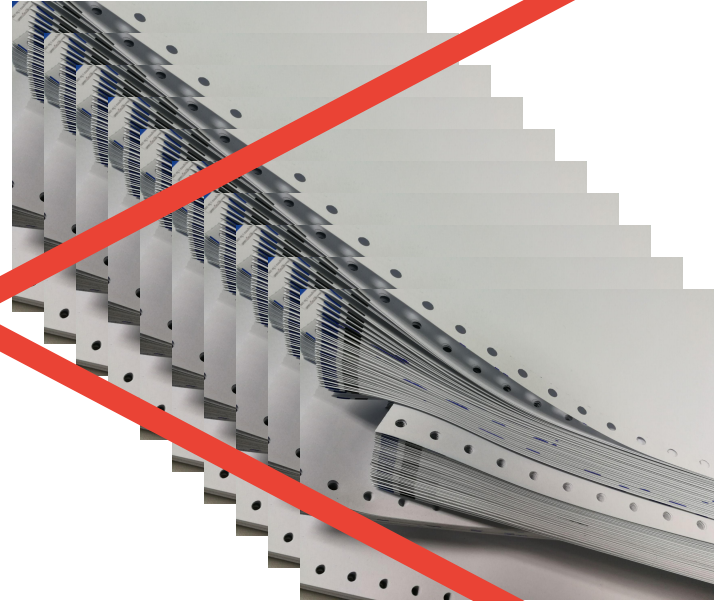
```
{ $match: { "user_name": "xdg" } }
```





`{$match: {"user_name": "xdg"}}`





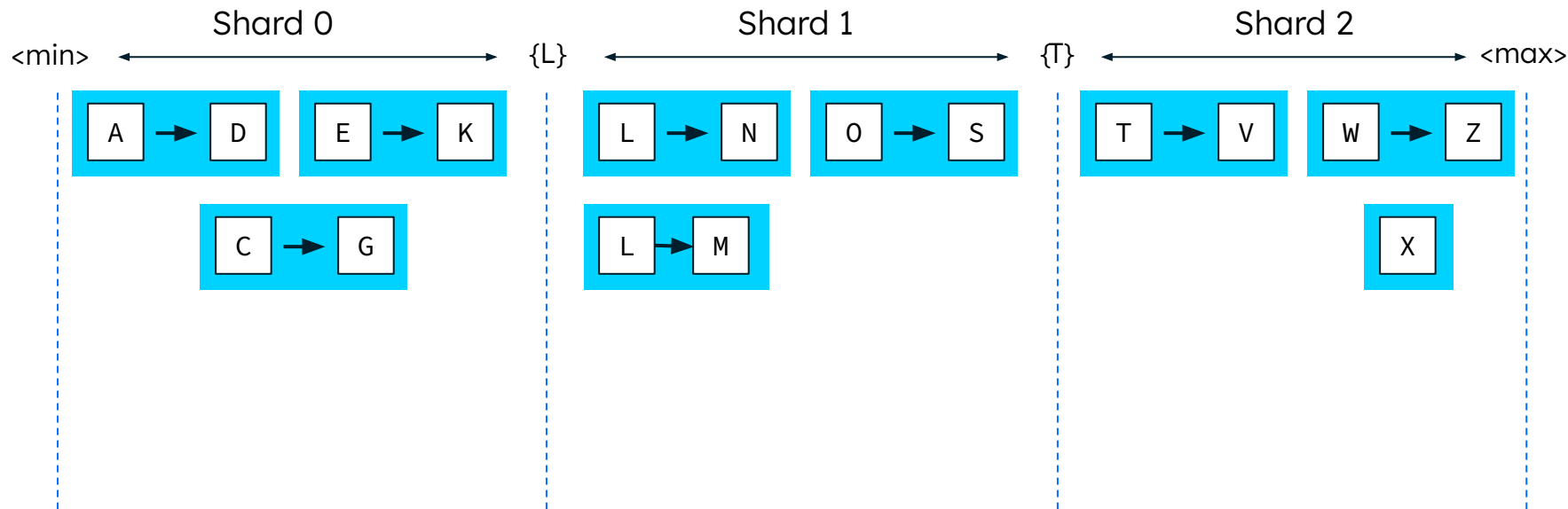


Change #3:

Auto-rebalancing



How did the overlaps happen?

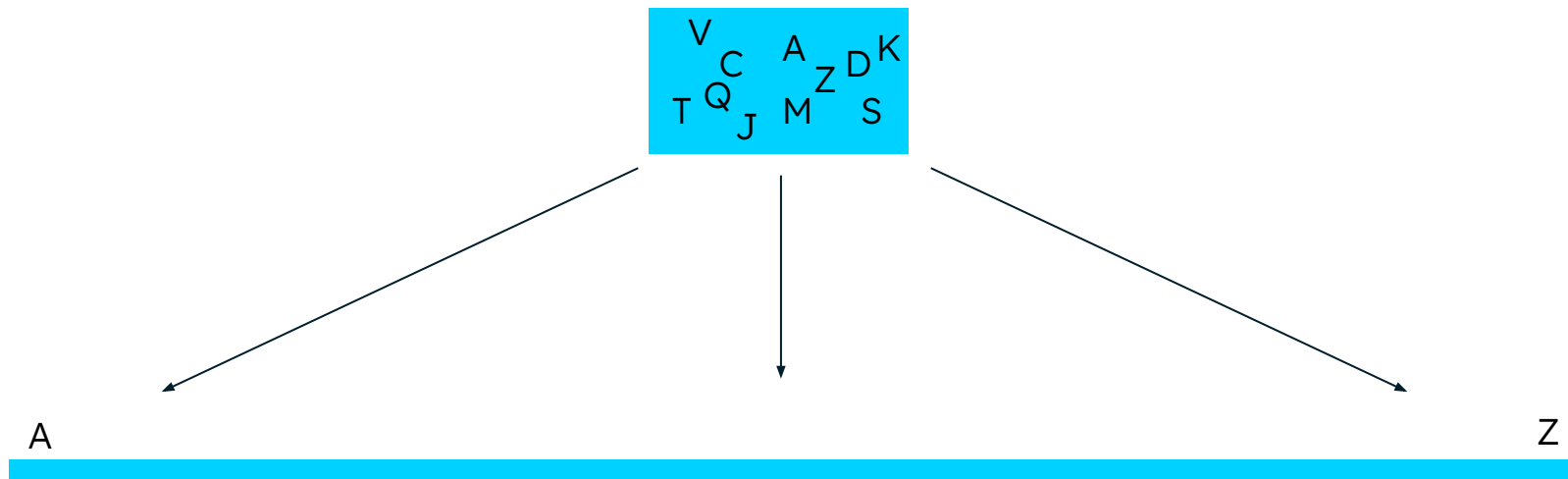




Continuous ingestion

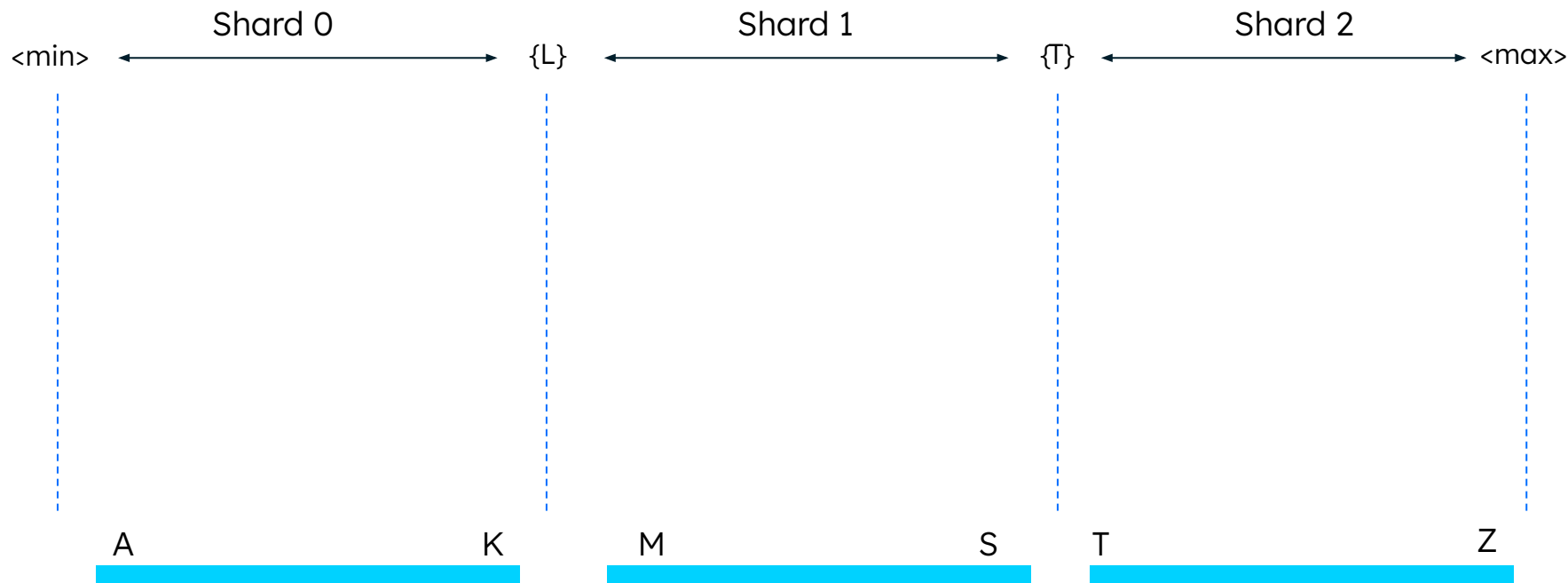


Sort new data along the keyspace



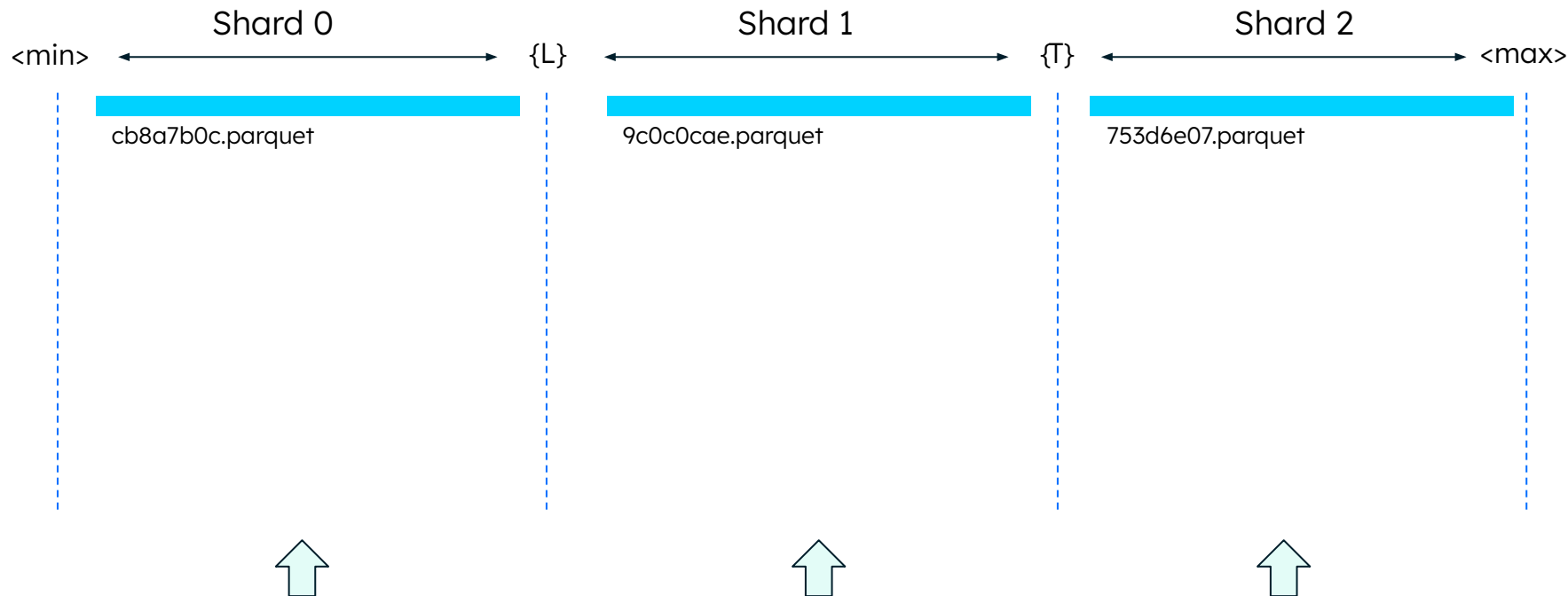


Split along shard boundaries



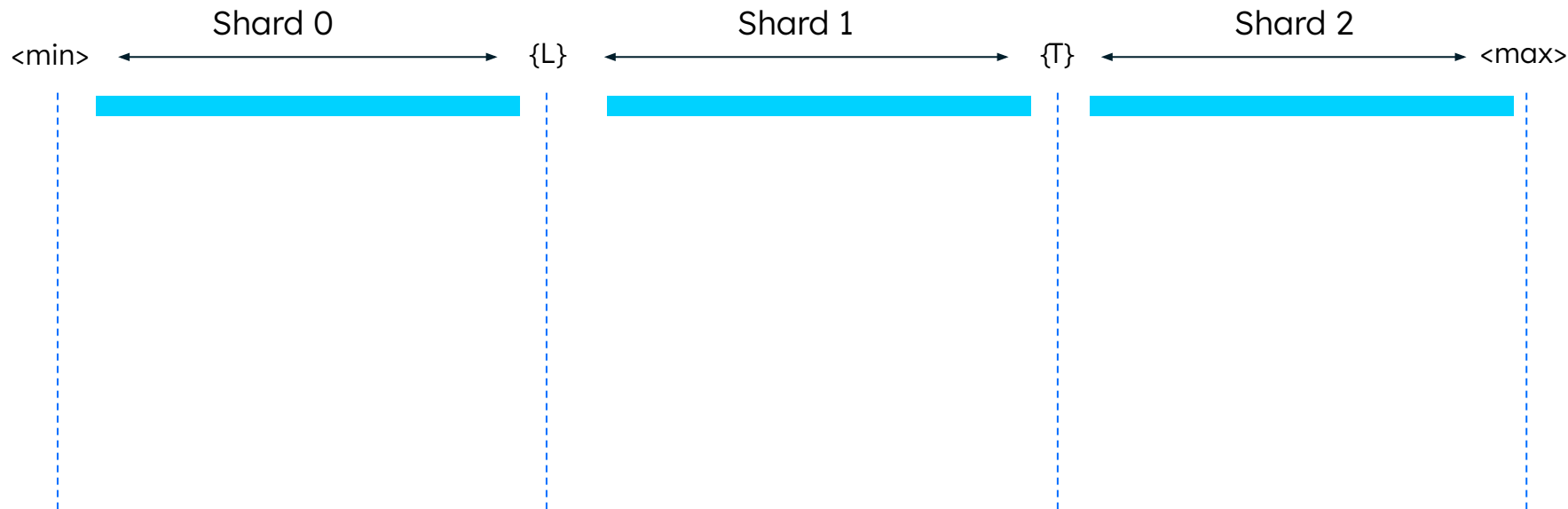


Upload to S3



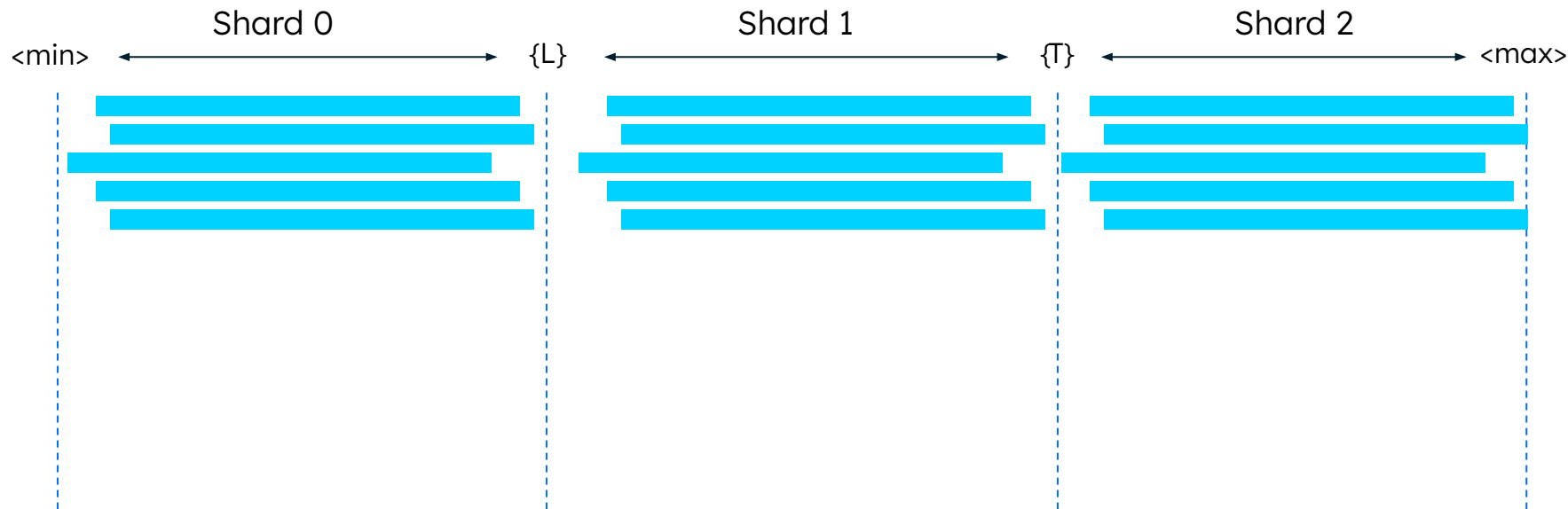


Files might be small





Small files accumulate over time





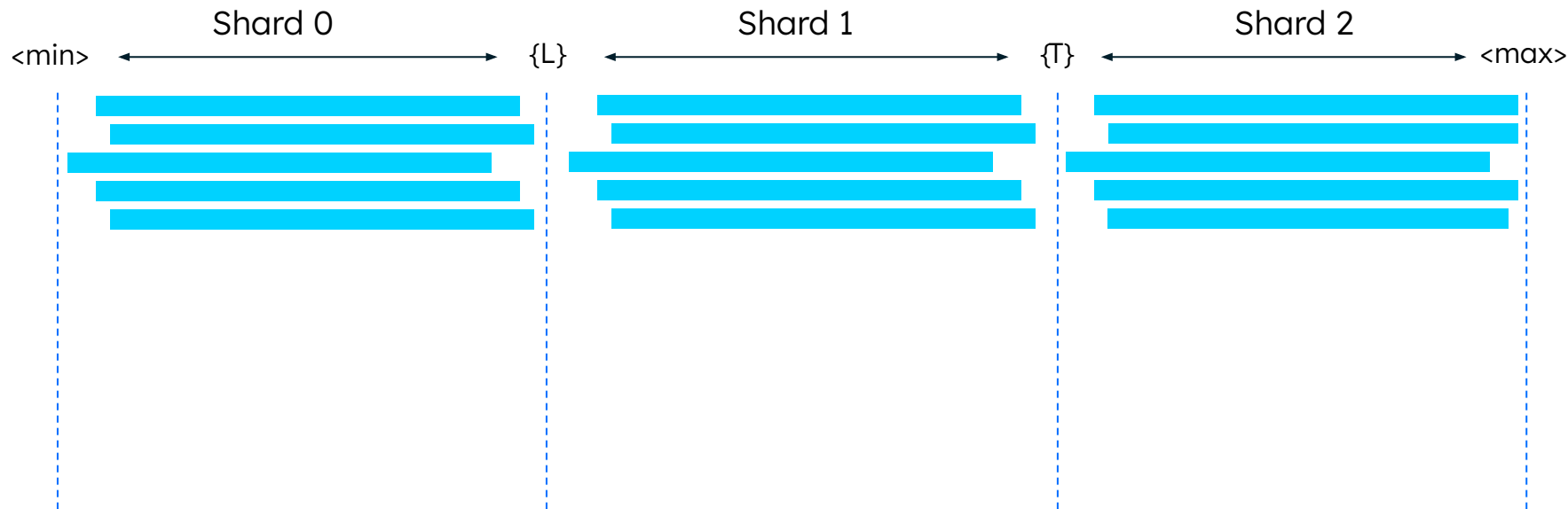
Auto-rebalancing!



Heuristic triggers

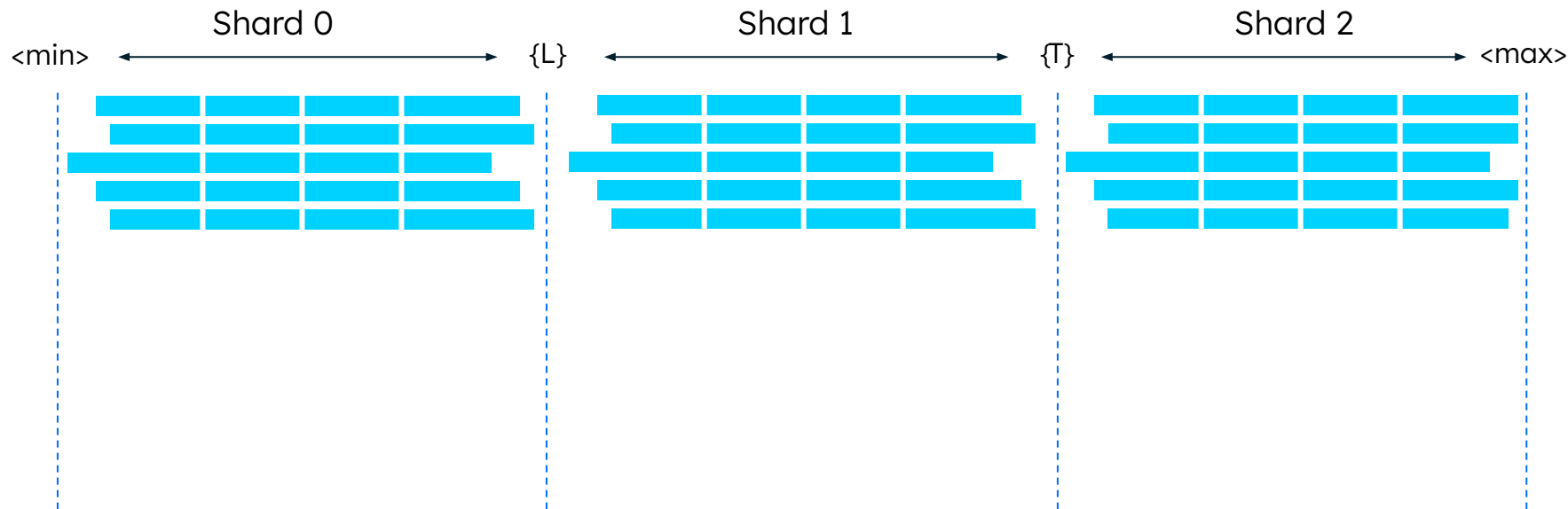


Small files? → Compaction



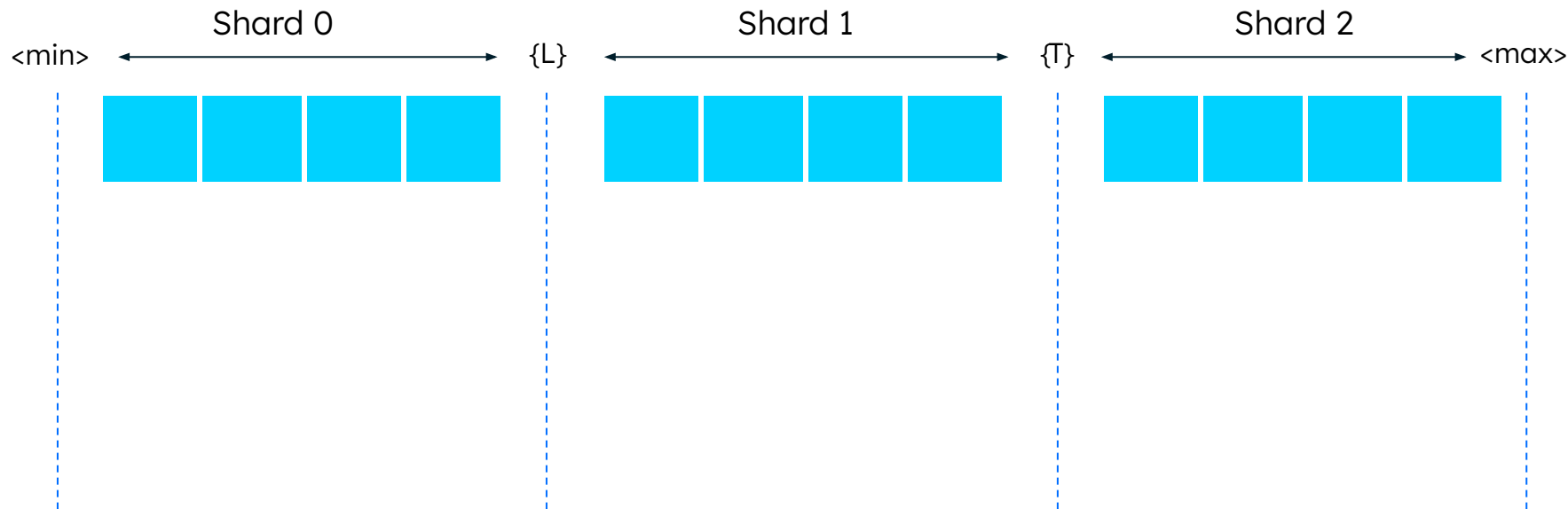


Split files...





...sort and combine

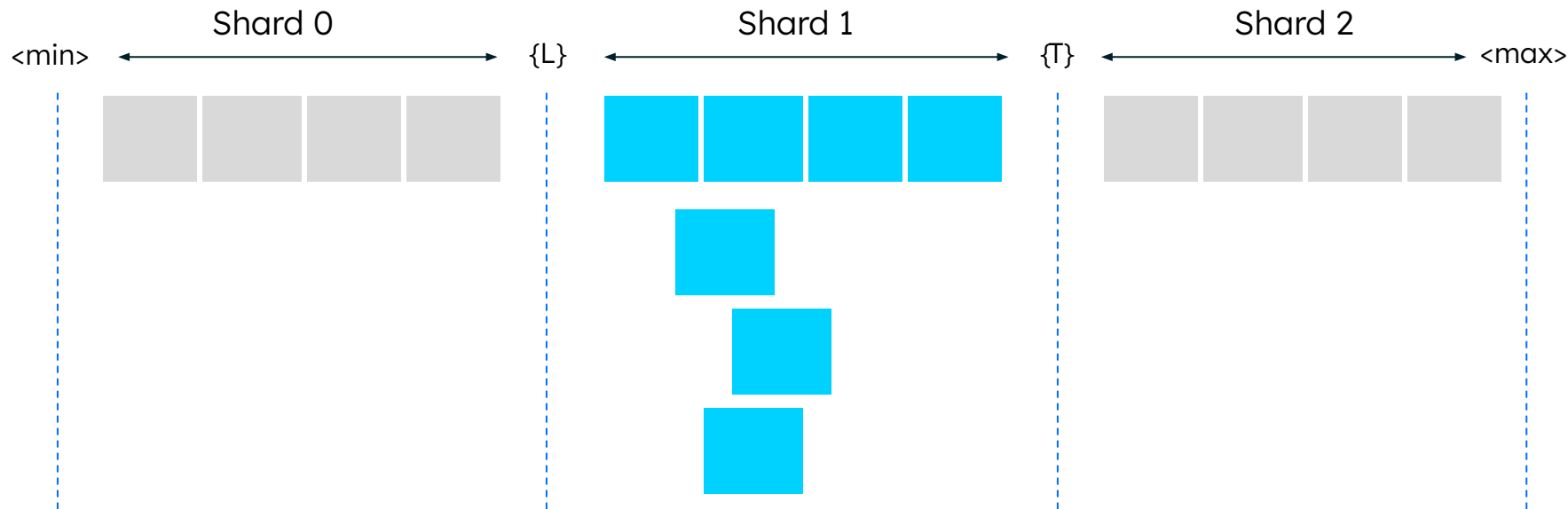




Works shard
by shard

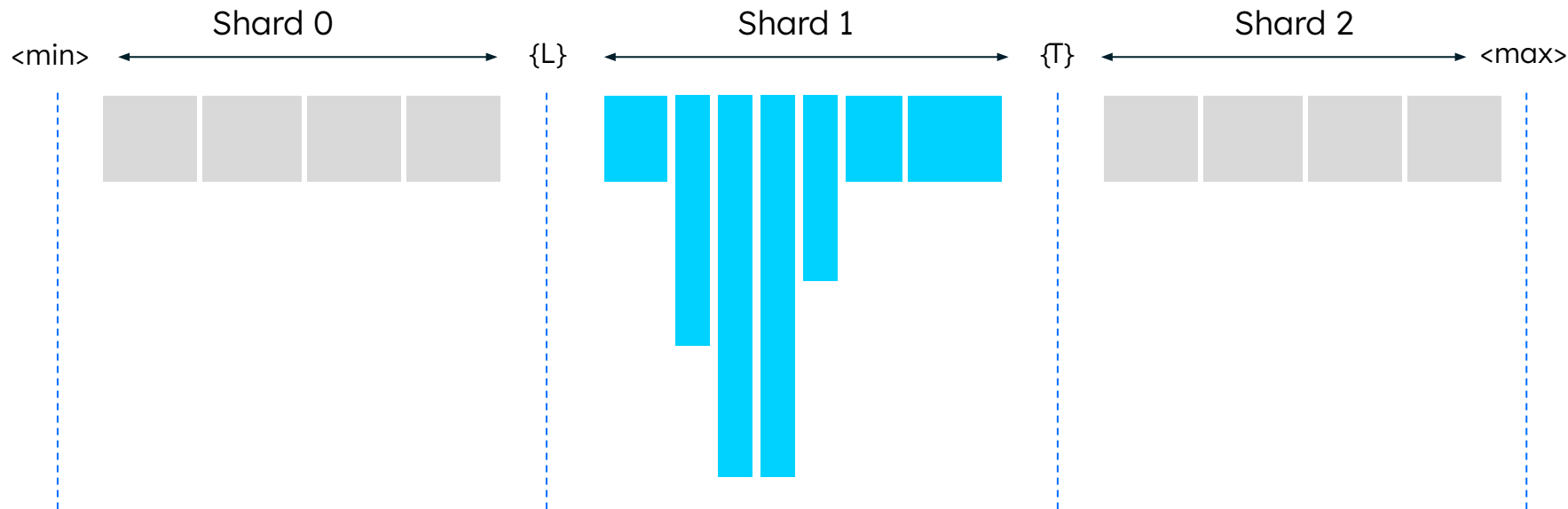


Hot shard? → Rebalance



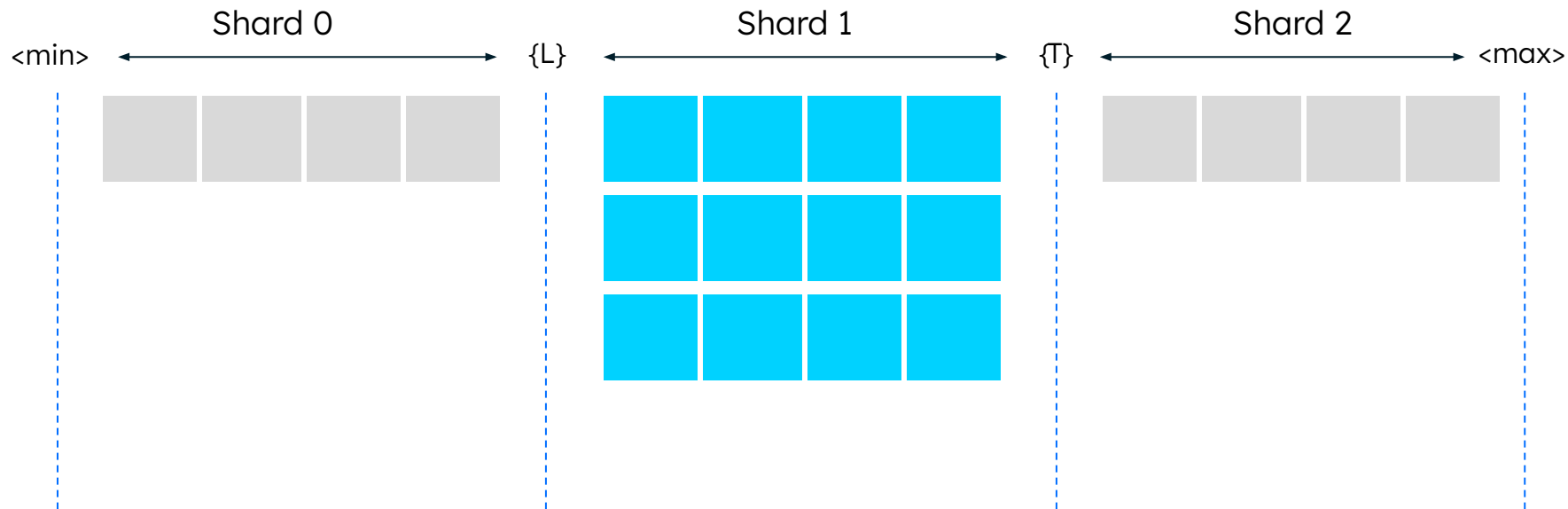


Rewrite with narrower ranges



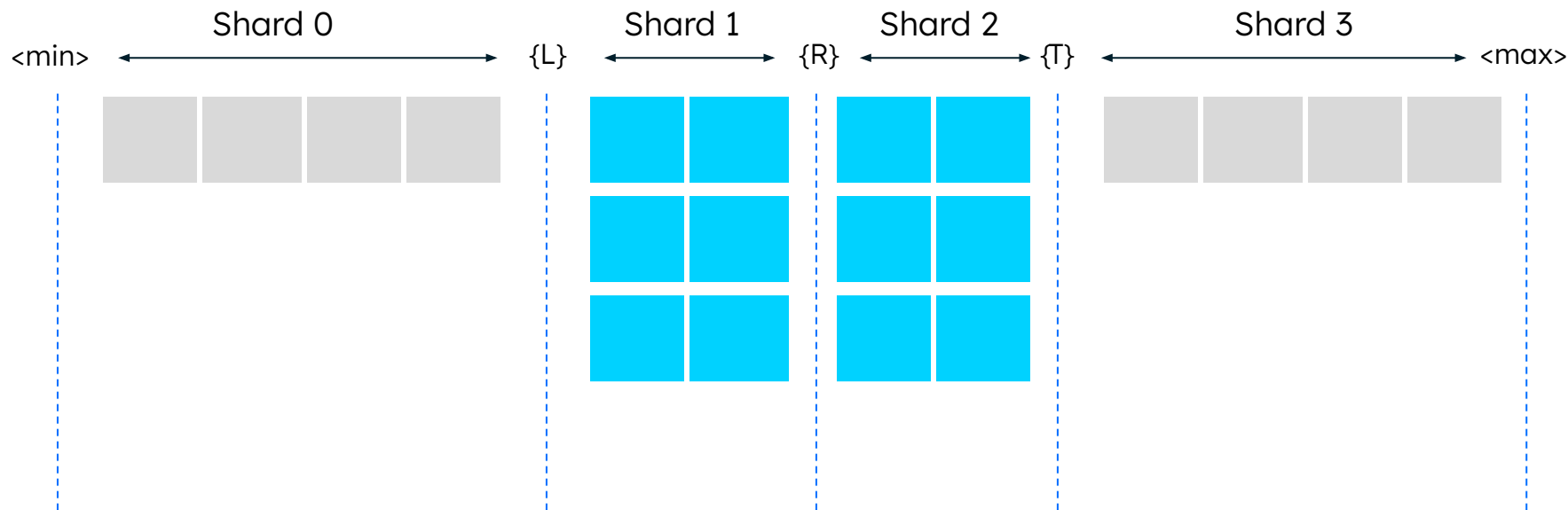


Big shard?





Big shard? → Split it





Continuous optimization

In summary...







Minimize...

Number of network trips

Fetching unneeded documents

Fetching unneeded fields



Better performance

Optimized storage

Metadata catalog

Auto-rebalancing



Roadmap

Available now...

Snapshot ingestion
(Preview)



Snapshot ingestion

Ingest Cloud Backup Snapshots

High performance, low cost historical data

Join across history for timeseries analysis



Roadmap

Available now...

Snapshot ingestion
(Preview)

Coming soon...

Migrate online archive
to optimized storage

Future...

New ingestion sources



Curious? Try it out!

The screenshot shows the MongoDB Atlas web interface. The top navigation bar includes 'Mongo...', 'Access Manager', 'Billing', 'All Clusters', 'Admin', 'Get Help', and a user profile 'D'. Below this, a secondary bar shows 'MDBW2...', 'ACTIVE', and tabs for 'Atlas', 'Realm', and 'Charts'. The left sidebar is divided into sections: 'DEPLOYMENT' (with 'Database' and 'Data Lake' - the latter is circled in red and labeled 'PREVIEW'), 'DATA SERVICES' (with 'Triggers', 'Data API' - labeled 'PREVIEW', and 'Data Federation'), and 'SECURITY' (with 'Database Access', 'Network Access', and 'Advanced'). The main content area features a blue informational banner, a heading 'Focus on your data and analytics instead of managing infrastructure', a descriptive paragraph 'Manage project-level dataset stored in columnar cloud object storage for fast analytic performance', and a large green button labeled 'Create Data Lake Pipeline'. A red arrow points from the 'Data Lake' sidebar item to this button.

Mongo... Access Manager Billing All Clusters Admin Get Help D

MDBW2... ACTIVE Atlas Realm Charts

DEPLOYMENT

Database

Data Lake PREVIEW

DATA SERVICES

Triggers

Data API PREVIEW

Data Federation

SECURITY

Database Access

Network Access

Advanced

Please go to [Data Federation](#) to see past Data Lakes you created.

Focus on your data and analytics instead of managing infrastructure

Manage project-level dataset stored in columnar cloud object storage for fast analytic performance

Create Data Lake Pipeline





Data Lake $\neq$ Boxes



Atlas Data Lake → Answers

Thank you for
your time