



Caching Framework for Exabyte-Scale Data Lakes

(COMMUNITY OVER CODE 2023)

Jing Zhao, Uber

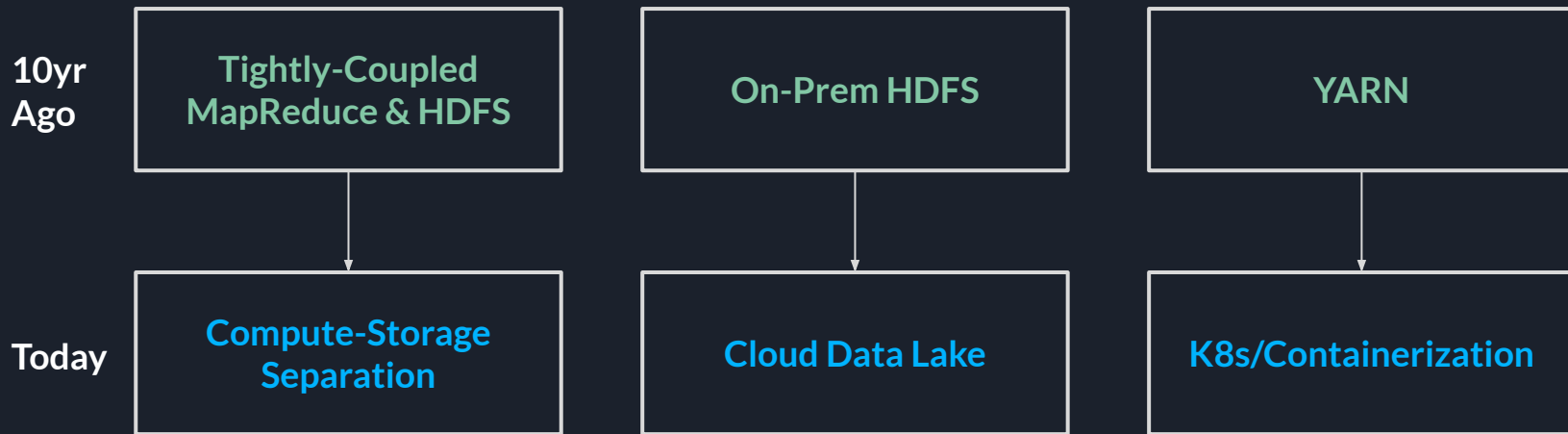
Hope Wang, Alluxio



Agenda

- Data locality challenges in large-scale data lakes
- Alluxio: An open-source caching framework
 - Improve both performance and cost efficiency
 - Distributed cache and embedded cache
- Uber: Using cache for exabyte-scale data lake
 - HDFS DataNode local cache for high-density HDD adoption
 - Presto local cache

The Evolution of the Modern Data Stack



More Elastic, Cheaper, Easier to Manage, More Scalable



We are Losing Data Locality

Today

Compute-Storage
Separation

Cloud Data Lake

K8s/Containerization

Data locality is missing ...

The Implications of Losing Data Locality

Slow & Expensive



Slow and inconsistent data access performance



Fast-growing cloud storage costs, including API and egress costs



High data operation costs when migrating to the cloud

Complex Platform



Data copies, synchronization costs



Multiple APIs necessitate integration and application rewrites



On-premise, cloud, hybrid, multi-cloud environments all have different environment properties

10%

Of Your Data is Hot Data



Caching's Values



**Boost
Performance**



**Prevent
Network
Congestion**

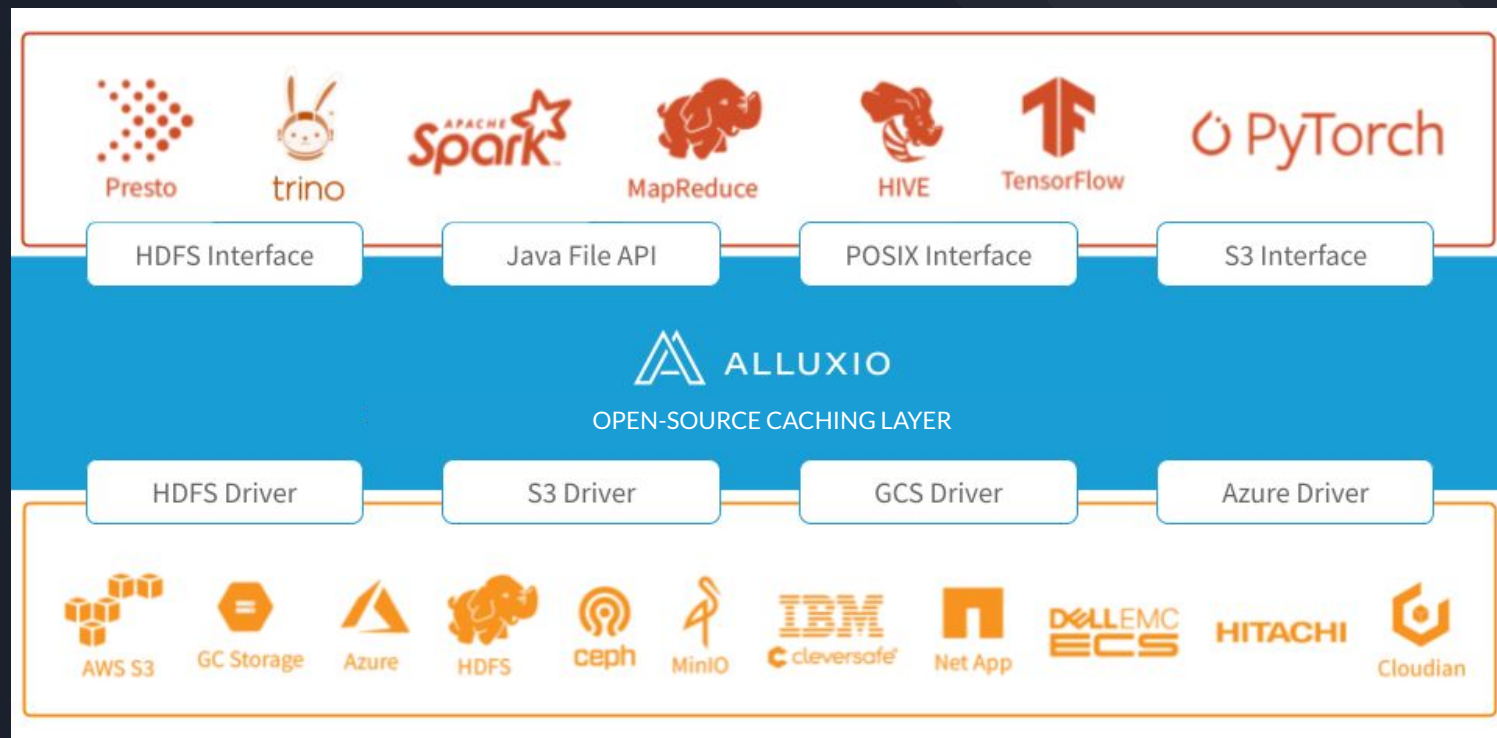


Save Costs



**Offload
Under
Storage**

Alluxio: A Critical Caching Framework in the Open-Source Data Stack



Open-Source Started From UC Berkeley AMPLab in 2014

Join the
conversation on
Slack
alluxio.io/slack



1,200+ contributors
& growing

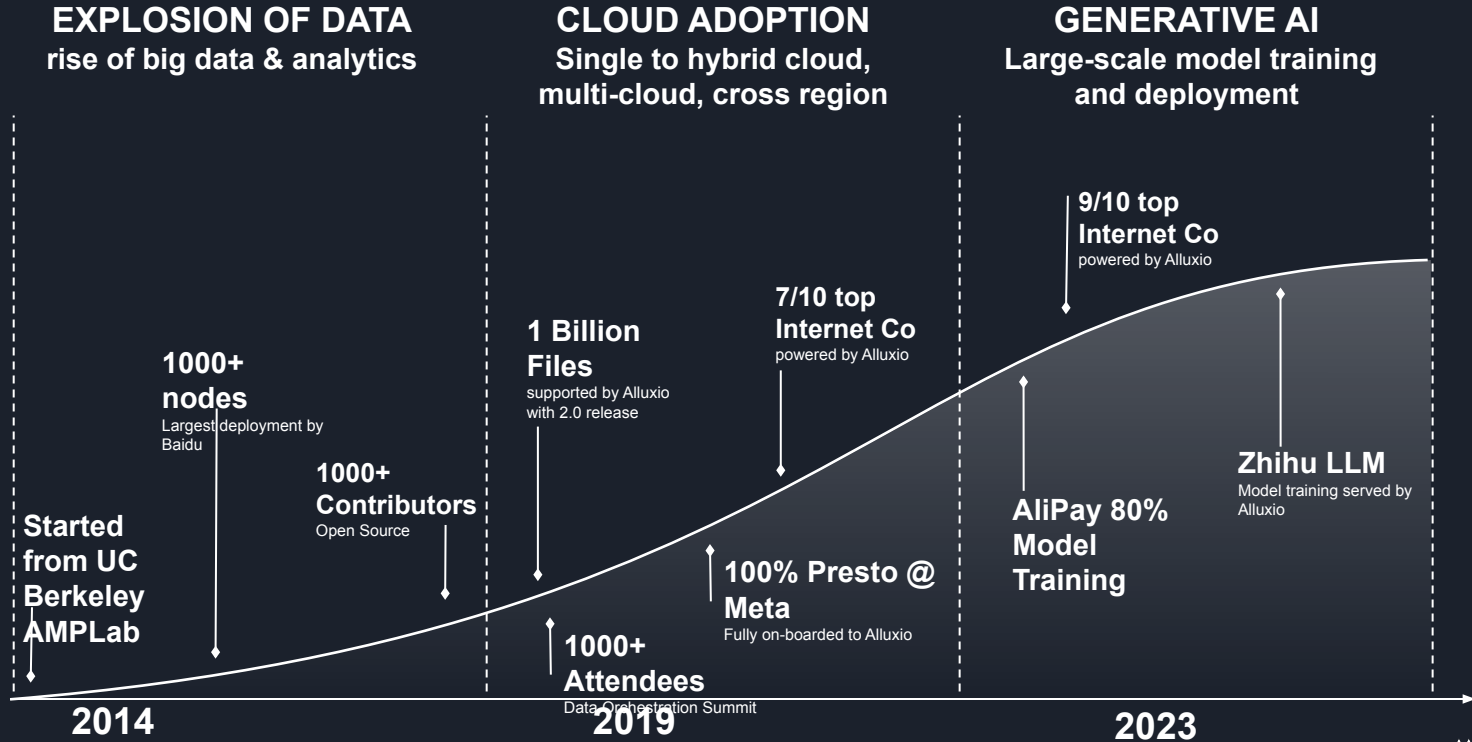
11,000+ Slack
Community Members



Top 10 Most Critical Java
Based Open Source Project

GitHub's Top 100 Most
Valuable Repositories
Out of 96 Million

Alluxio Technology Journey



Powered by Alluxio

TECHNOLOGY

INTERNET



PUBLIC CLOUD PROVIDERS



GENERAL



E-COMMERCE



OTHERS



FINANCIAL SERVICES



TELCO & MEDIA



[LEARN MORE](#)

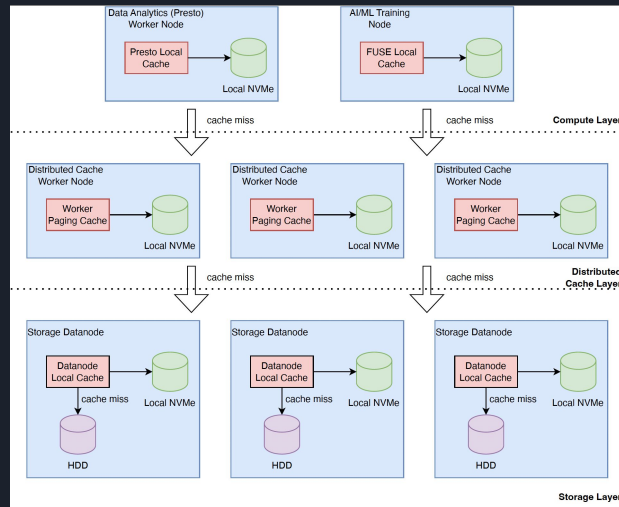
A Caching Framework to Fit Different Needs

Alluxio Embedded/Local Cache

- Run as a library in the application processes (Presto, HDFS DataNode)
- Leverage local disk NVMe or memory
- When the size of hot data fits local disks

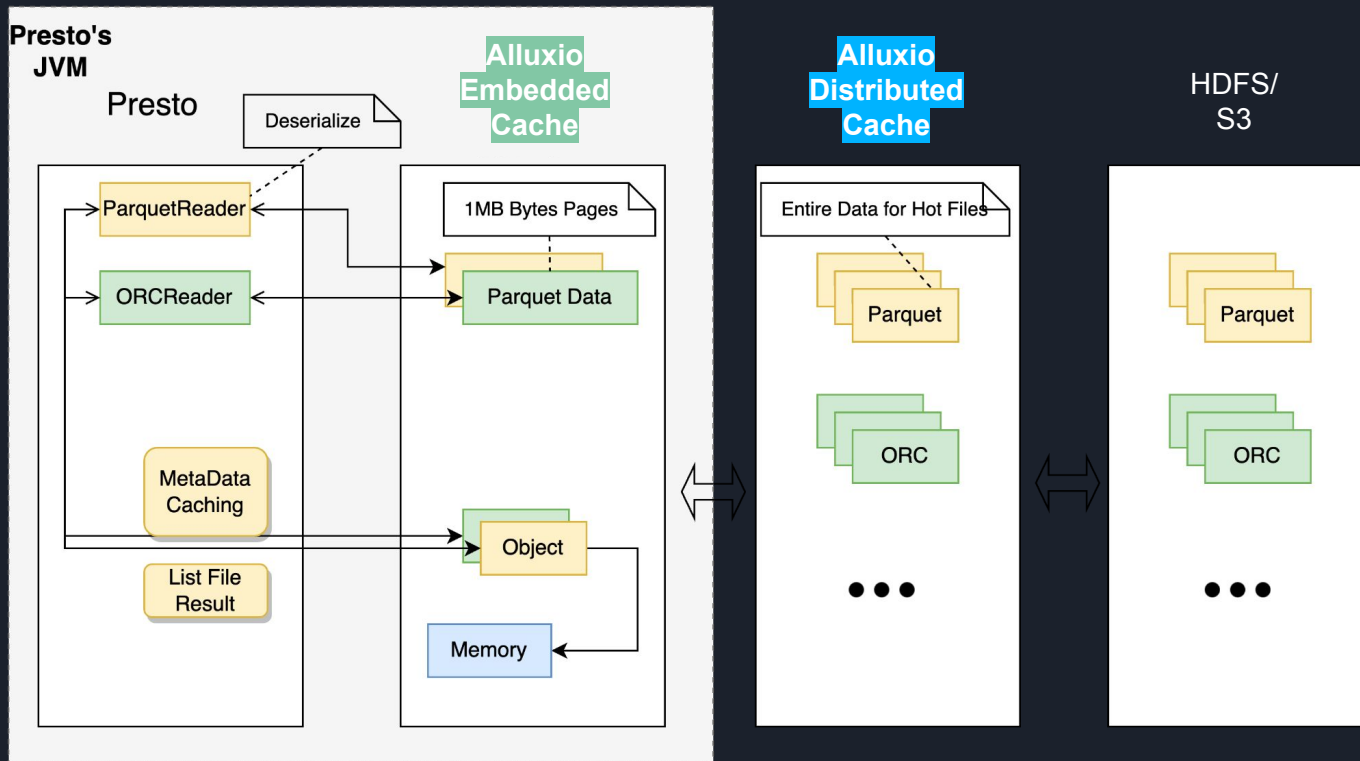
Alluxio Distributed Cache

- Standalone cache service shareable across applications
- Cache capacity scales horizontally



Multi-level Caching

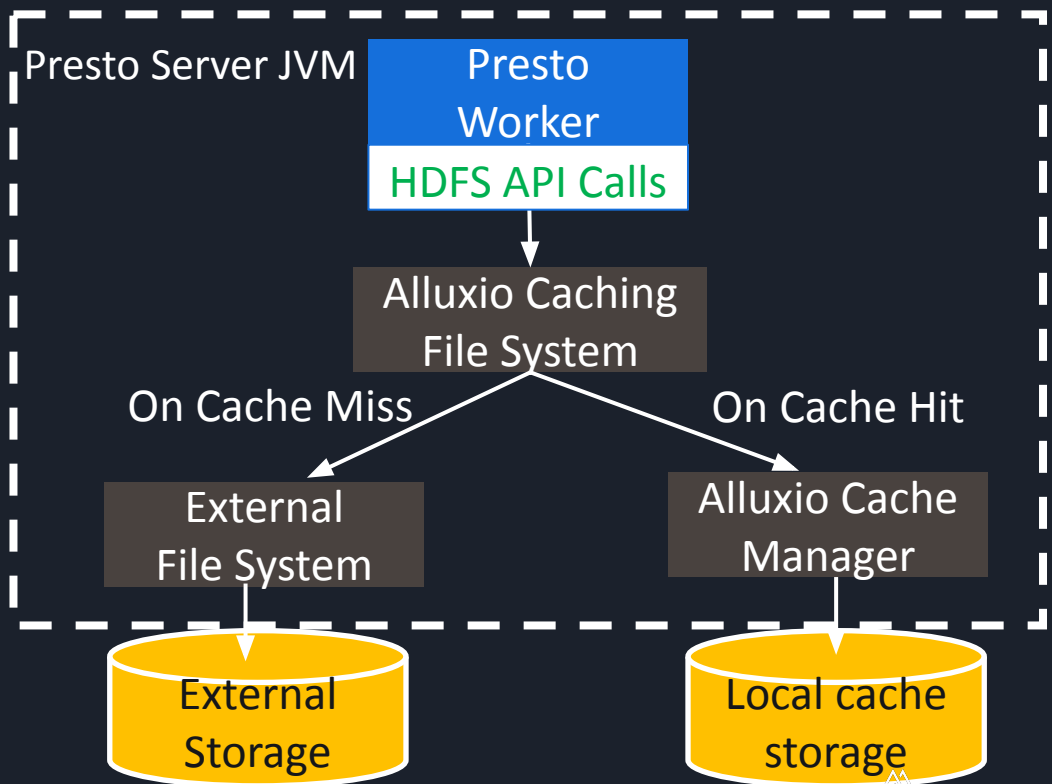
L1 Embedded Cache+L2 Distributed Cache



Alluxio Embedded Cache in Presto

Embedded
Cache

- Battle tested in Uber, Meta, Tiktok and etc.
- Support Iceberg, Hudi, Delta Lake and Hive tables
- Support varied file format such as Paquet, ORC and CSV
- Fully optimized for local NVMe storages





Alluxio Embedded Cache Data Management

Embedded
Cache

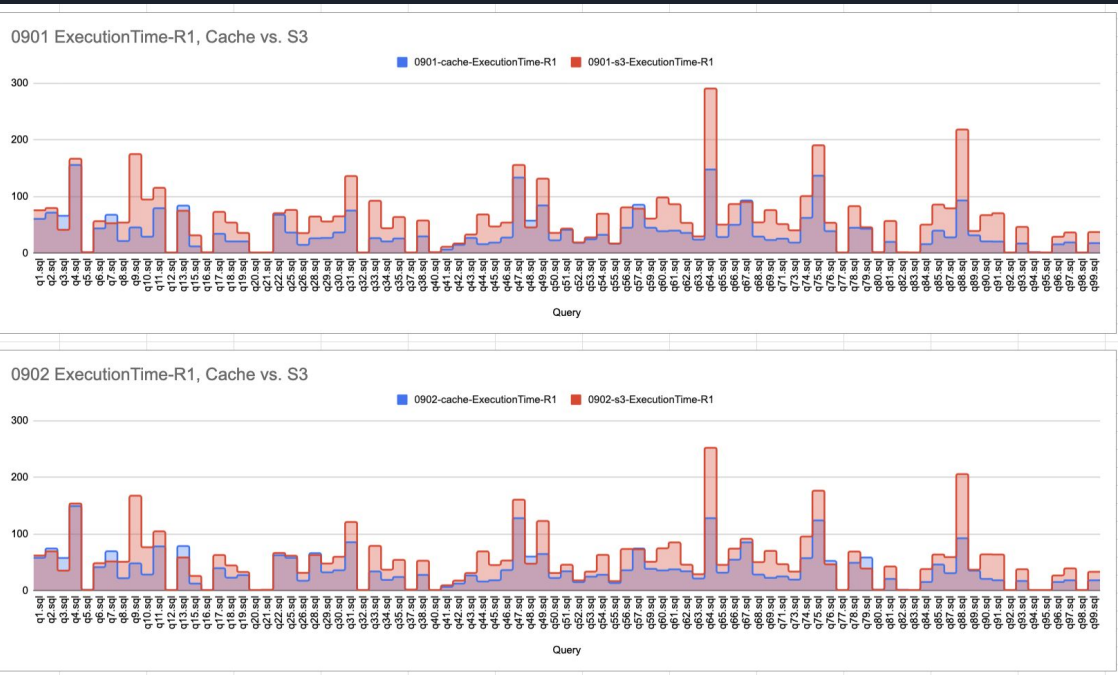
Alluxio embedded cache provides cache eviction and admission

- Support LRU and FIFO cache eviction policy
- Support customized cache admission policy
- Support TTL
- Support data quota

TPC-DS Benchmark of Presto+Embedded Cache

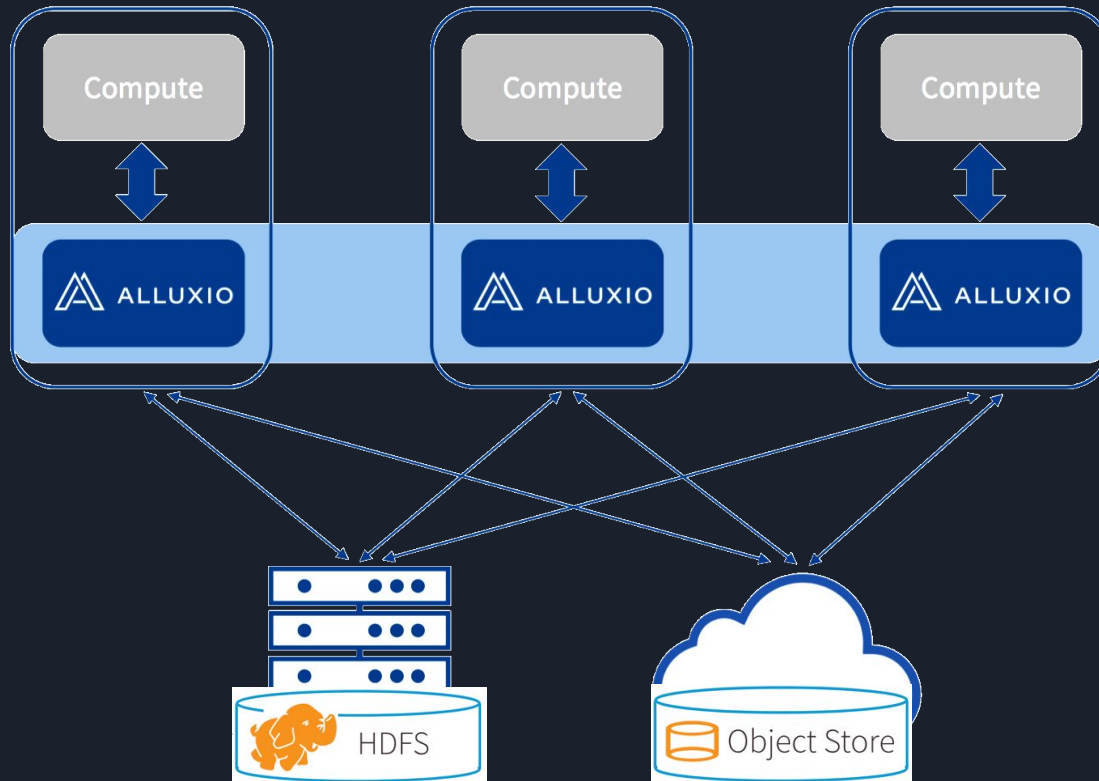
Embedded
Cache

- Accelerate 70 out of 73 queries that lasting $> 2s$
- Total running time q1-q99 improved by 63% on average



Alluxio Distributed Cache

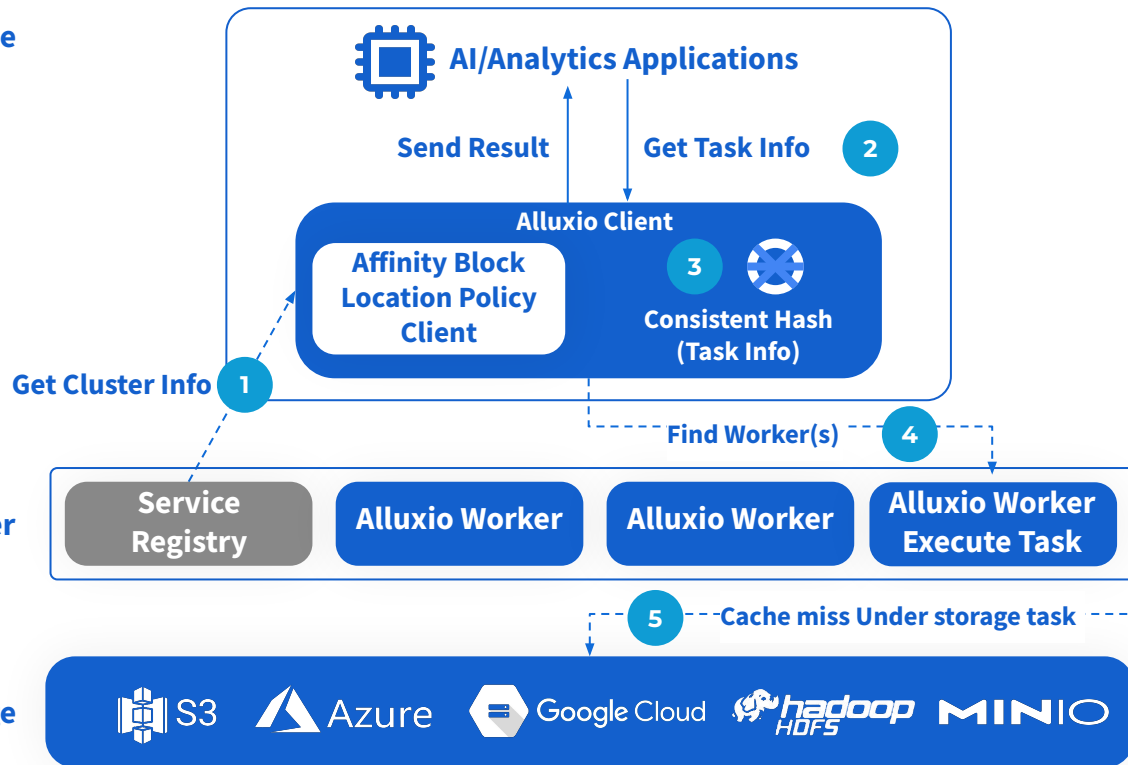
Distributed
Cache



Alluxio Cluster Architecture

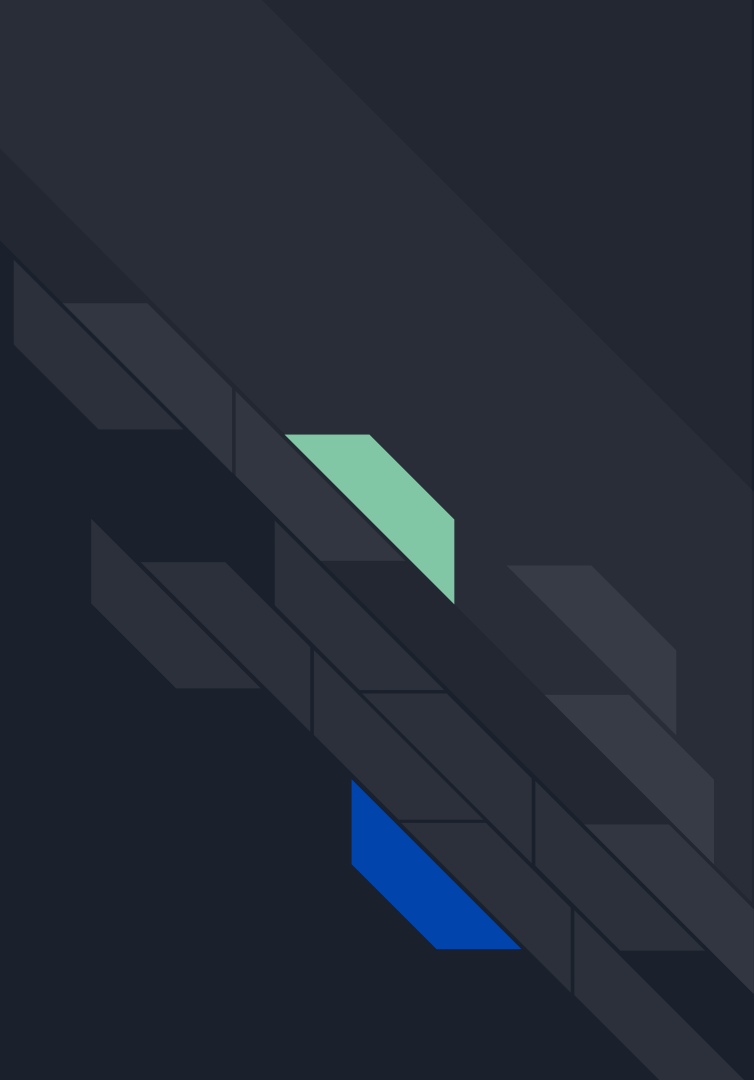
Distributed
Cache

Compute Node



Uber: Using Cache for Exabyte-scale Data Lake

Uber



Data informs **every decision** at Uber



Community
Operations



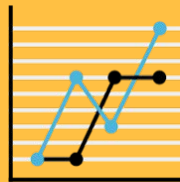
Marketplace
Pricing



Eats



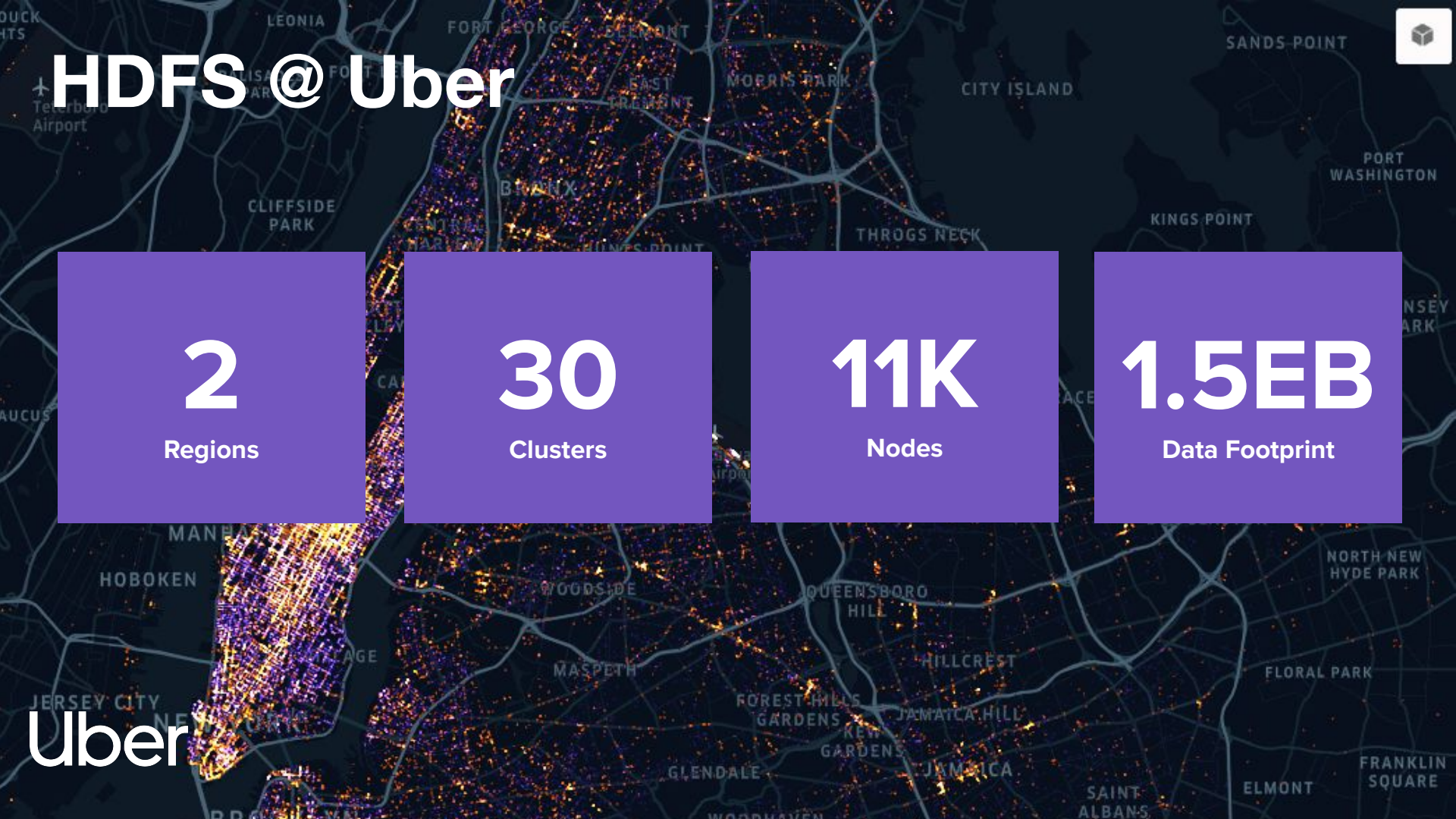
Compliance



Growth Marketing



Data Science



HDFS @ Uber

2

Regions

30

Clusters

11K

Nodes

1.5EB

Data Footprint

Uber

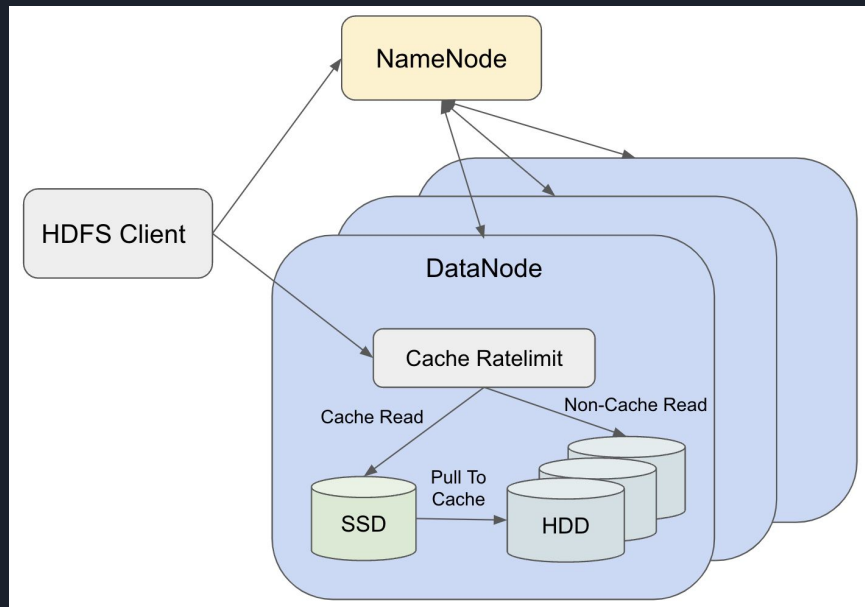
Adopting High-Density HDD in HDFS

- Capacity per Host: 4TB * 24 → 16TB * 35
- Efficiency: >50% HW cost reduction
- Challenges
 - **DataNode IO performance**
 - HDFS reliability (blast radius)
- Opportunities
 - Traffic focuses on a small group of extremely hot blocks
 - Top 10K blocks attracted >90% read traffic

Host	Total reads	Total writes	Number of Blocks stored in the host	Number of blocks being read	Average Block Size	Capacity Usage	Read traffic on top 10k blocks (1 hour time window)
Host1	13.5M	3.3 K	1,074,622	84769	380 MB	77%	89%
Host2	12.8M	4.7 K	633,923	59376	330 MB	79.89%	94%
Host3	11.3M	275 K	479,544	49317	300 MB	79.86%	91%
Host4	8.5M	4.6 K	247,206	31048	300 MB	82.85%	99%
Host5	14.3M	45 K	463,819	79958	160 MB	81.62%	99%

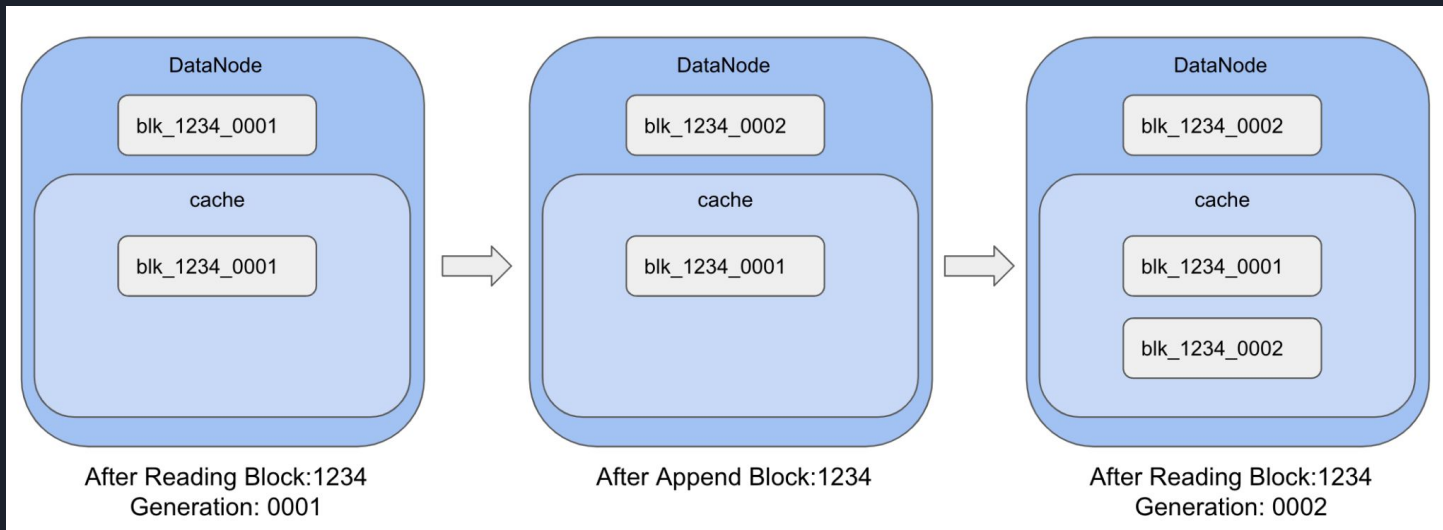
DataNode Local Cache

- Build a local cache within DataNode
 - 4TB NVMe SSD disk
 - Based on DataNode local traffic
- Leverage **Alluxio** for cache management
 - Page-level cache
 - 1MB default page size matches traffic pattern
- Exclude Non-Client Reads



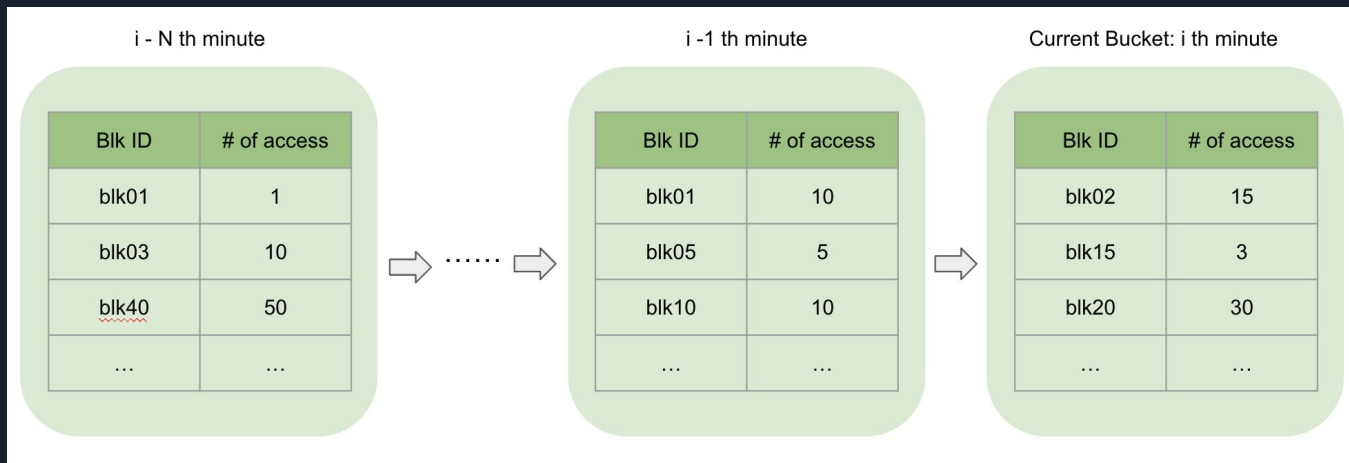
Support Append to a Block

- HDFS blocks may not be immutable due to Append ops
- Leverage HDFS generation stamp to achieve “snapshot isolation”



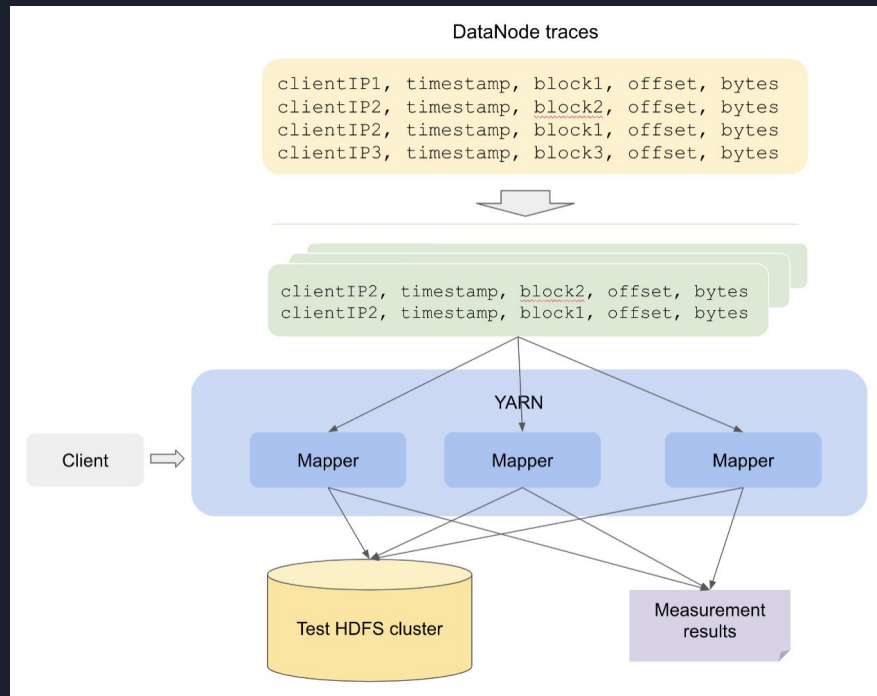
Rate Limiter

- 1:140 SSD-to-HDD ratio
- Need mechanisms to control what data can be loaded into Cache
 - Cache hit rate
 - SSD write endurance
- Track block access frequency within a sliding time window



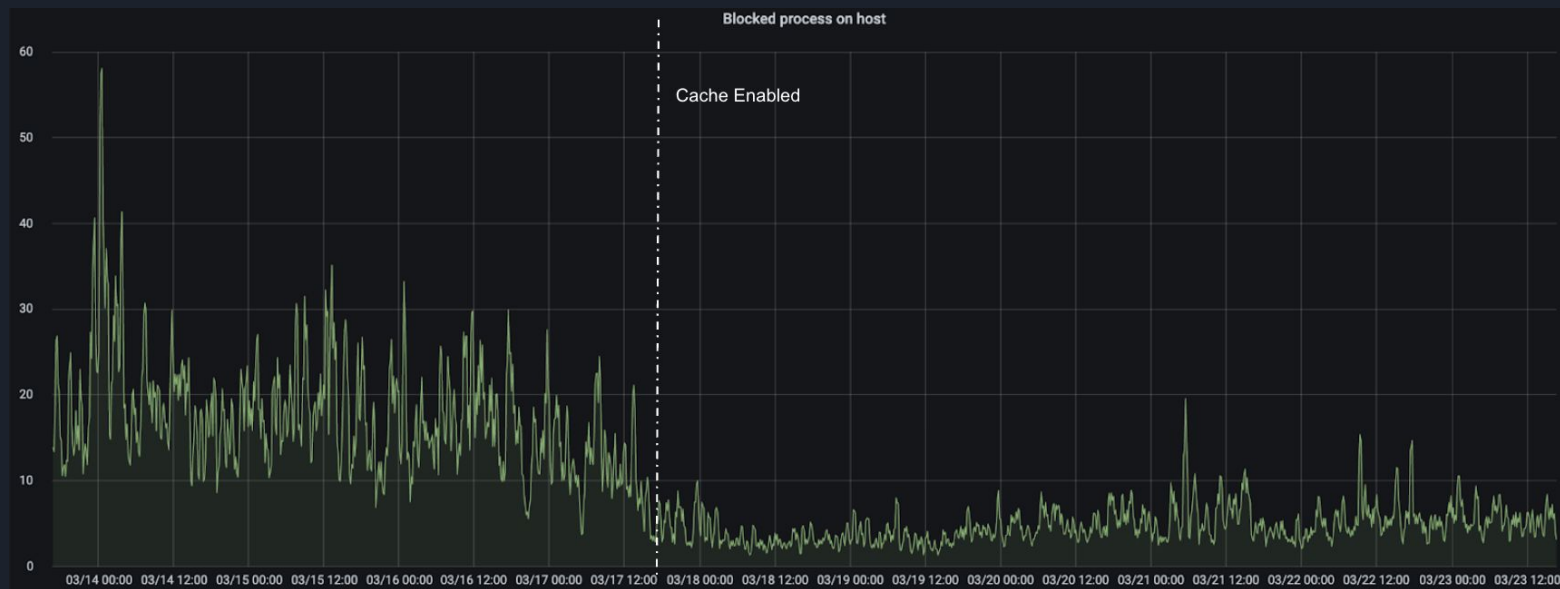
Evaluation and Rollout

- Performance evaluation: DataNode trace replay
 - Test with scale
 - Reproduce production traffic profile
- Alluxio cache “shadow mode”
 - A simulated cache mode to collect metrics

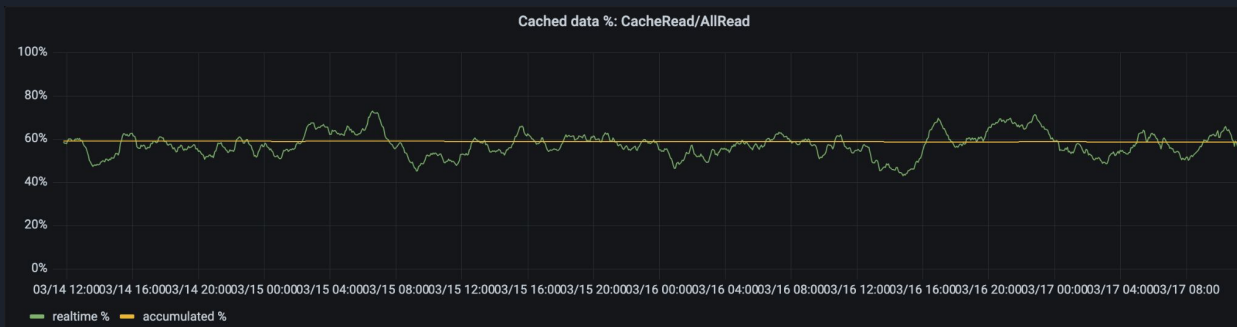


Current Status

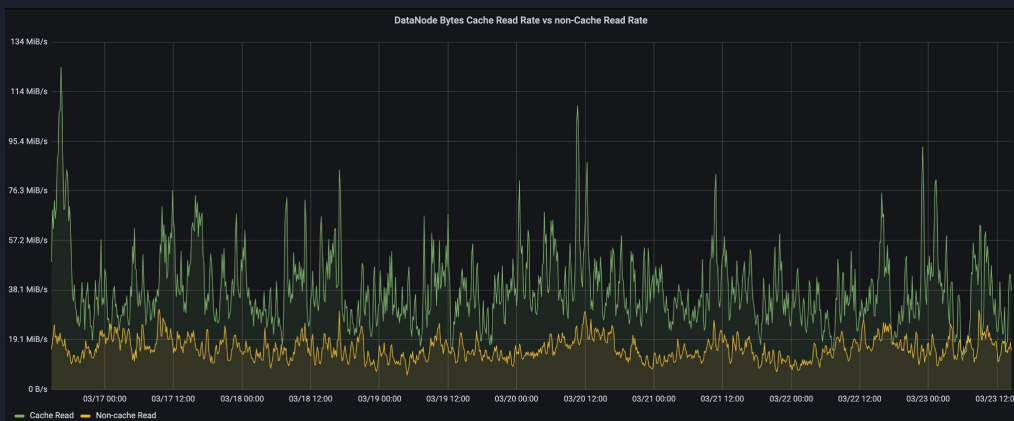
- Deployed in all major production clusters (1200 DataNodes)



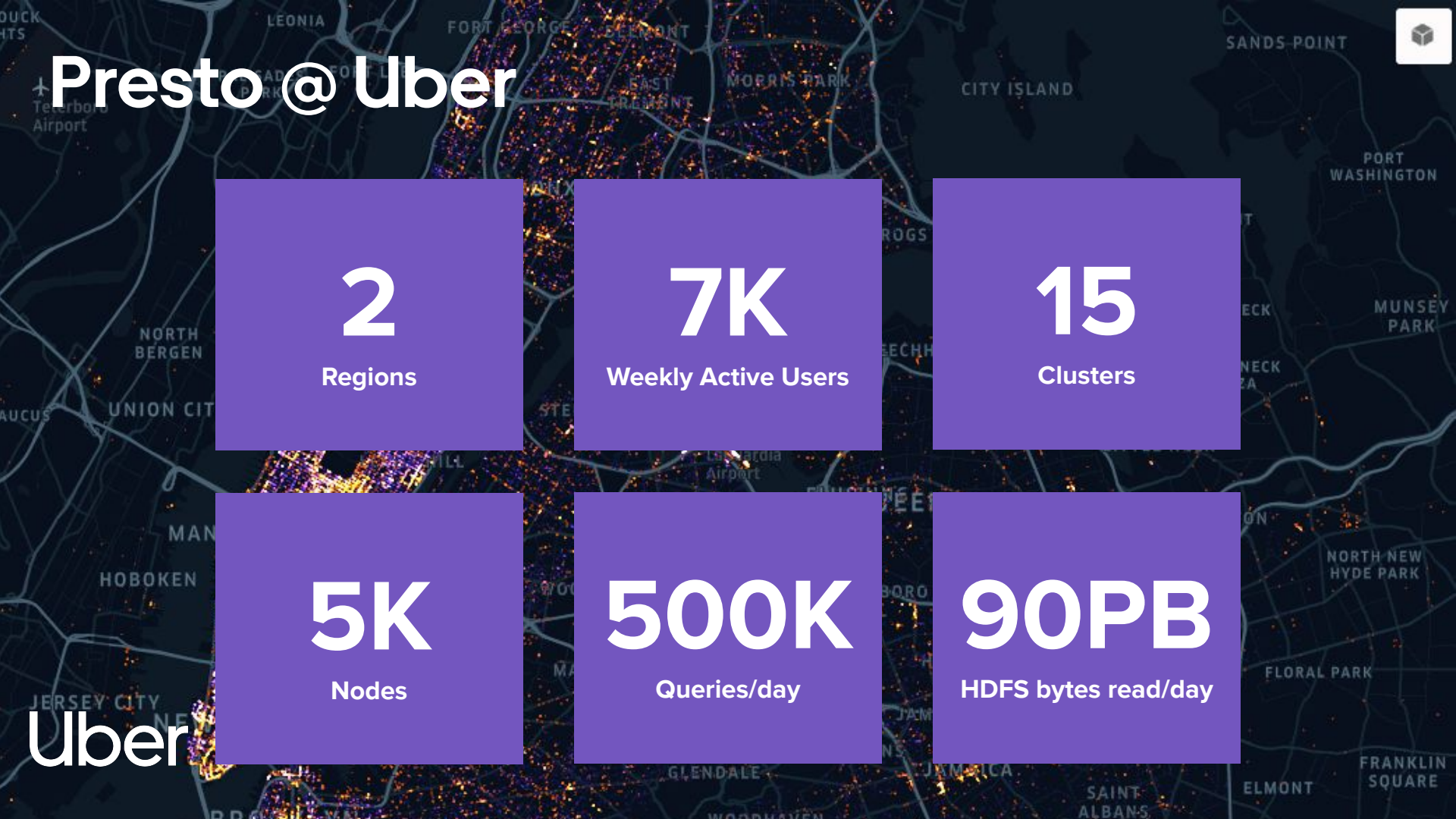
Current Status



Offload 60% of IO from HDD



Cache read throughput is significantly higher, nearly twice that of non-cache read throughput



Presto @ Uber

2

Regions

7K

Weekly Active Users

15

Clusters

5K

Nodes

500K

Queries/day

90PB

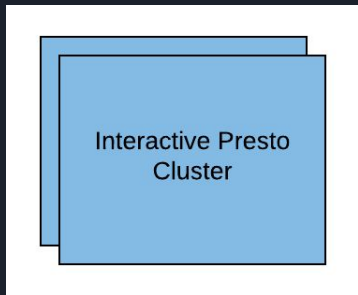
HDFS bytes read/day

Uber

Workloads

Interactive

Ad hoc queries



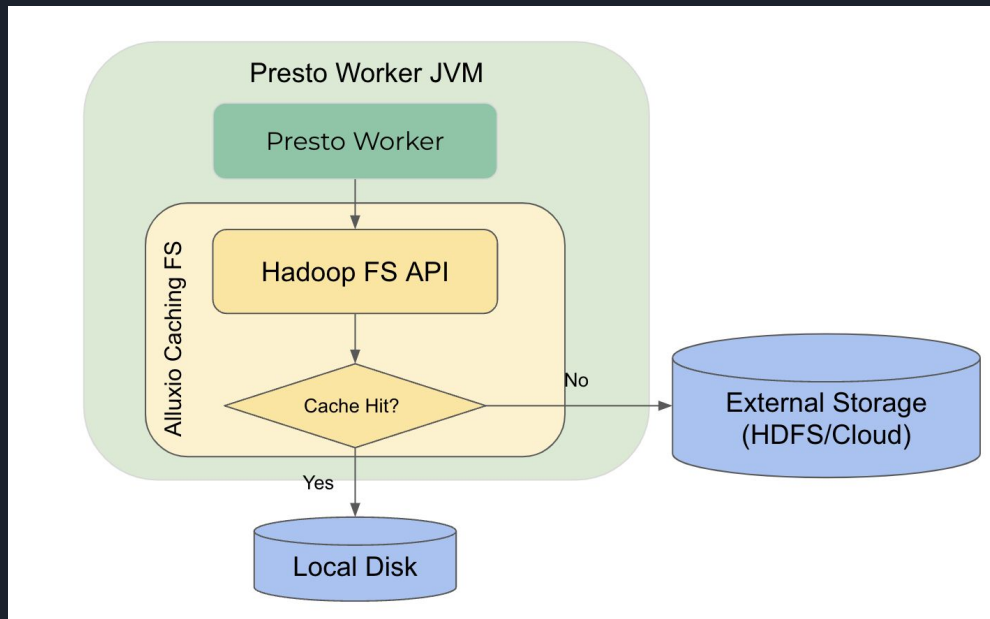
Batch

Scheduled



Presto Local Disk Cache

- Leverage Presto workers' local NVMe SSD disks
- Selective caching based on table
- Leverage Alluxio for cache management in Presto worker
- Forced cache TTL to meet compliance requirements





Key Challenges @ Uber and Solutions

- **Realtime Partition Updates**

- **Problem:** Tables/partitions are constantly changing, naive caching could cause returning stale data.
- **Solution:** Integrated HDFS file modification time as part of cache key

- **Cluster Membership Change**

- **Problem:** Presto nodes always leave and join cluster, causing file read to route to wrong node and thus cache miss.
- **Solution:** Introduced consistent hashing to tie files to certain Presto worker nodes.



Key Challenges @ Uber and Solutions

- **Cache Size Restriction**
 - **Problem:** Available cache space is limited compared to total data that needs to be scanned.
 - **Solution:** Introduced cache filter to selectively cache only certain hot data (top accessed tables).
- **Tiny Reads**
 - **Problem:** Uber's Presto traffic often involves a series of tiny consecutive reads, leading to degraded behavior in Alluxio.
 - **Solution:** Introduced buffer reads in Alluxio caching.



Presto Local SSD Cache Onprem - Today

- **Deployment**

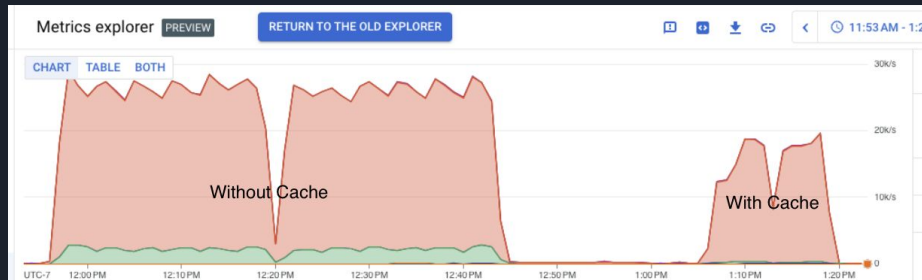
- Deployed to Presto production since 2022
- Clusters run with local cache
 - 5 clusters (out of 11 batch clusters)
 - 1500 nodes (43% of all batch nodes in primary region)

- **Improvement**

- ~**13%** Presto batch HDFS read offloaded to cache
- Input read latency reduced by ~**44%** compared to HDFS reads

Presto Local SSD Cache On Cloud - Initial Evaluation

- Deployed to Presto on cloud for initial evaluation
- Cost reduction AND performance improvement
 - **>80% reduction** of # of read requests to GCS
 - **228 s → 50 s reduction** of P90 query latency



THANK YOU

Jing Zhao, Uber

Hope Wang, Alluxio

