

Blizzard: Fast, Cloud-scale Block Storage for Cloud-oblivious Applications

Microsoft®
Research



James Mickens, Jeremy Elson,
Edmund B. Nightingale,
Darren Gehring, Krishna
Nareddy



THE UNIVERSITY
of
WISCONSIN
MADISON



Asim Kadav
Vijay Chidambaram

**Carnegie
Mellon
University**



Bin Fan



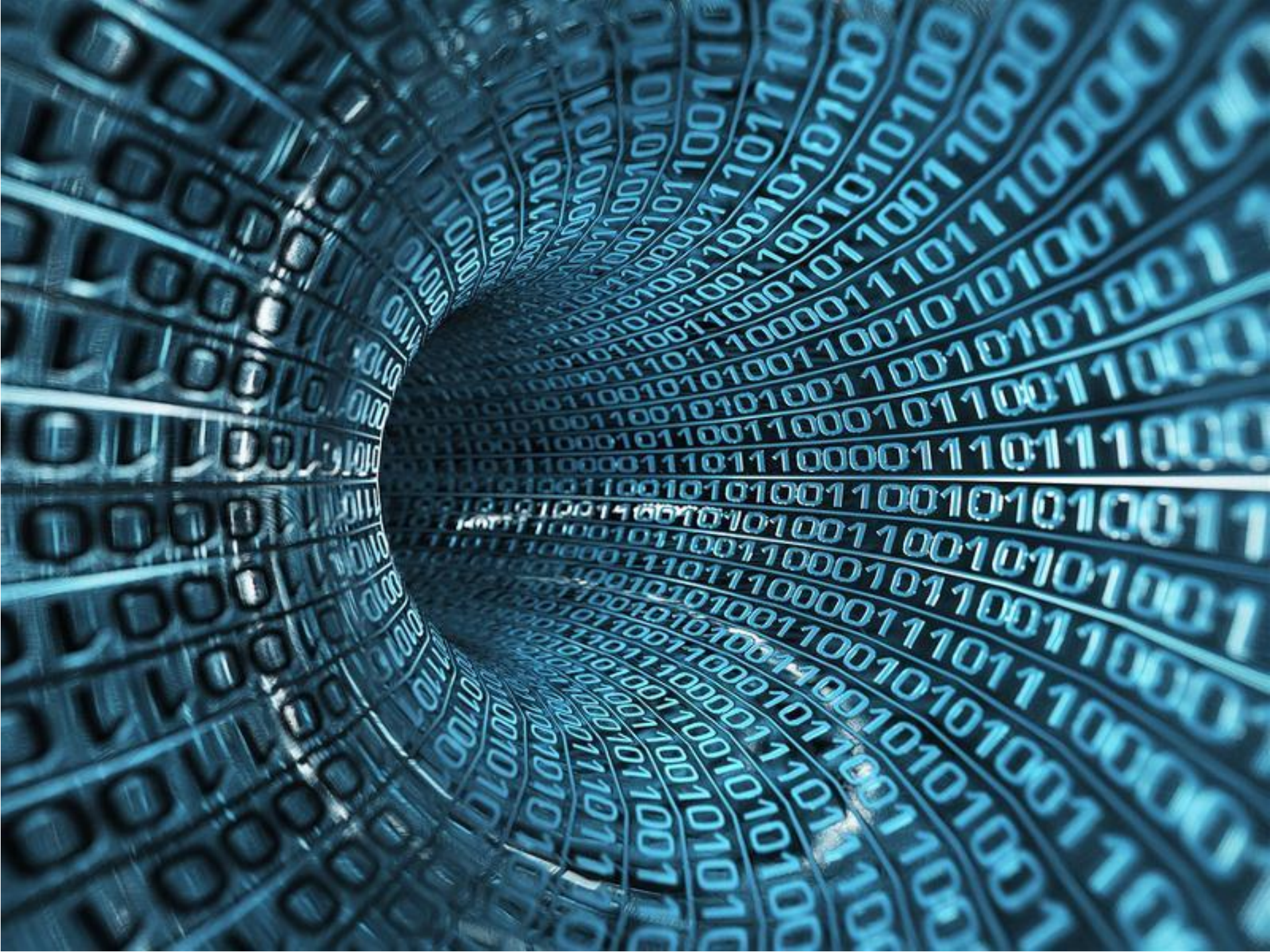
JOHNS HOPKINS



Osama Khan



IT'S, UH, THE FUTURE







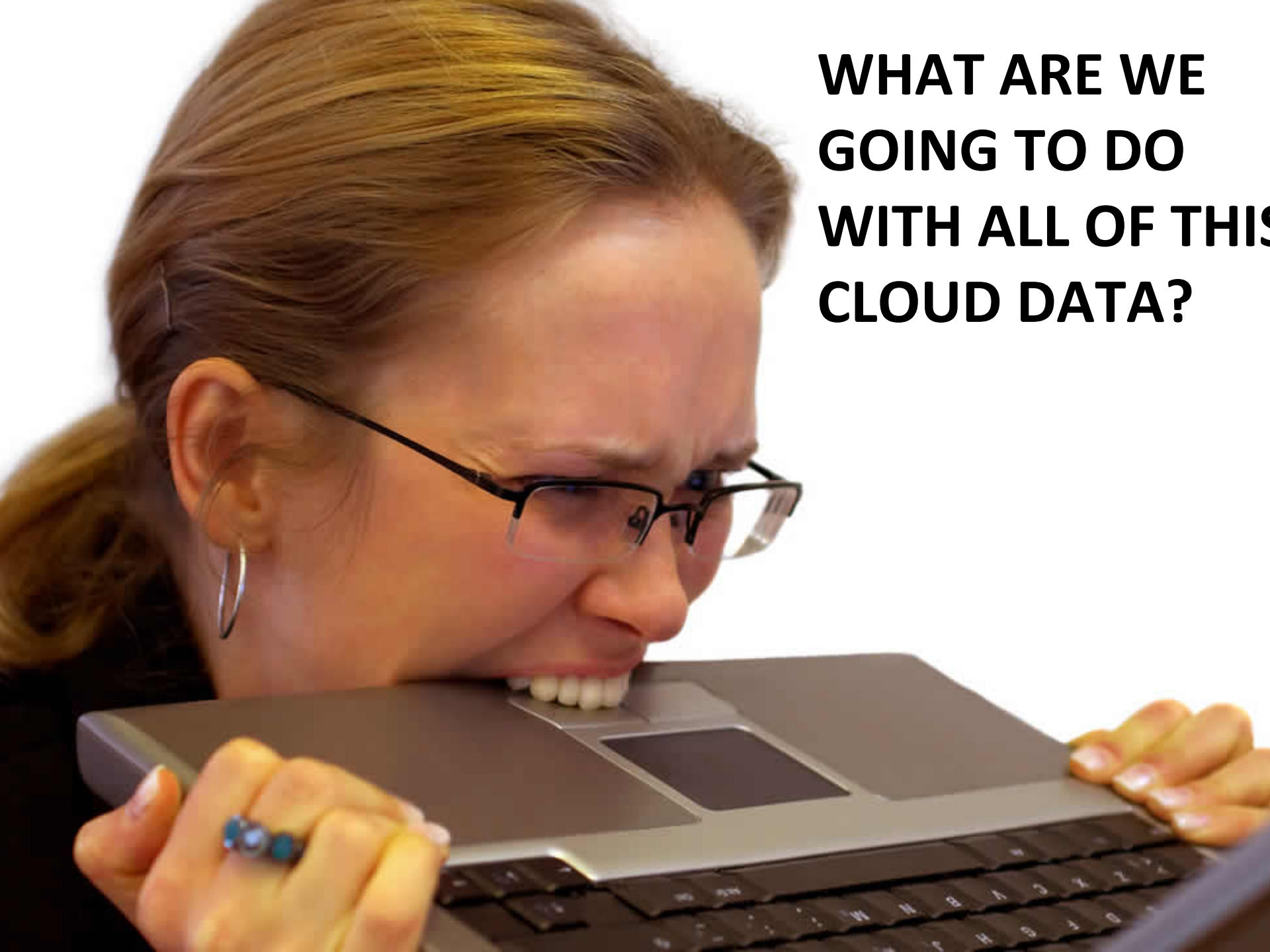
BIG

DATA



real-time DATA

**WHAT ARE WE
GOING TO DO
WITH ALL OF THIS
CLOUD DATA?**





**CLOUD
DATA**



OCEAN

My Goal

- Take unmodified POSIX/Win32 applications . . .
- Run those applications in the cloud . . .
- On the same hardware used to run big-data apps . . .
- . . . and give them cloud-scale IO performance!

PostgreSQL



Microsoft®
Exchange 2013



My Goal

- Take unmodified POSIX/Win32 applications . . .
- Run those applications in the cloud . . .
- On the same hardware used to run big-data apps . . .
- . . . and give them cloud-scale IO performance!

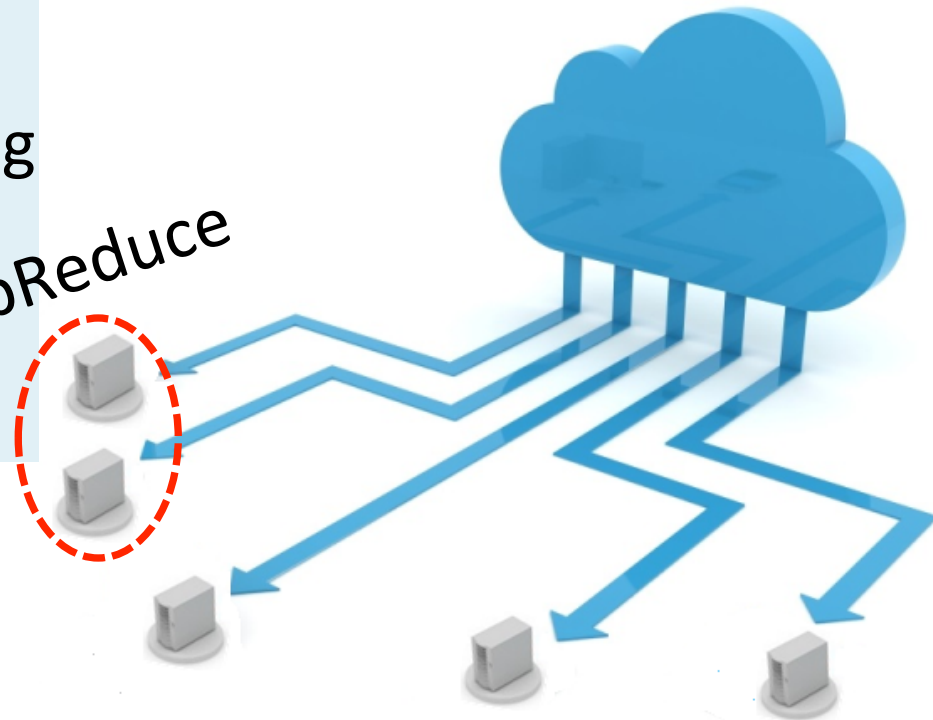
PostgreSQL throughput > 1000 MB/s



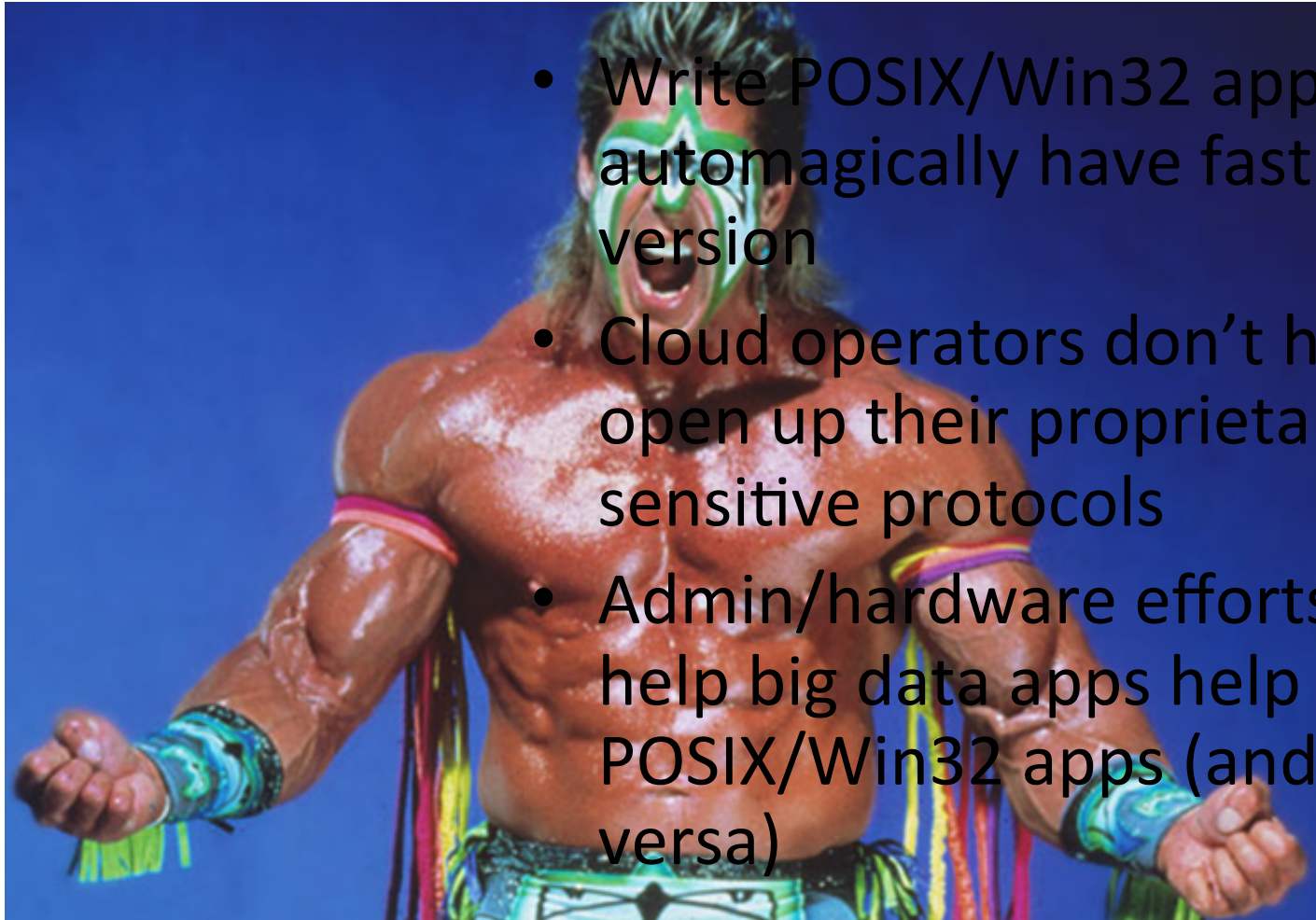
Scale-out architecture using commodity parts

Transparent failure recovery

MapReduce



Why Do I Want To Do This?



- Write POSIX/Win32 app once, automagically have fast cloud version
- Cloud operators don't have to open up their proprietary or sensitive protocols
- Admin/hardware efforts that help big data apps help POSIX/Win32 apps (and vice versa)

BECAUSE THIS WOULD BE AMAZING

Our Solution: Blizzard

PostgreSQL



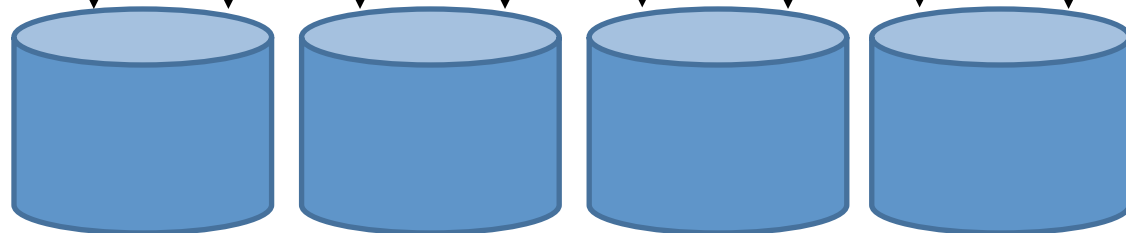
Microsoft®
Exchange 2013



Blizzard
virtual drive



Remote disks



Our Solution: Blizzard



Blizzard
virtual drive

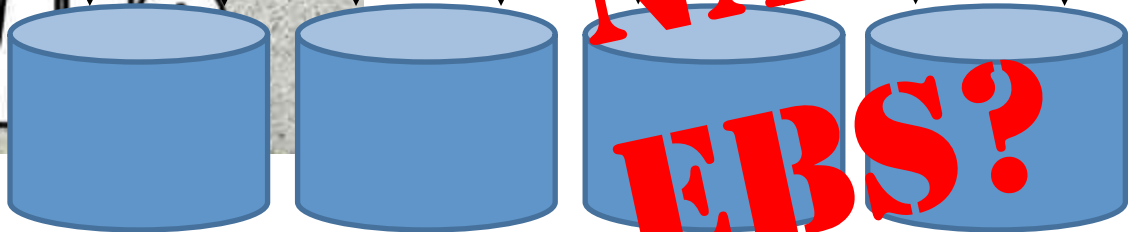


RAID?

NAS?

EBS?

Remote disks





**THE PROBLEM
IS YOU**

LISTEN



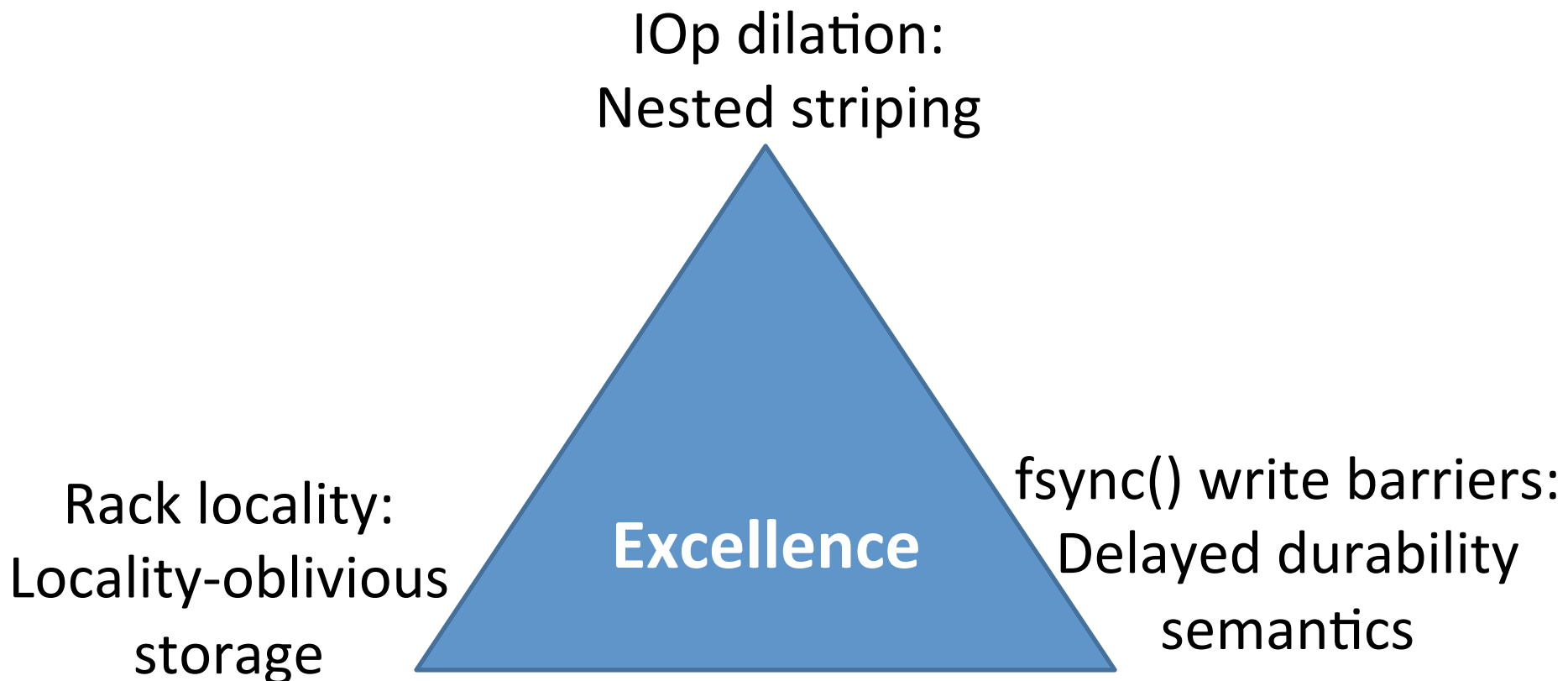
The naïve approach for implementing virtual disks does not **maximize spindle parallelism** for POSIX/Win32 applications which **frequently issue fsync() operations** to maintain consistency.

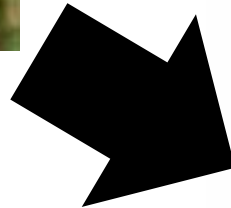
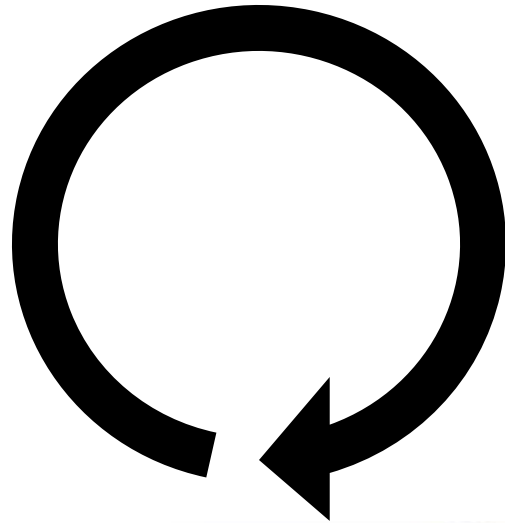
LISTEN



The naïve approach for implementing virtual disks does not **maximize spindle parallelism** for POSIX/Win32 applications which **frequently issue fsync() operations** to maintain consistency.

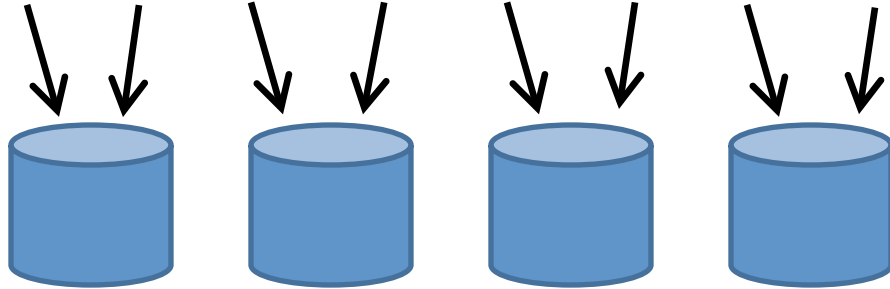
LISTEN







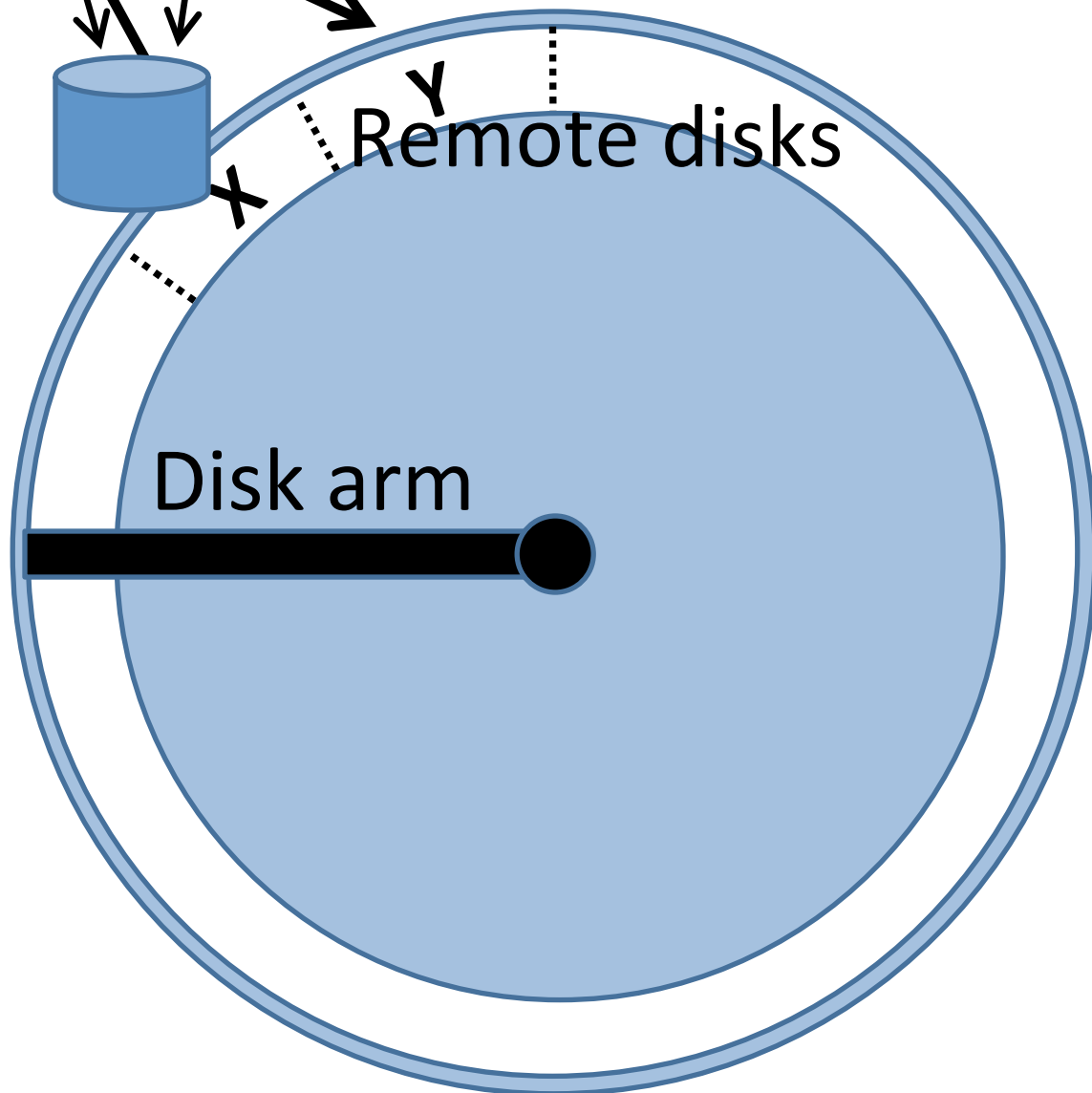
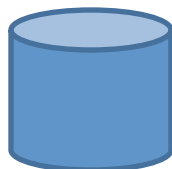
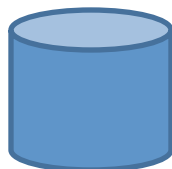
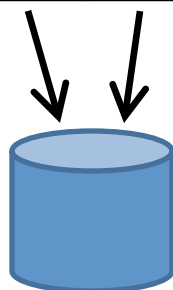
Virtual disk



Remote disks

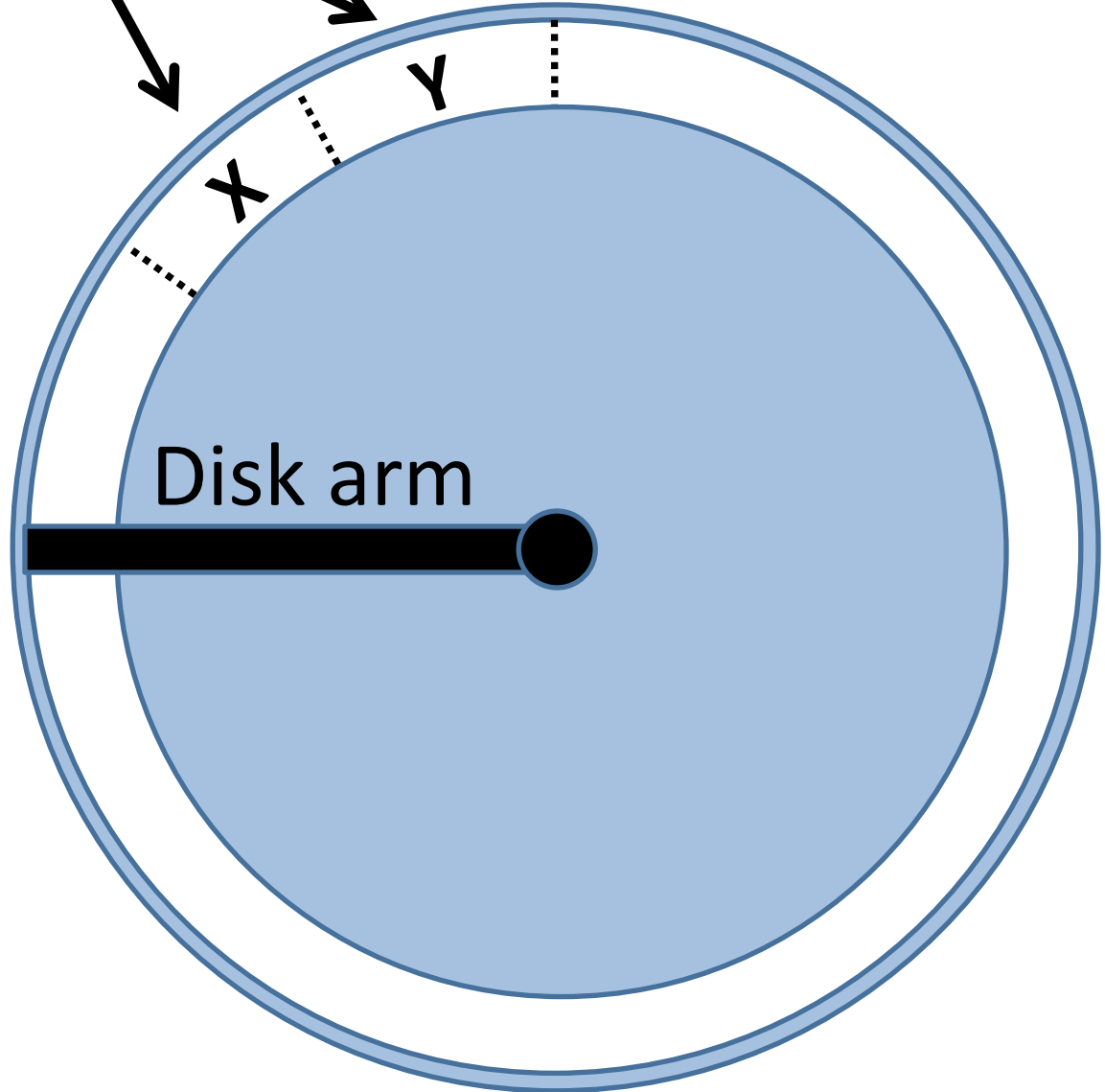
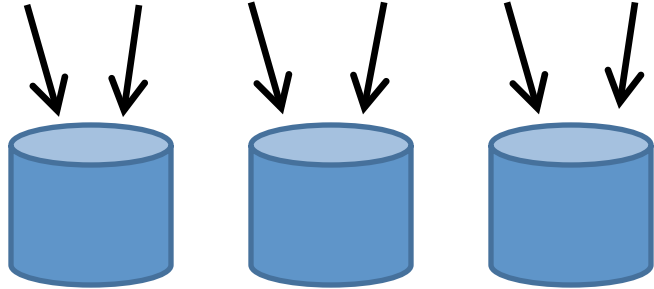
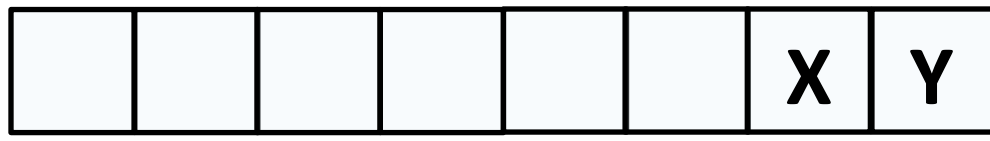


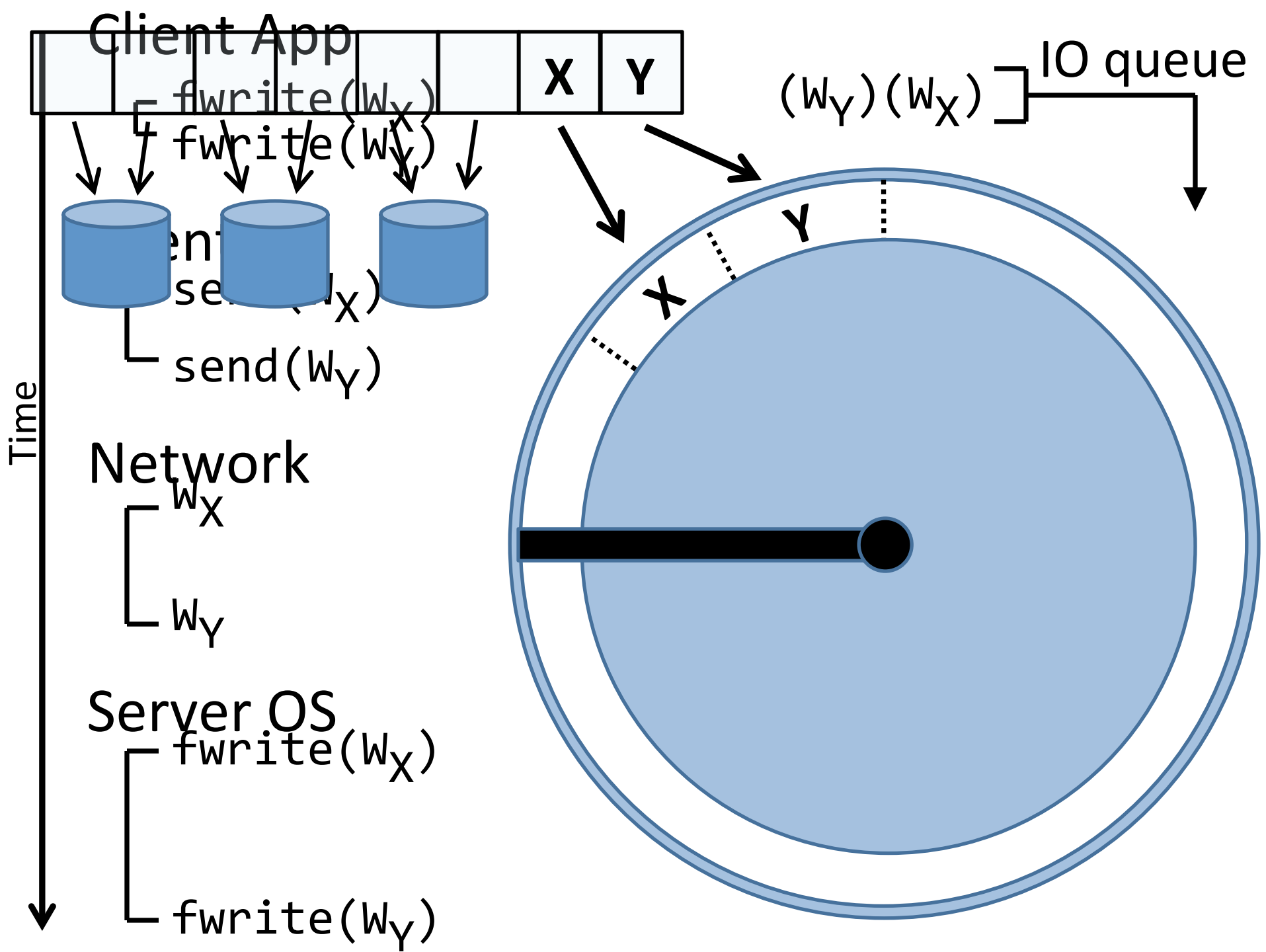
Virtual disk

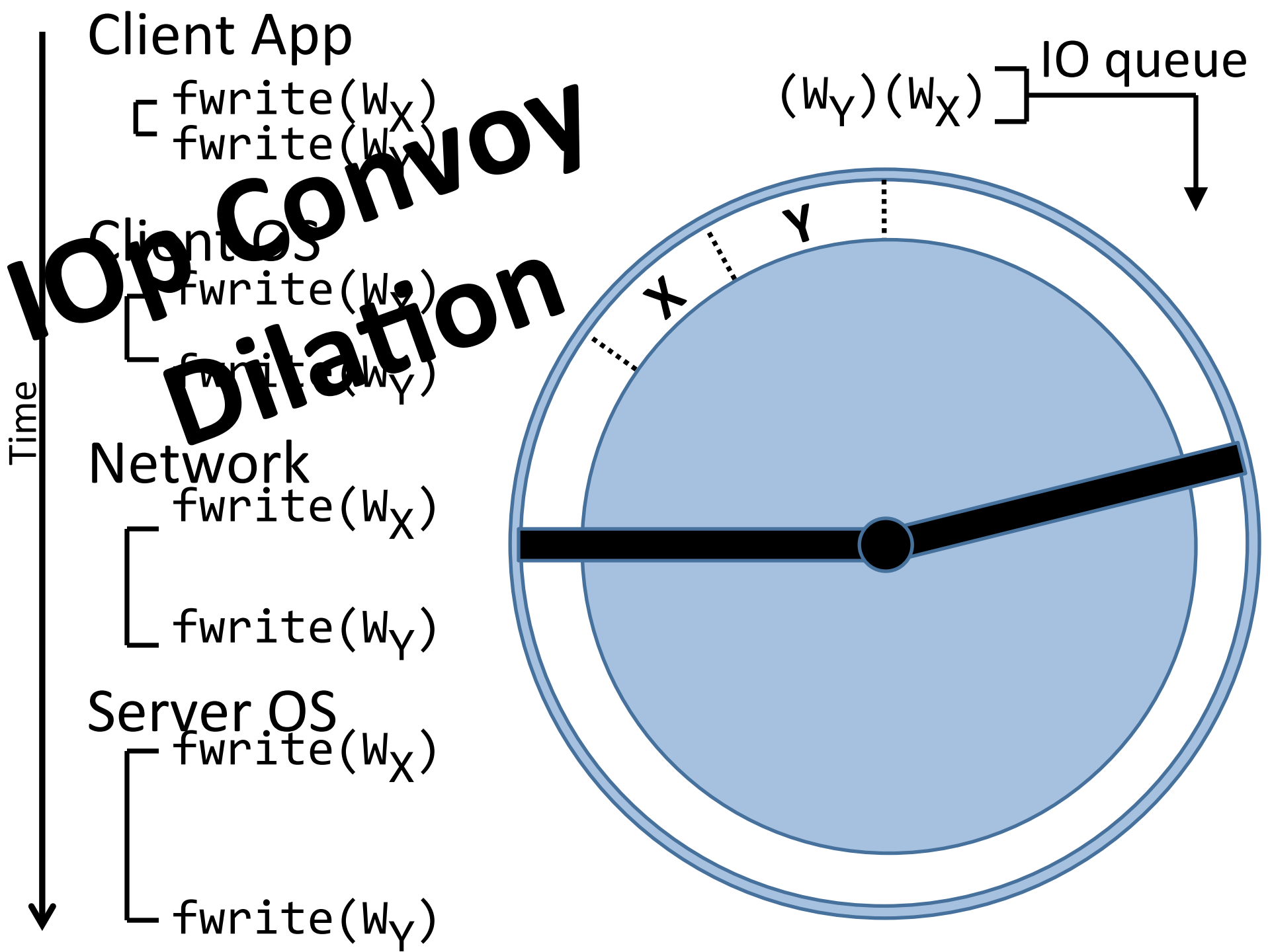


Remote disks

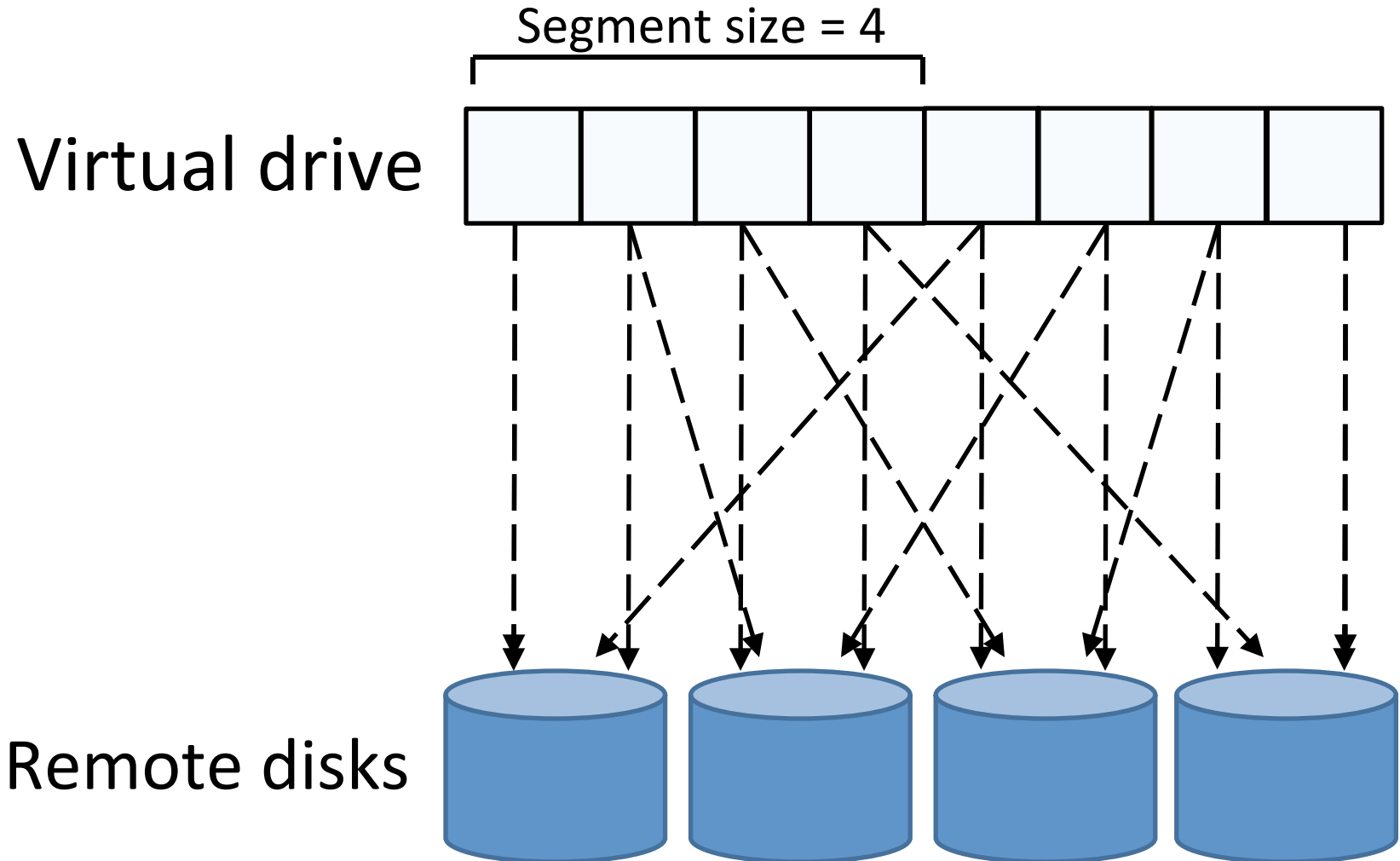
Disk arm



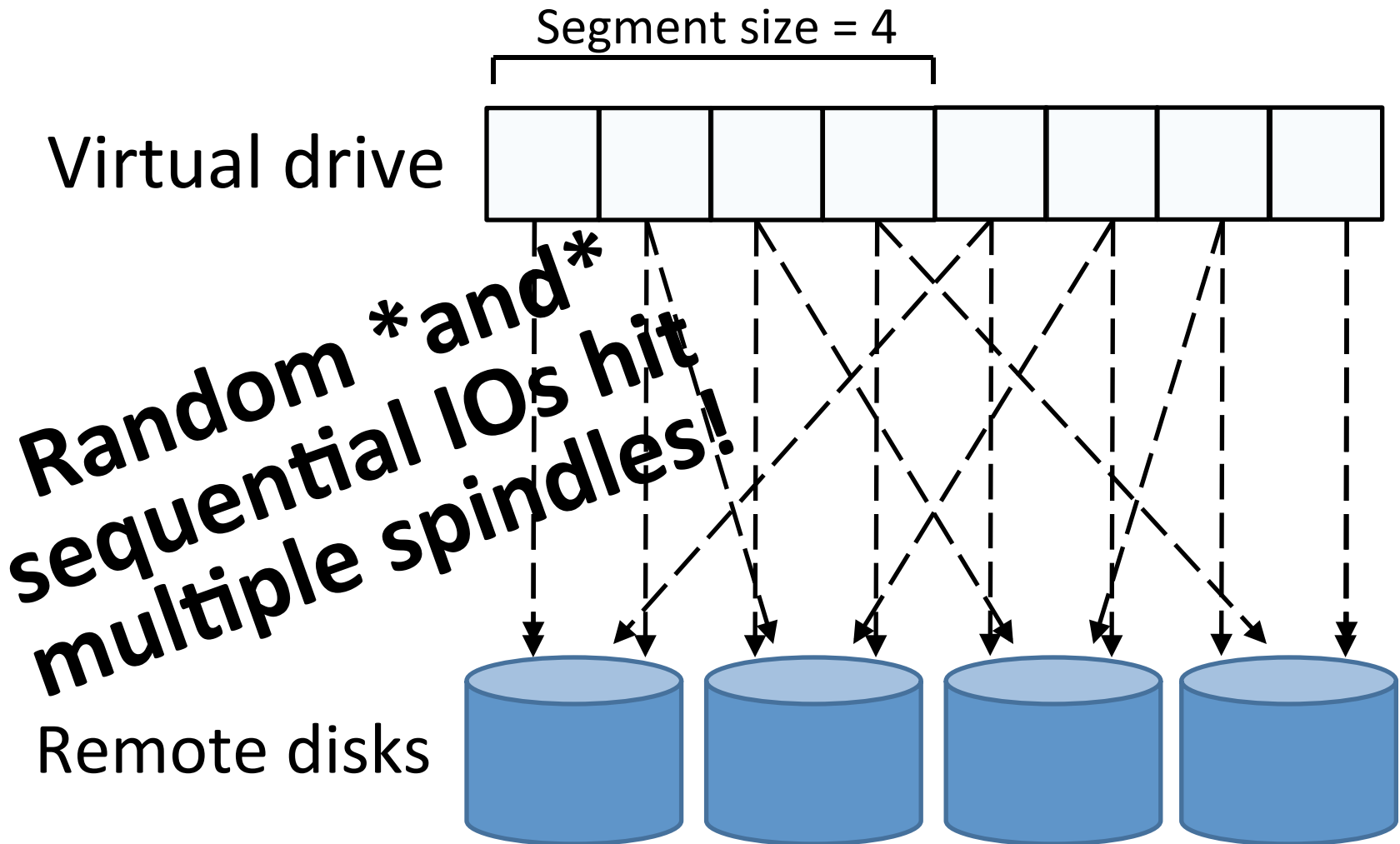




Fixing IOp Convoy Dilation



Fixing IOp Convoy Dilation



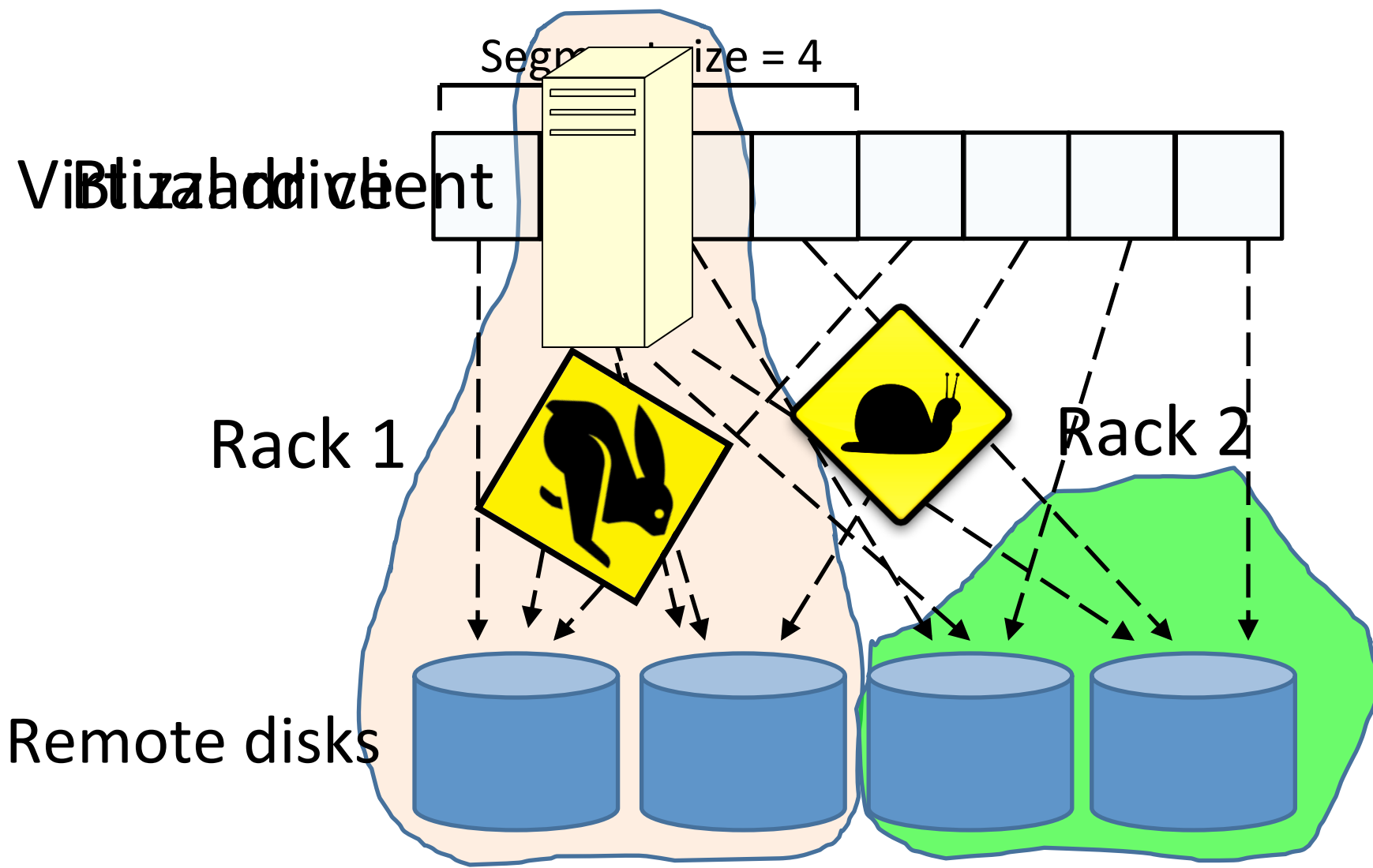


**TOTAL
VICTORY**

Rack Locality



Rack Locality



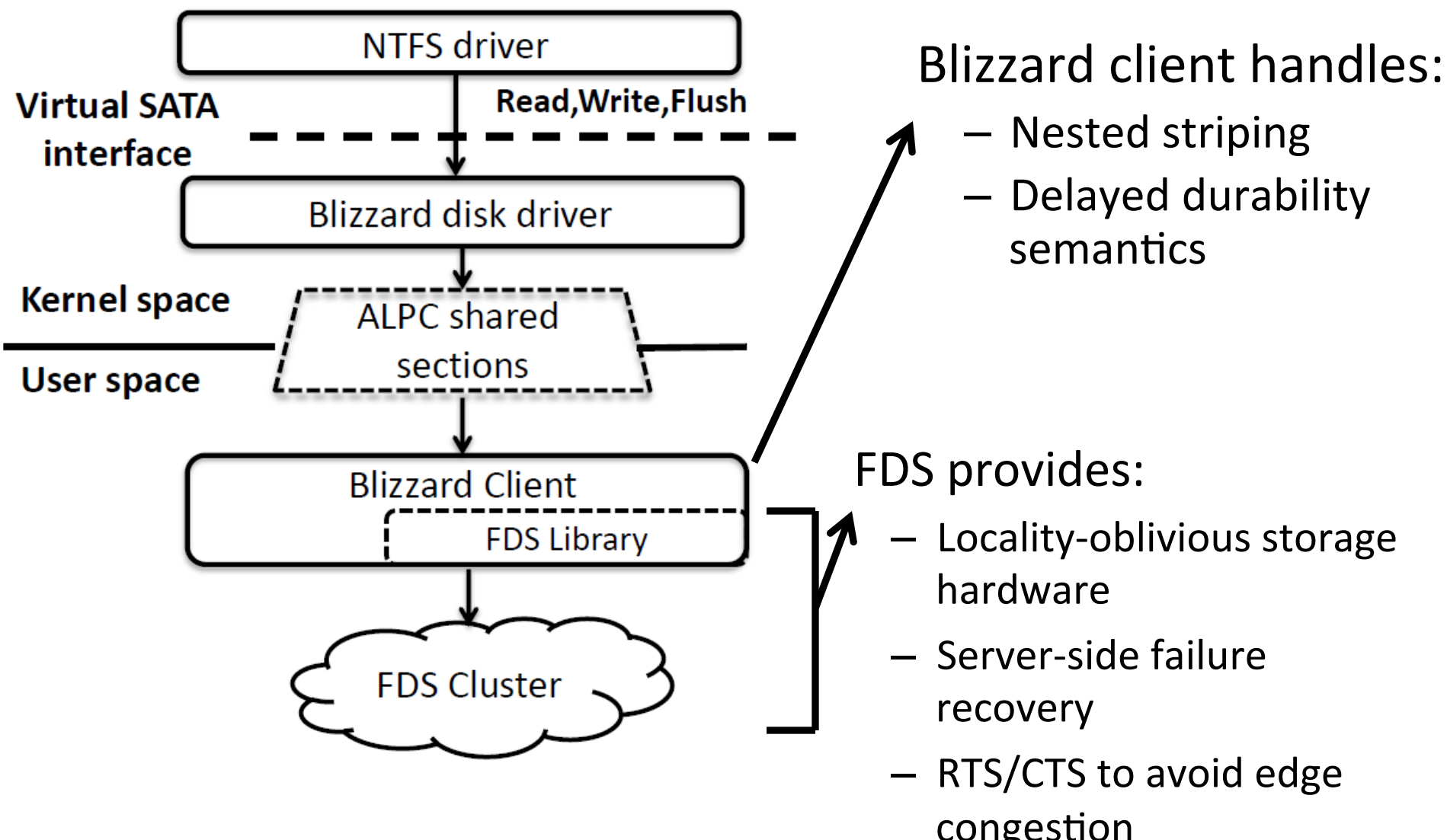


**AIN'T NOBODY
GOT TIME FOR
THAT**

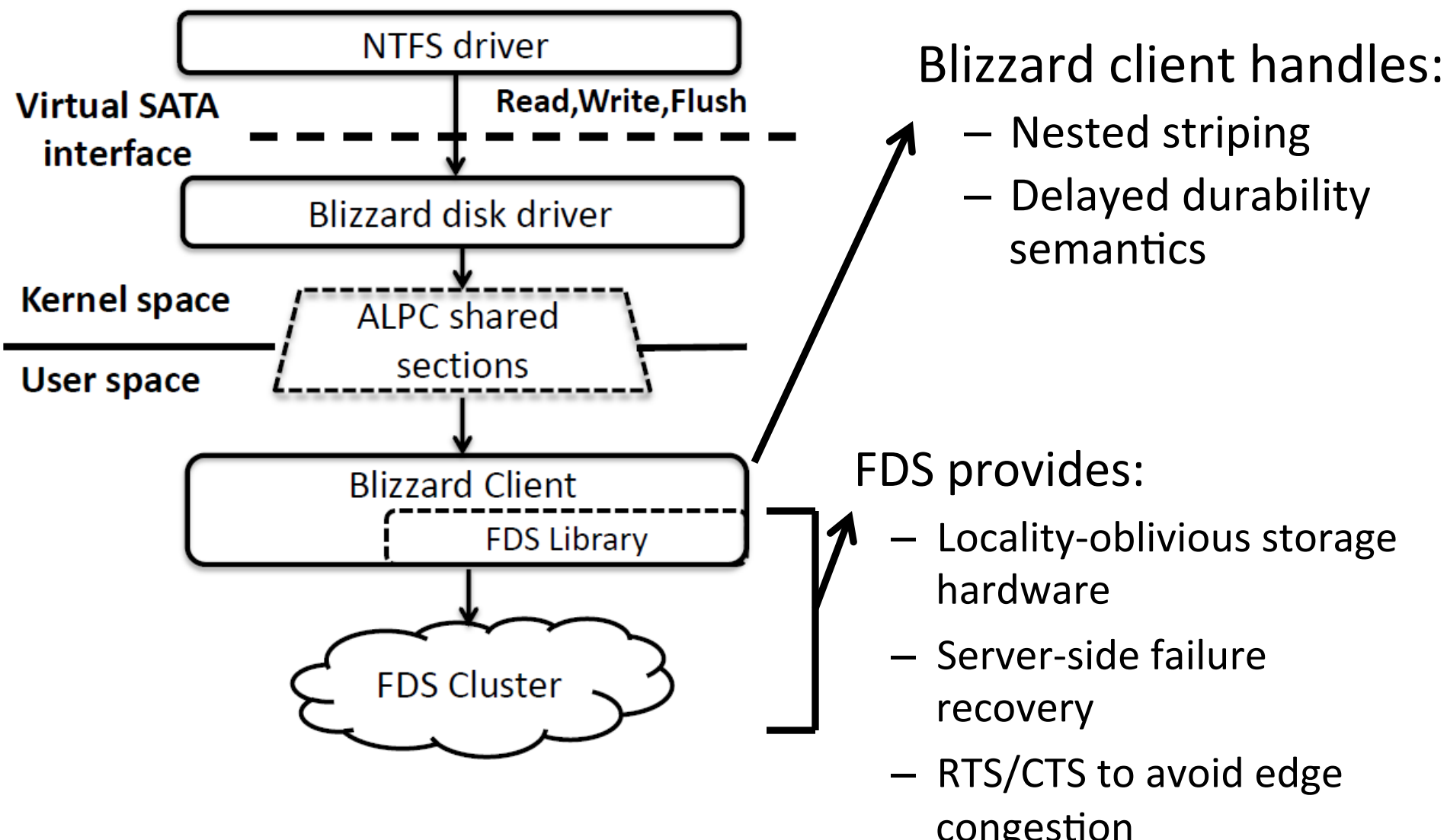
FDS To The Rescue (OSDI 2012)

- Hardware architecture
 - Full bisection bandwidth network (no oversubscription)
 - Allocate each disk enough network bandwidth to drive disk at full sequential speed (1 disk $\approx$ 128 MB/s $\approx$ 1Gbps)
- Result: locality-oblivious storage
 - Any client can access any disk as fast as local
 - **Enables aggressive striping**

Blizzard as FDS Client



Blizzard as FDS Client





**POSIX/Win32
apps**

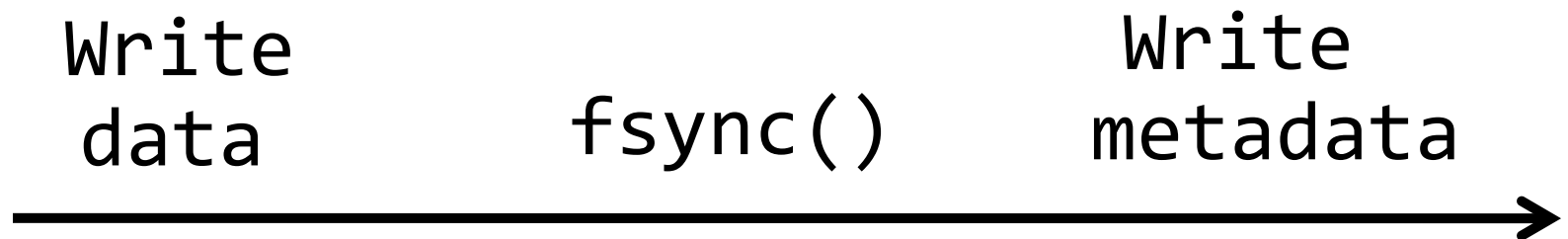
**Traditional
big-data apps**



ARE WE THERE YET?

The problem with `fsync()`

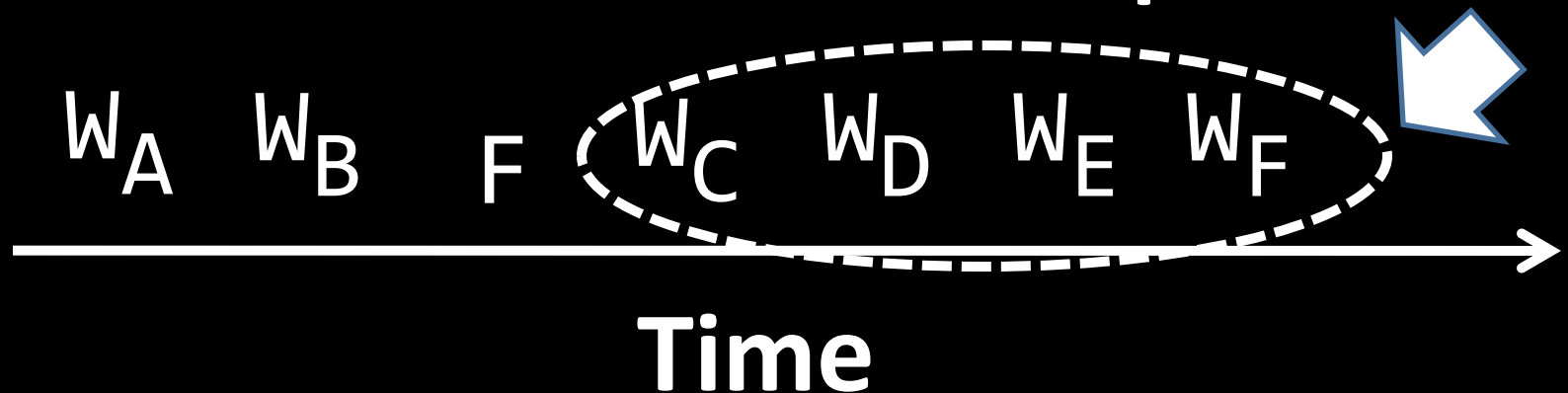
- Used by POSIX/Win32 file systems and applications to implement crash consistency
 - Disk only returns from `fsync()` when all prior writes have become durable
 - Ex: ensure data is written before metadata





WRITE BARRIERS RUIN BIRTHDAYS

Stalled operations
limit parallelism!

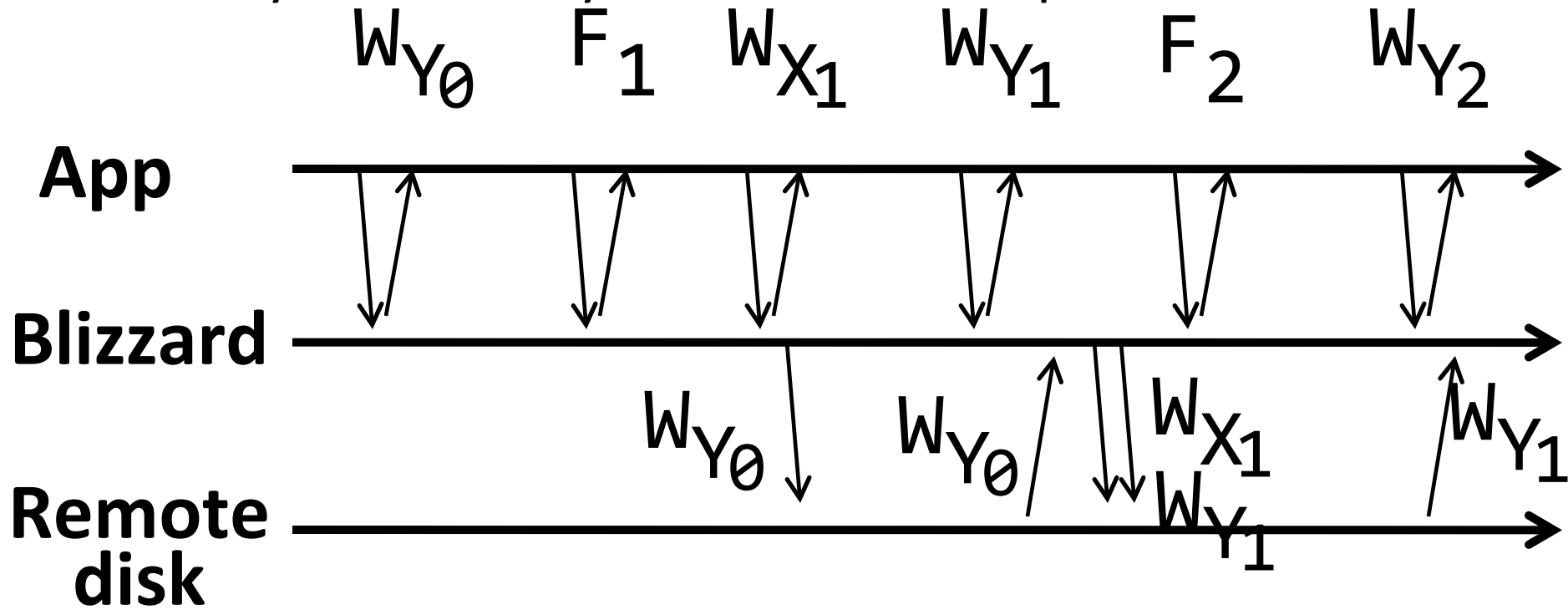


Delayed Durability

- Decouple durability from ordering
- Acknowledge `fsync()` immediately . . .
 - . . . but increment flush epoch
 - Tag writes with their epoch number,
asynchronously retire writes in epoch order

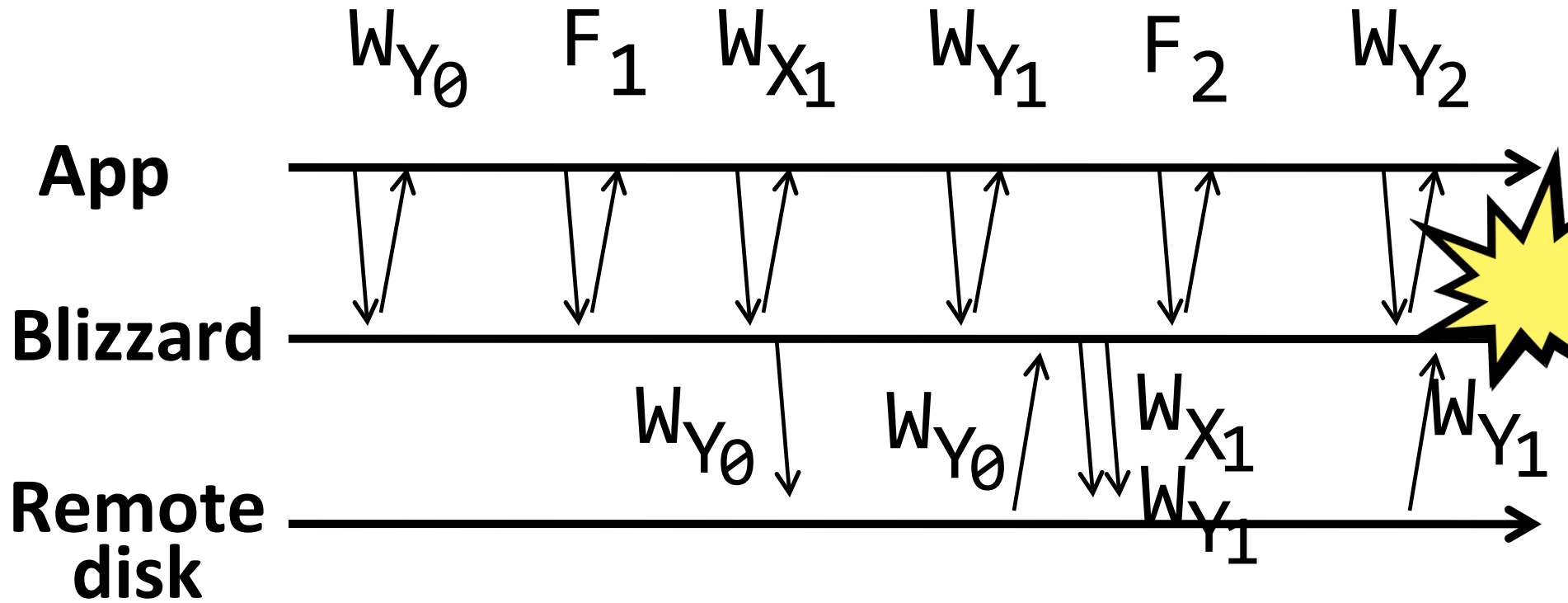
Delayed Durability

- Decouple durability from ordering
- Acknowledge `fsync()` immediately ...
 - ... but increment flush epoch
 - Tag writes with their epoch number, asynchronously retire writes in epoch order

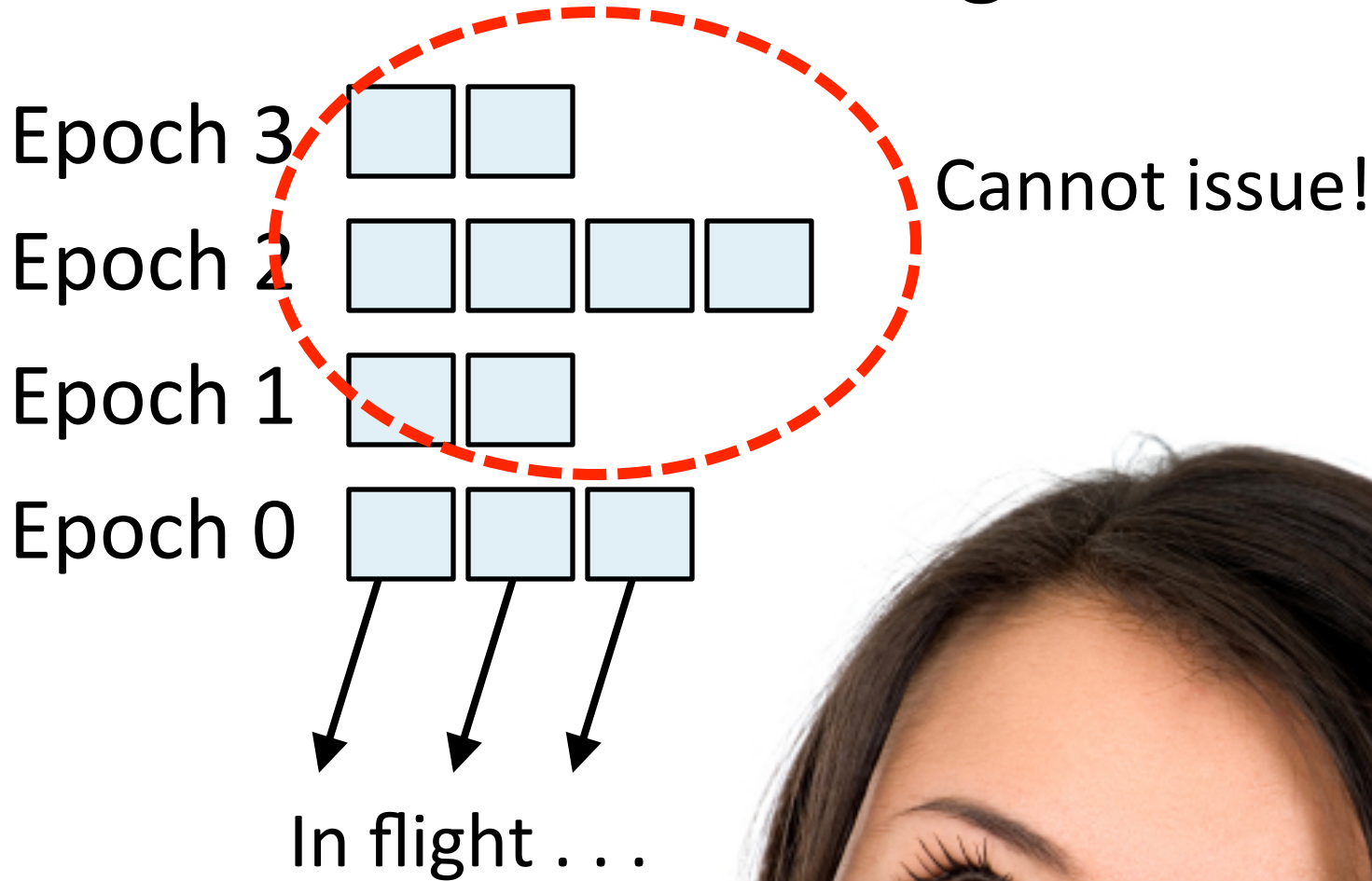


Delayed Durability

- All writes are acknowledged
- ... but only W_0 and W_1 are durable!
- Satisfies prefix consistency ordering
 - = All epochs up to $N-1$ are durable
 - = Acknowledge $f_{sync}()$ immediately ...
 - Some, all, or no writes from epoch N are durable
 - but tag writes with their epoch number
 - = No writes from later epochs are durable
- Prefix consistency and asynchronously retire writes in epoch N provides much better performance!



Isn't Blizzard buffering a lot of data?

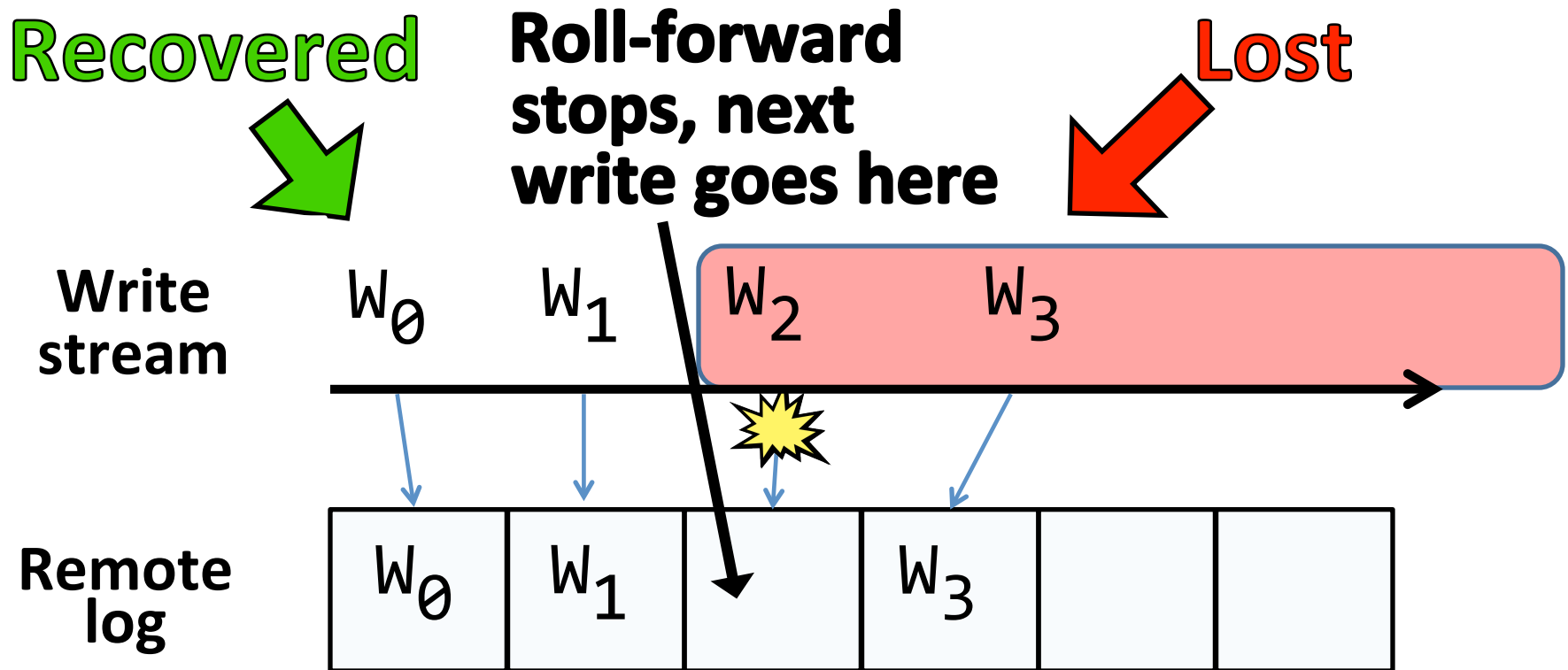


THERE ARE NO
SOLID-BASED
NEW CHEMICALS
EXTRACTION?
TECHNIQUES

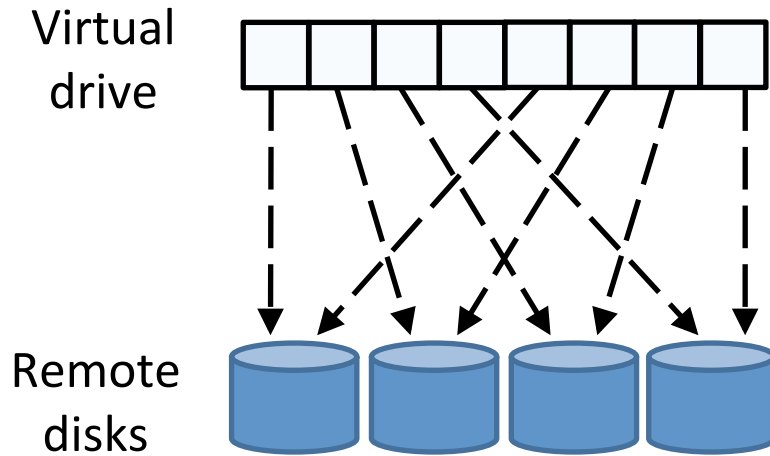


Log-based Writes

- Treat backing storage as a distributed log
 - Issue writes to log **immediately** and **in order**
 - On failure, roll forward from last checkpoint and stop when you find torn write, unallocated log block with old epoch number



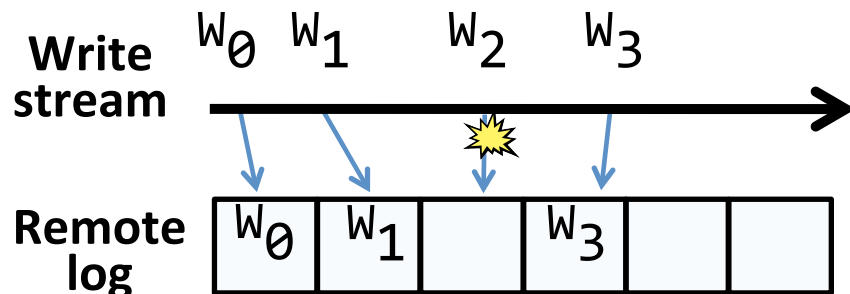
Summary of Blizzard's Design



FDS

- Problem: I/O Dilation
- Solution: Nested striping

- Problem: Rack locality constrains parallelism
- Solution: Full-bisection networks, match disk and network bandwidth



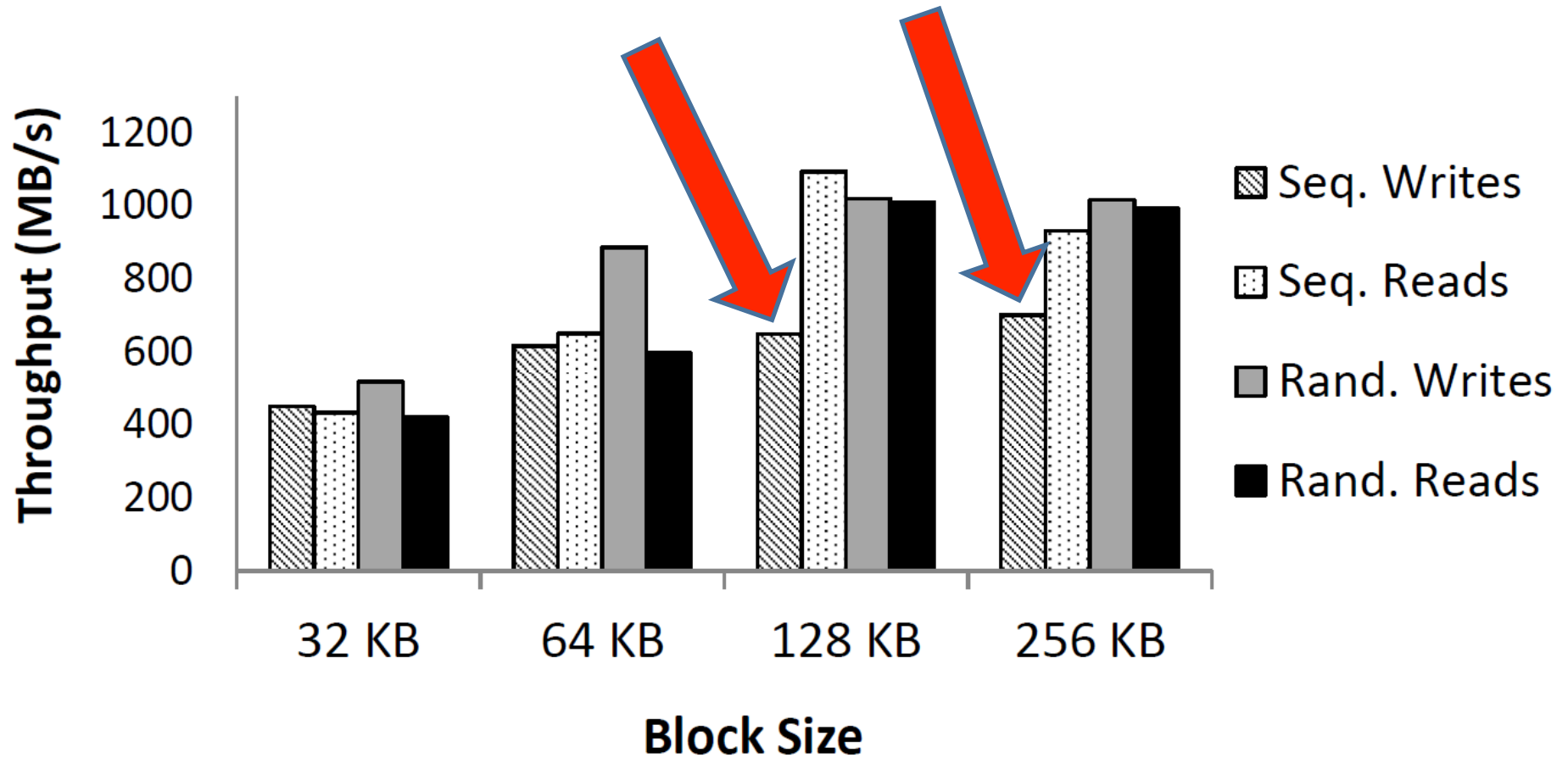
- Problem: Evil fsync()s
- Solution: Delayed durability

PROOF OF YOUR EXCELLENCE

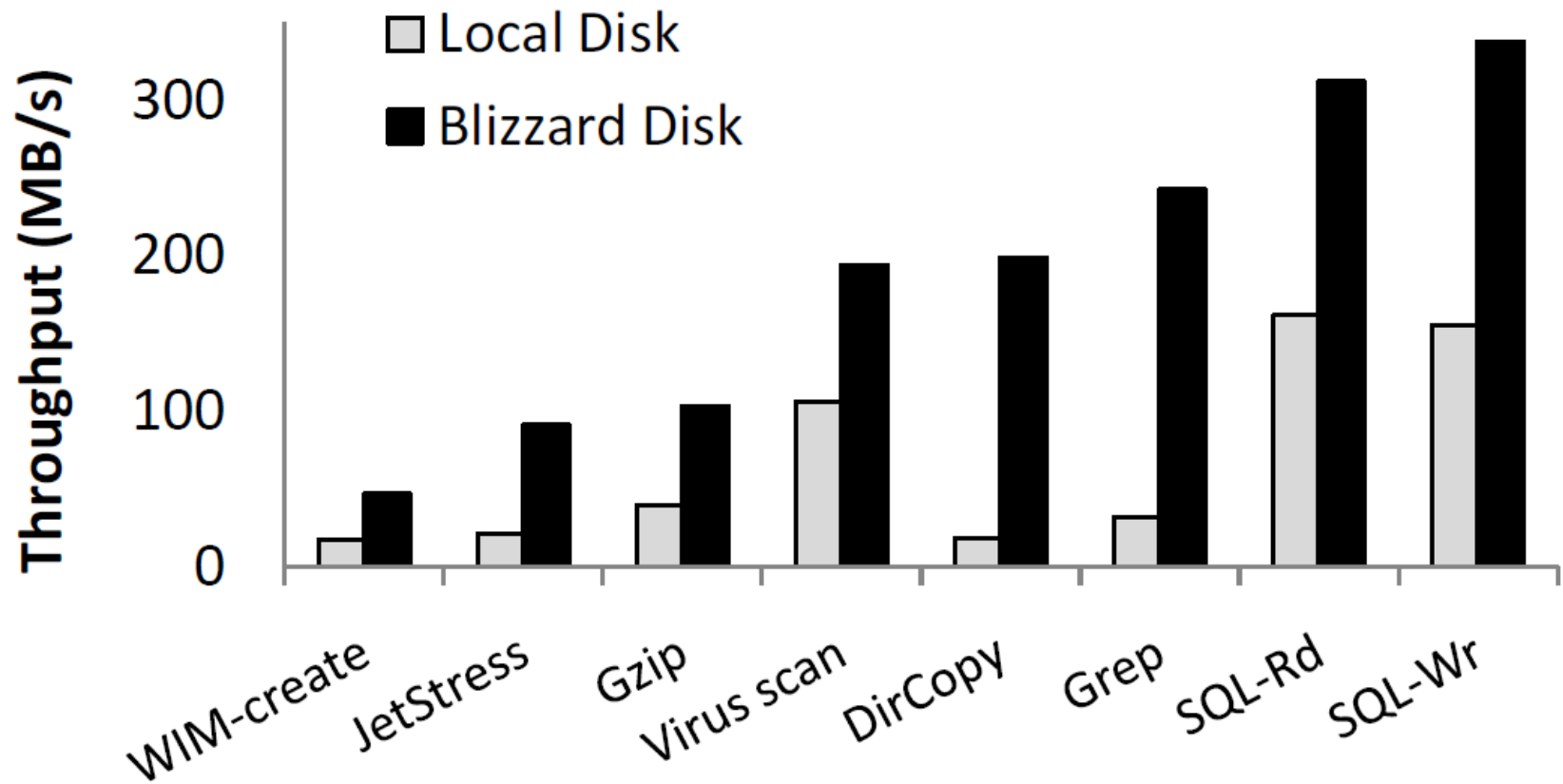


I DO NOT SEE IT

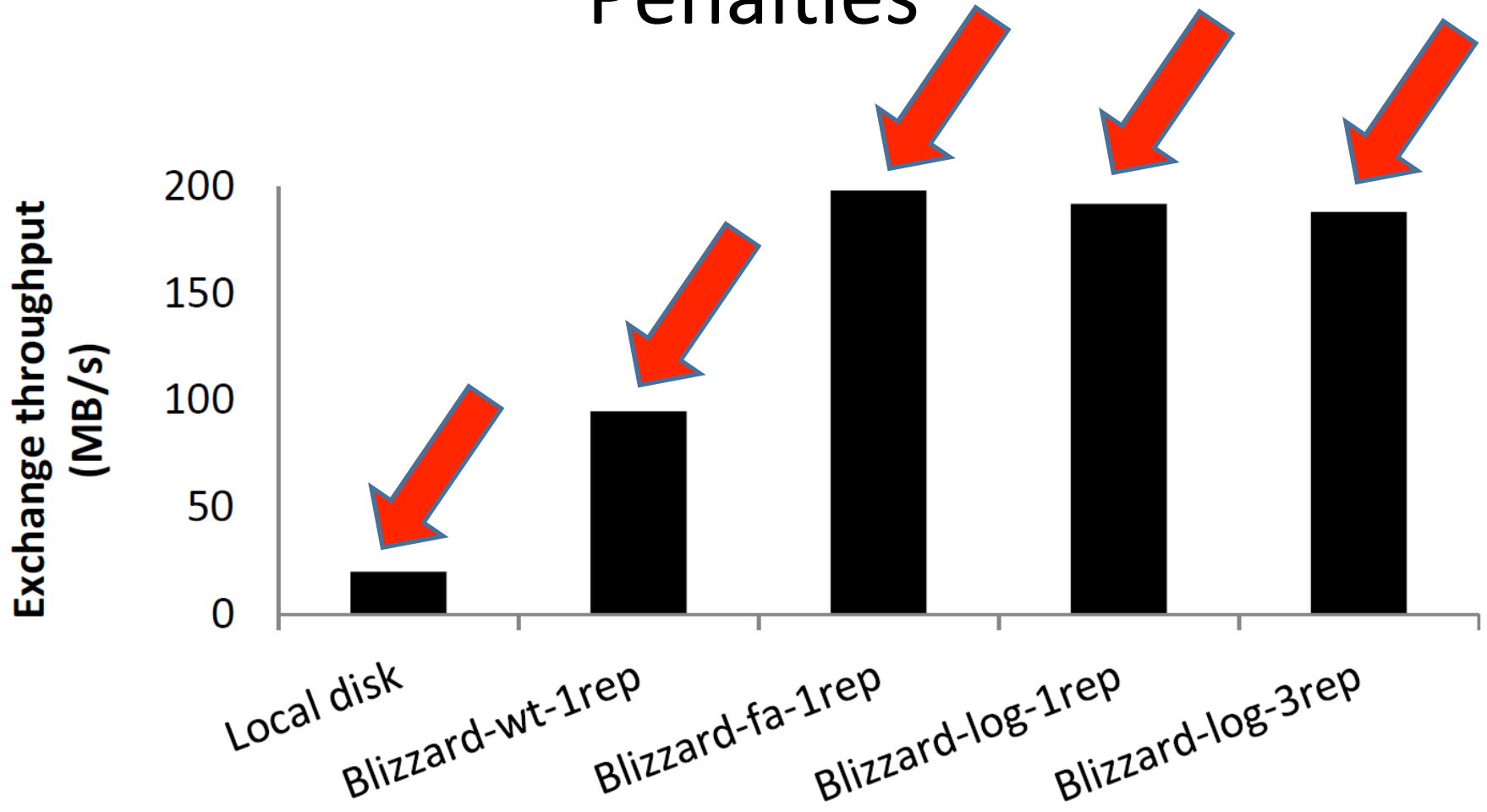
Throughput Microbenchmark



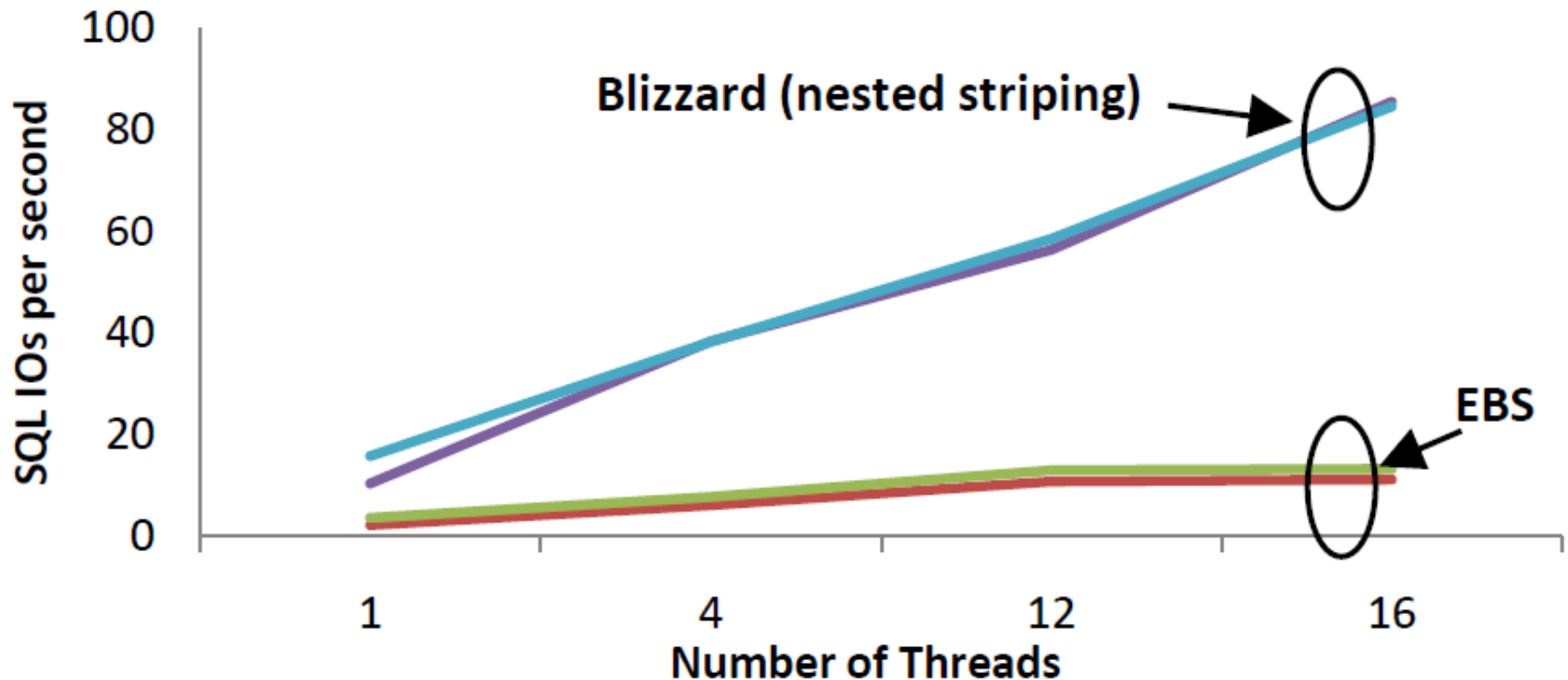
Application Macrobenchmarks



Delayed Durability: Hiding Replication Penalties



Blizzard vs EBS: Write IOps and Read IOps



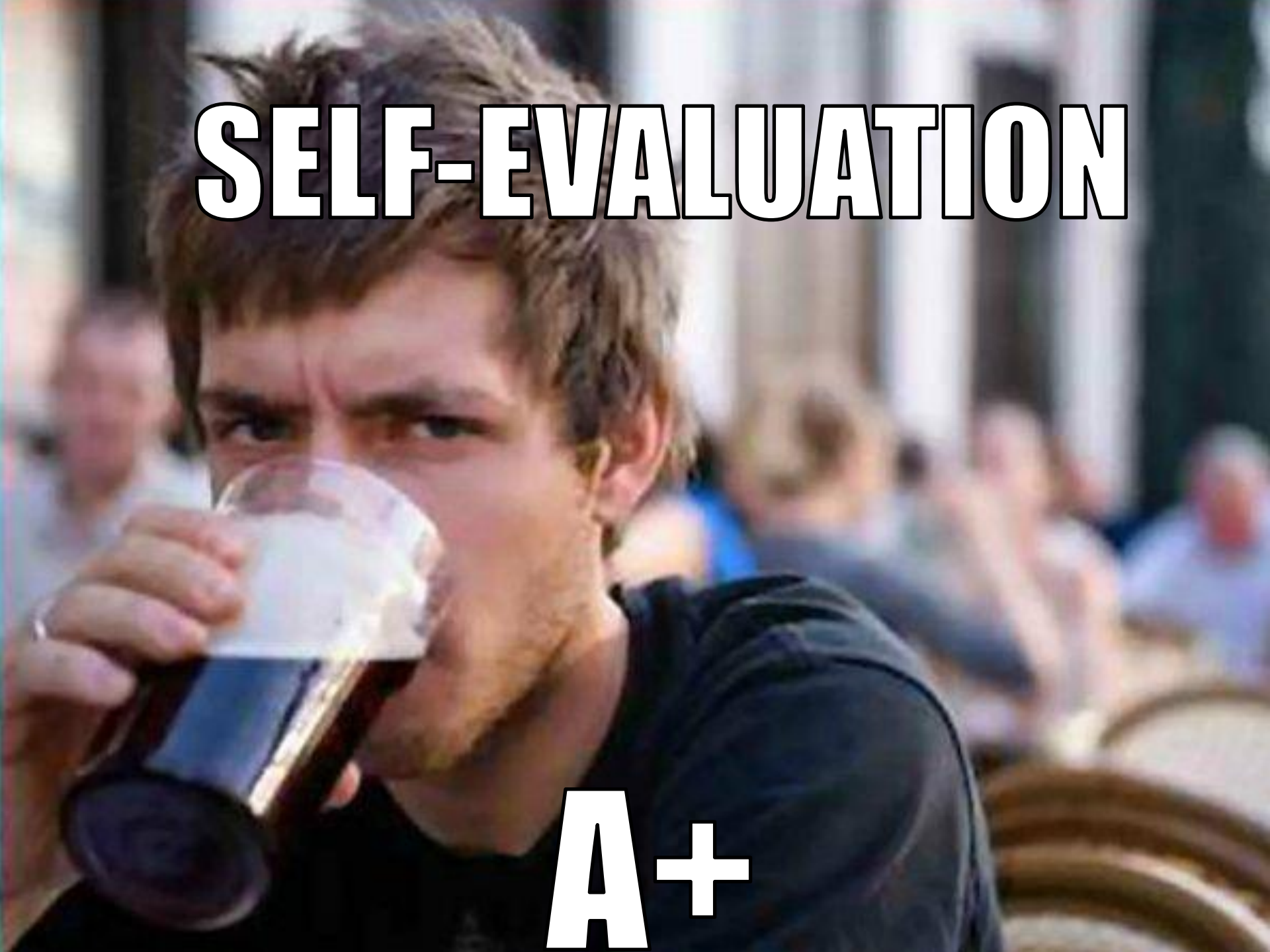
Related Work

Virtual drives

- iSCSI, AoE
- Petal
- EBS, Azure Drive
- Salus

File systems

- BlueSky, pNFS
- OptFS
- BPFS

A young man with brown hair and a slight beard is shown in a close-up, drinking from a large glass of dark beer with a thick head of white foam. He has a satisfied, almost smug expression on his face. The background is a blurred outdoor setting, possibly a pub or festival, with other people visible in the distance.

SELF-EVALUATION

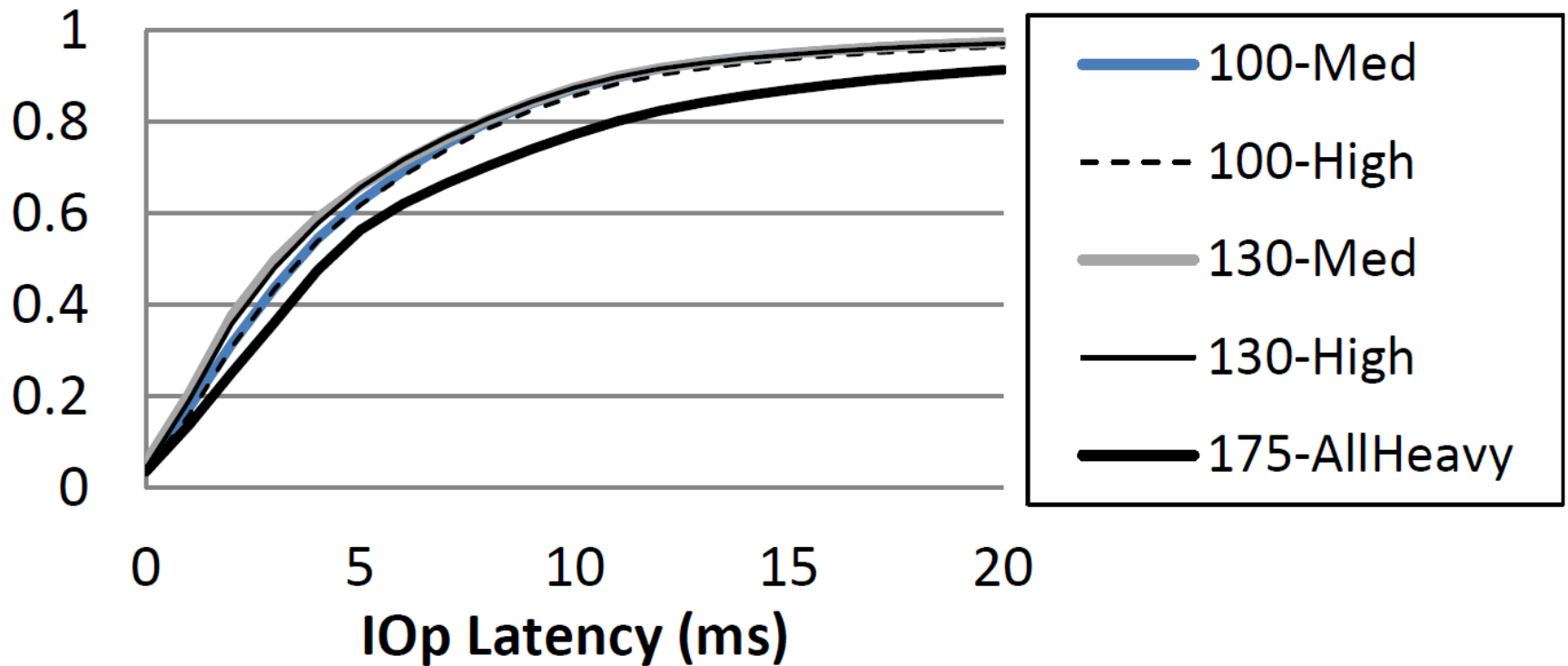
A+

Conclusions

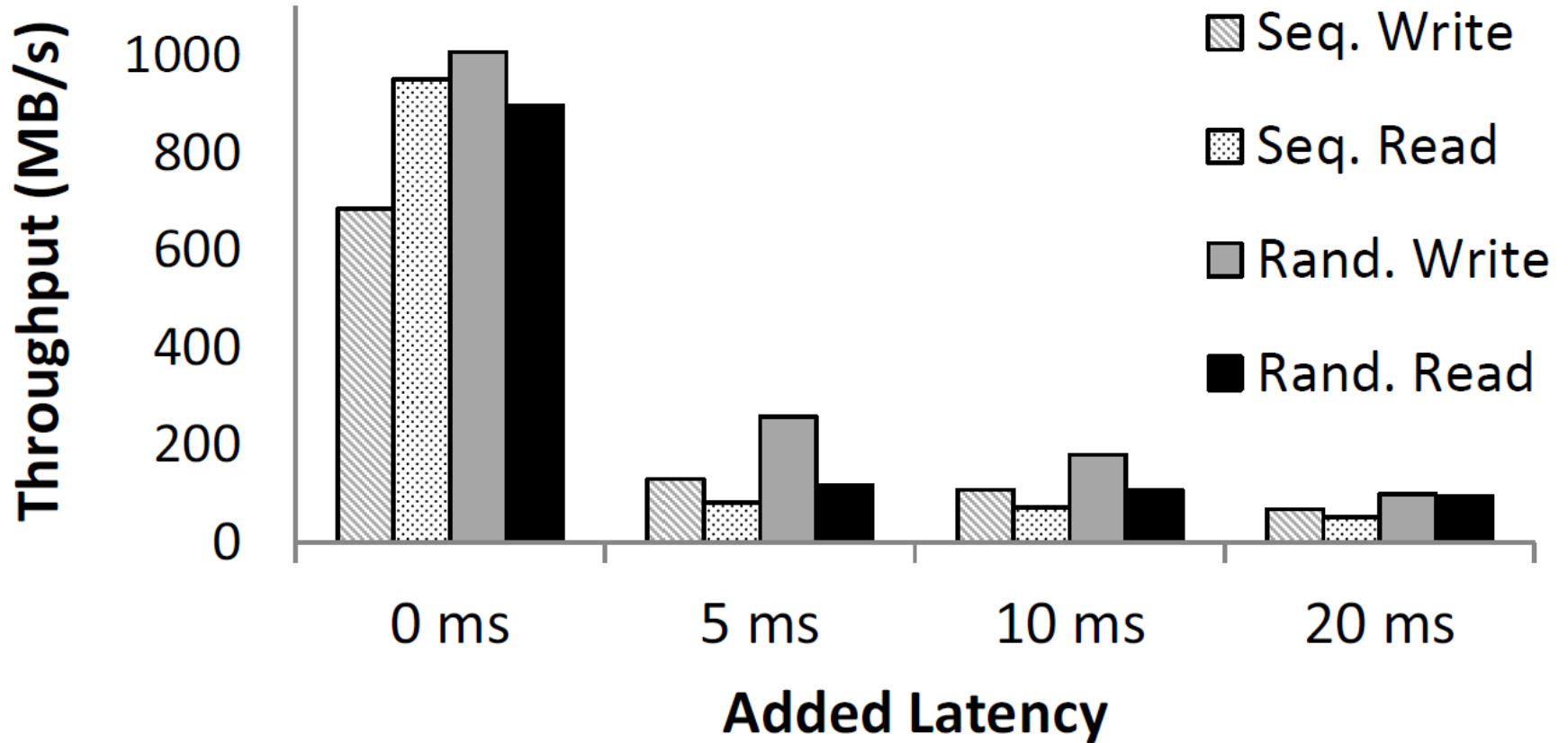
- Unmodified POSIX/Win32 apps can have cloud-scale IO!
 - Nested striping
 - FDS-style hardware substrate
 - Delayed durability semantics
- Raw perf: 1000+ MB/s
- 2x—10x app-level speedups



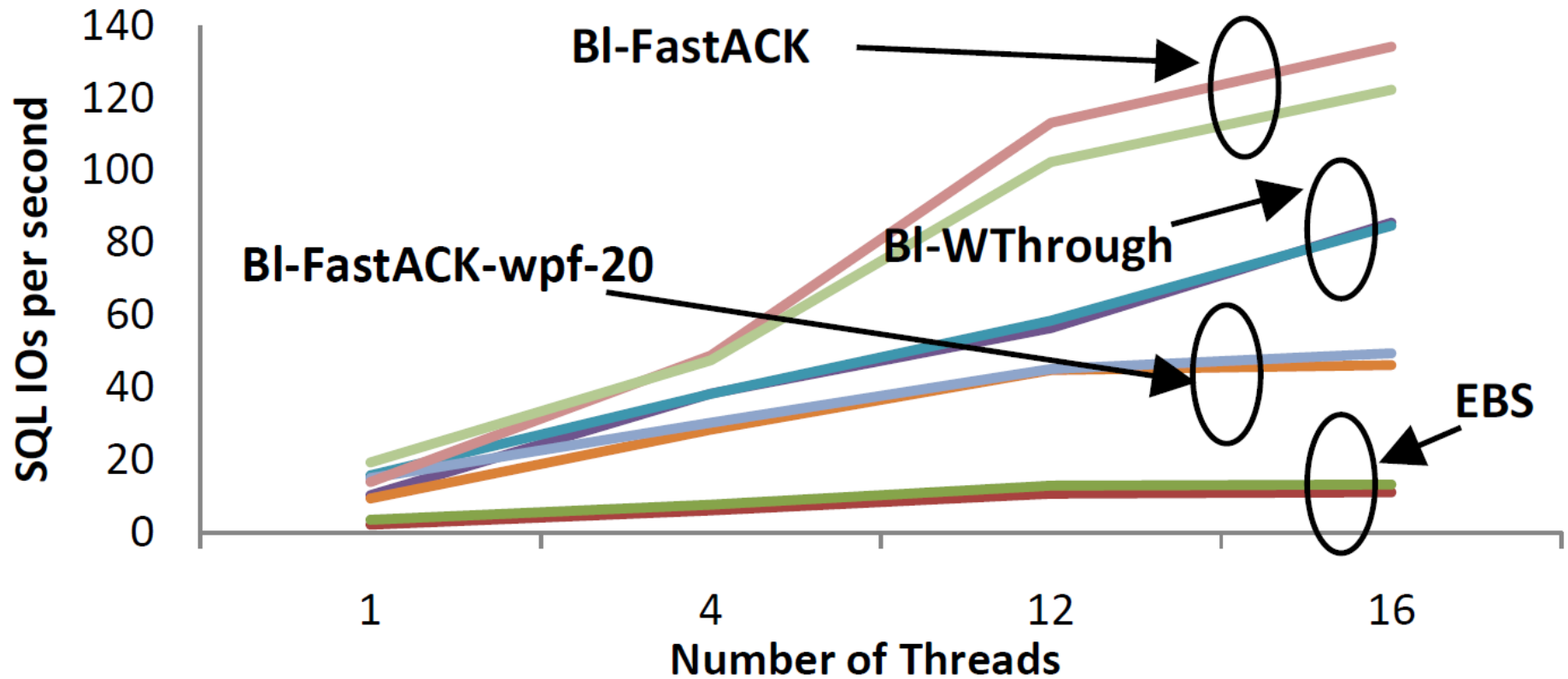
IOP Latency



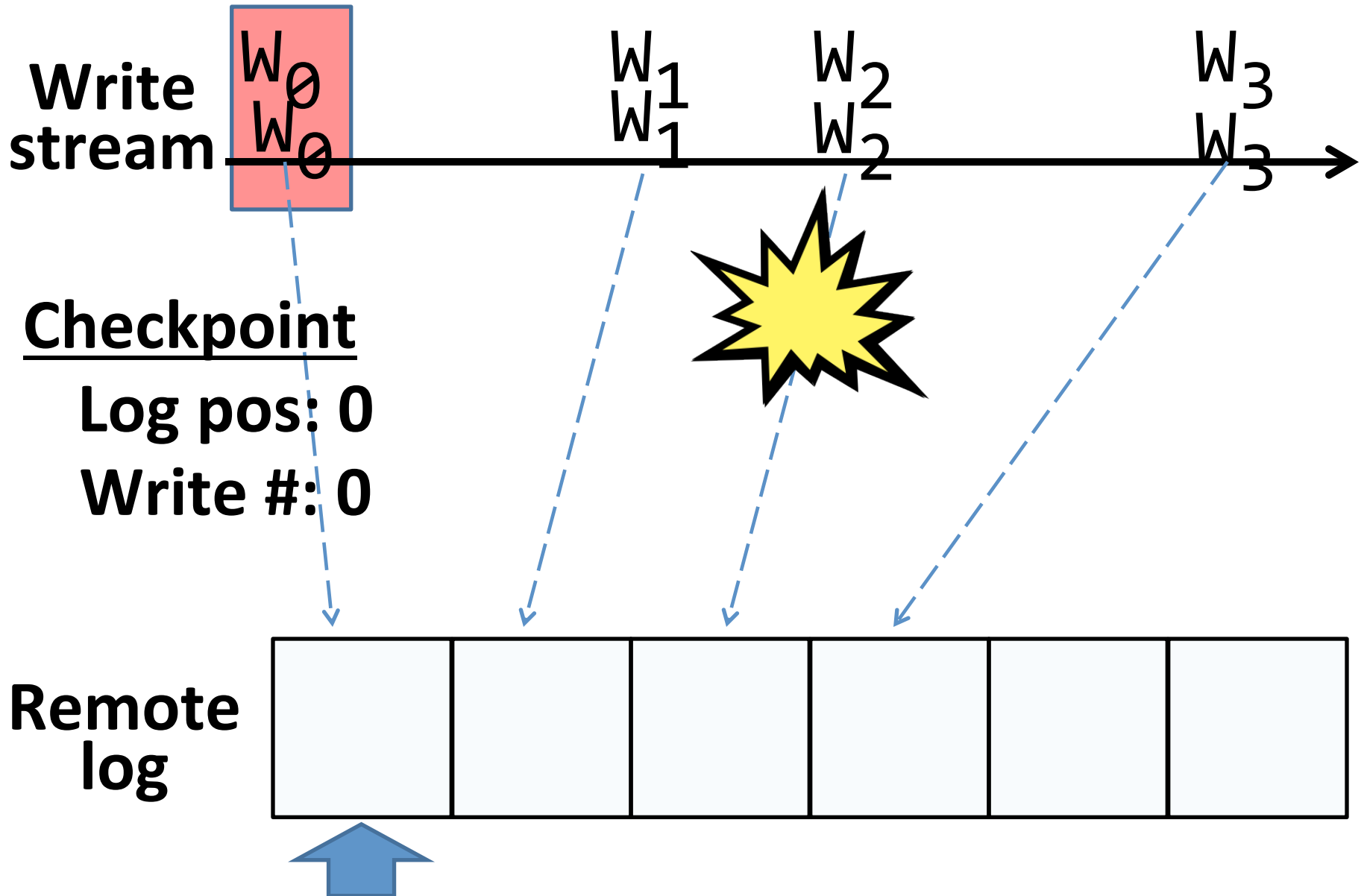
RTT Sensitivity



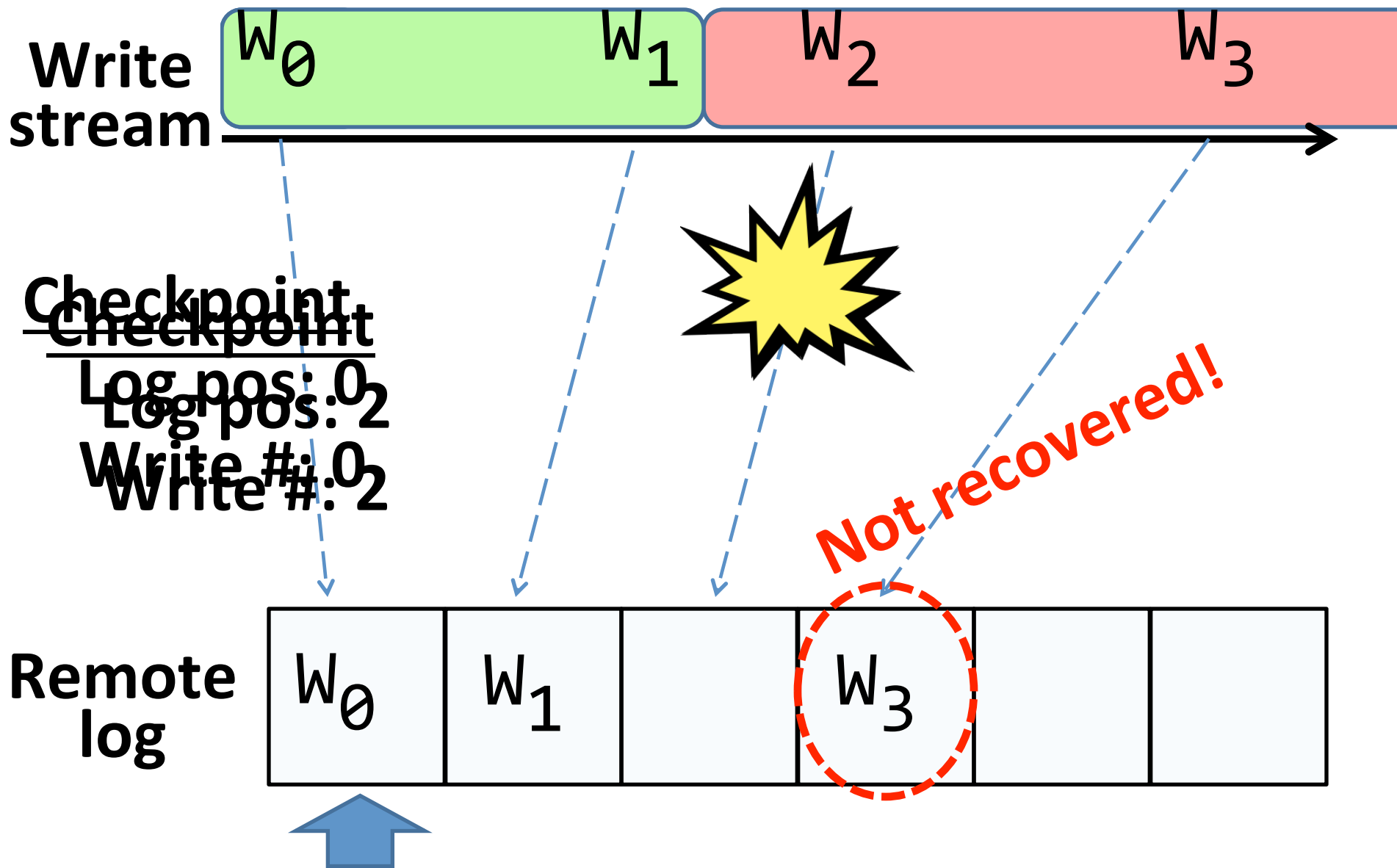
Blizzard vs. EBS



Log-Based Recovery



Recovery



Additional Details

- Blizzard maps each virtual block to backing physical block in the log
 - Allocation map included in checkpoints
- To avoid I/Op dilation, use random permutation to determine next log position!

Prior example: 0, 1, 2, 3, 4, ...

What Blizzard

really does: $X_{n+1} = (aX_n + c) \bmod m$

(Checkpoint a, c, m, X_n)