

Scale it easy: YDB's high performance in a nutshell

Evgenii Ivanov

Senior software engineer, Yandex



Co-organizer

Yandex

About myself

2

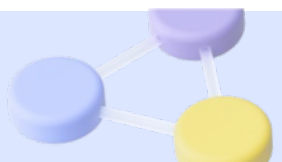
- YDB developer
- Outside YDB I enjoy aerial photography, spending time with my family and reading



Costs of a wrong DB choice

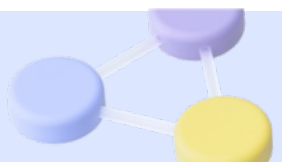
3

- Customer's data loss
- Increased expenses
- Low performance
- Inability to scale



Goals of this talk

- Demonstrate the high availability and superior performance of distributed databases
- Illustrate scenarios where YDB is the optimal choice for your data management needs
- Delve into intriguing technical aspects underpinning YDB's high performance

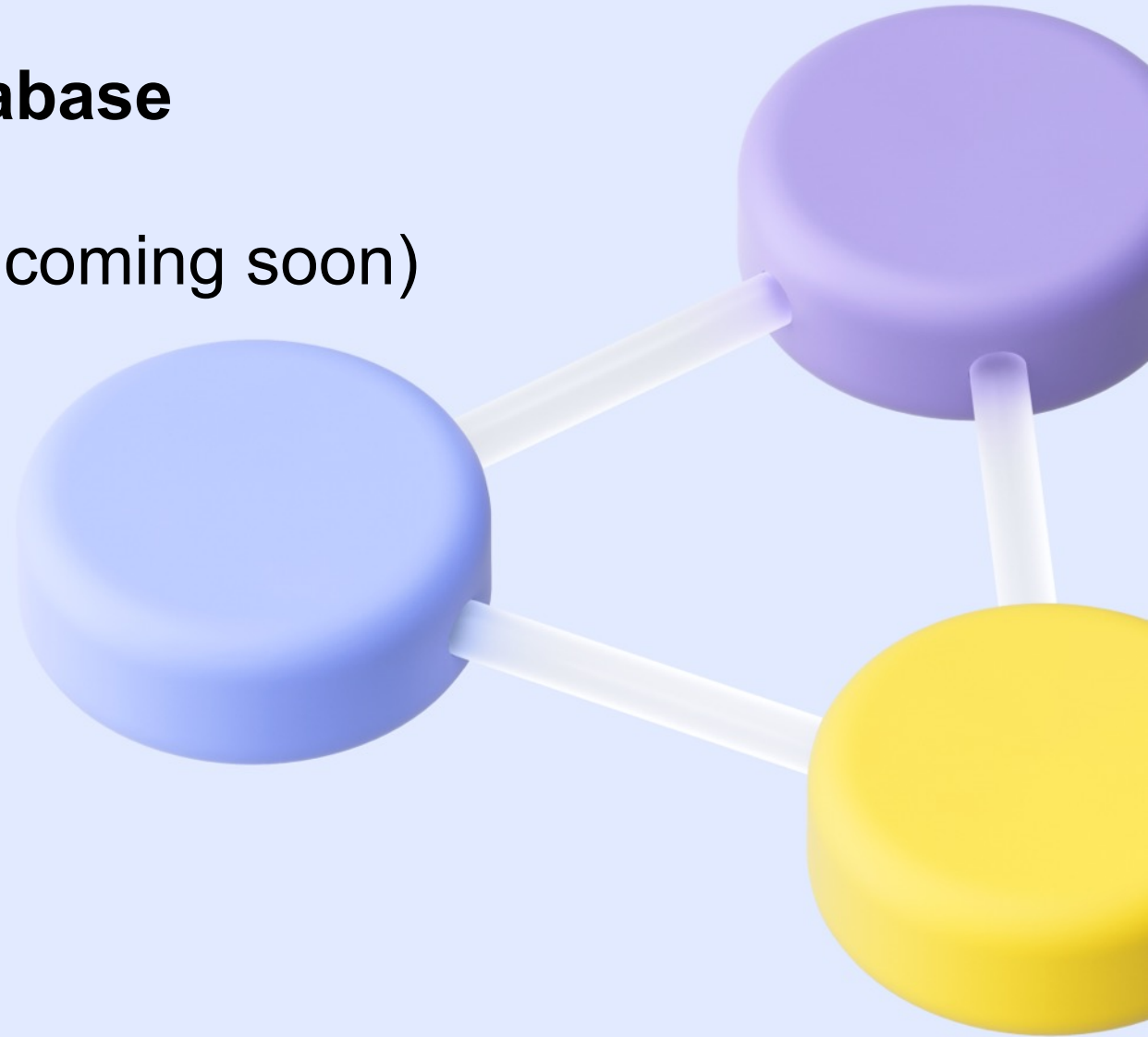


YDB

5

Open-Source Distributed SQL Database

- Relational DB (mainly OLTP, OLAP coming soon)
- Clusters with thousands of servers
- Apache 2.0 license
- Star [ydb-platform](#) on GitHub



YDB is ...

Strictly consistent

- CAP-theorem – we choose CP
- Serializable transaction execution



YDB is ...

7

Highly available and fault tolerant

- Multiple availability zones (AZ): automatic recovery
- YDB is read-write available even after losing AZ and a rack simultaneously



YDB is ...

A mission critical database

- 365x24x7 (366x24x7 when needed)
- Zero maintenance window



What is performance?



- Throughput
- Latency

Usually at YDB we are focusing on achieving max possible throughput with 50 ms latency limit on 99 percentile



Performance has no goal, only path

10

- New hardware
- New algorithms
- New challenges and applications
- Never ending improvements



Performance saves money

- The better DB performance is – the less hardware you need
- With the high availability and scalability you simply earn more



01

Performance tips and tricks for everybody



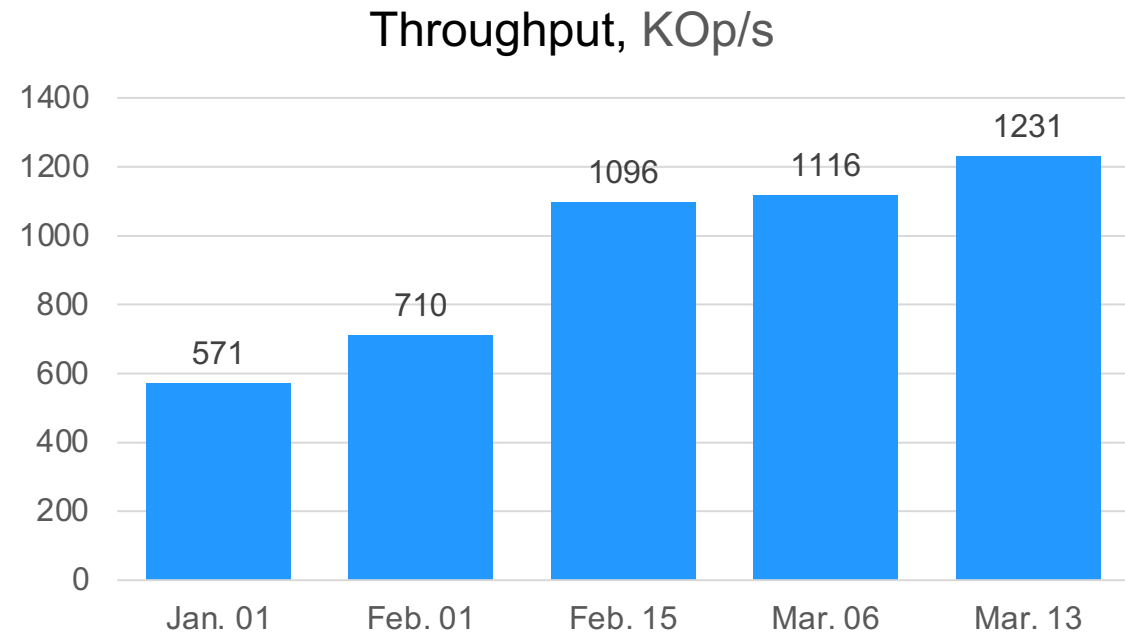
Co-organizer

Yandex

The source of tips is our path

13

- There are a lot of architectural decisions
- Some system configuration tweaks
- Anti bottlenecks campaign
- Lots of flamegraph sessions and experiments



Whoa, can we still process more?

14

```
1 [|||||] 100.0% 33 [|||||] 100.0% 65 [|||||] 100.0% 97 [|||||] 100.0%
2 [|||||] 100.0% 34 [|||||] 100.0% 66 [|||||] 100.0% 98 [|||||] 100.0%
3 [|||||] 100.0% 35 [|||||] 100.0% 67 [|||||] 100.0% 99 [|||||] 100.0%
4 [|||||] 100.0% 36 [|||||] 100.0% 68 [|||||] 100.0% 100 [|||||] 100.0%
5 [|||||] 100.0% 37 [|||||] 98.8% 69 [|||||] 100.0% 101 [|||||] 100.0%
6 [|||||] 100.0% 38 [|||||] 100.0% 70 [|||||] 100.0% 102 [|||||] 100.0%
7 [|||||] 100.0% 39 [|||||] 100.0% 71 [|||||] 100.0% 103 [|||||] 100.0%
8 [|||||] 100.0% 40 [|||||] 100.0% 72 [|||||] 100.0% 104 [|||||] 100.0%
9 [|||||] 100.0% 41 [|||||] 100.0% 73 [|||||] 100.0% 105 [|||||] 100.0%
10 [|||||] 100.0% 42 [|||||] 100.0% 74 [|||||] 100.0% 106 [|||||] 100.0%
11 [|||||] 100.0% 43 [|||||] 100.0% 75 [|||||] 100.0% 107 [|||||] 100.0%
12 [|||||] 100.0% 44 [|||||] 100.0% 76 [|||||] 100.0% 108 [|||||] 100.0%
13 [|||||] 100.0% 45 [|||||] 100.0% 77 [|||||] 100.0% 109 [|||||] 100.0%
14 [|||||] 100.0% 46 [|||||] 100.0% 78 [|||||] 100.0% 110 [|||||] 100.0%
15 [|||||] 100.0% 47 [|||||] 100.0% 79 [|||||] 100.0% 111 [|||||] 100.0%
16 [|||||] 100.0% 48 [|||||] 100.0% 80 [|||||] 100.0% 112 [|||||] 100.0%
17 [|||||] 100.0% 49 [|||||] 100.0% 81 [|||||] 100.0% 113 [|||||] 100.0%
18 [|||||] 100.0% 50 [|||||] 100.0% 82 [|||||] 100.0% 114 [|||||] 100.0%
19 [|||||] 100.0% 51 [|||||] 100.0% 83 [|||||] 100.0% 115 [|||||] 100.0%
20 [|||||] 100.0% 52 [|||||] 100.0% 84 [|||||] 100.0% 116 [|||||] 100.0%
21 [|||||] 100.0% 53 [|||||] 100.0% 85 [|||||] 100.0% 117 [|||||] 100.0%
22 [|||||] 100.0% 54 [|||||] 100.0% 86 [|||||] 100.0% 118 [|||||] 100.0%
23 [|||||] 100.0% 55 [|||||] 100.0% 87 [|||||] 100.0% 119 [|||||] 100.0%
24 [|||||] 100.0% 56 [|||||] 100.0% 88 [|||||] 100.0% 120 [|||||] 100.0%
25 [|||||] 100.0% 57 [|||||] 100.0% 89 [|||||] 100.0% 121 [|||||] 100.0%
26 [|||||] 100.0% 58 [|||||] 100.0% 90 [|||||] 100.0% 122 [|||||] 63.4%
27 [|||||] 100.0% 59 [|||||] 100.0% 91 [|||||] 100.0% 123 [|||||] 100.0%
28 [|||||] 100.0% 60 [|||||] 100.0% 92 [|||||] 100.0% 124 [|||||] 100.0%
29 [|||||] 100.0% 61 [|||||] 100.0% 93 [|||||] 100.0% 125 [|||||] 100.0%
30 [|||||] 100.0% 62 [|||||] 100.0% 94 [|||||] 100.0% 126 [|||||] 100.0%
31 [|||||] 100.0% 63 [|||||] 100.0% 95 [|||||] 100.0% 127 [|||||] 100.0%
32 [|||||] 100.0% 64 [|||||] 100.0% 96 [|||||] 100.0% 128 [|||||] 100.0%
Mem[|||||] 178G/503G Tasks: 71, 1373 thr; 128 running
Swp[|||||] 0K/0K Load average: 133.29 86.81 41.85
Uptime: 143 days(!), 05:26:42
```



Yes! CPU affinity

15

- No code changes – same binary
- Applicable to any multithreaded software and any DB



Yes! CPU affinity

- No code changes – same binary
- Applicable to any multithreaded software and any DB
- CPU caches are crucial for performance
- Important metric: instructions per cycle (IPC)



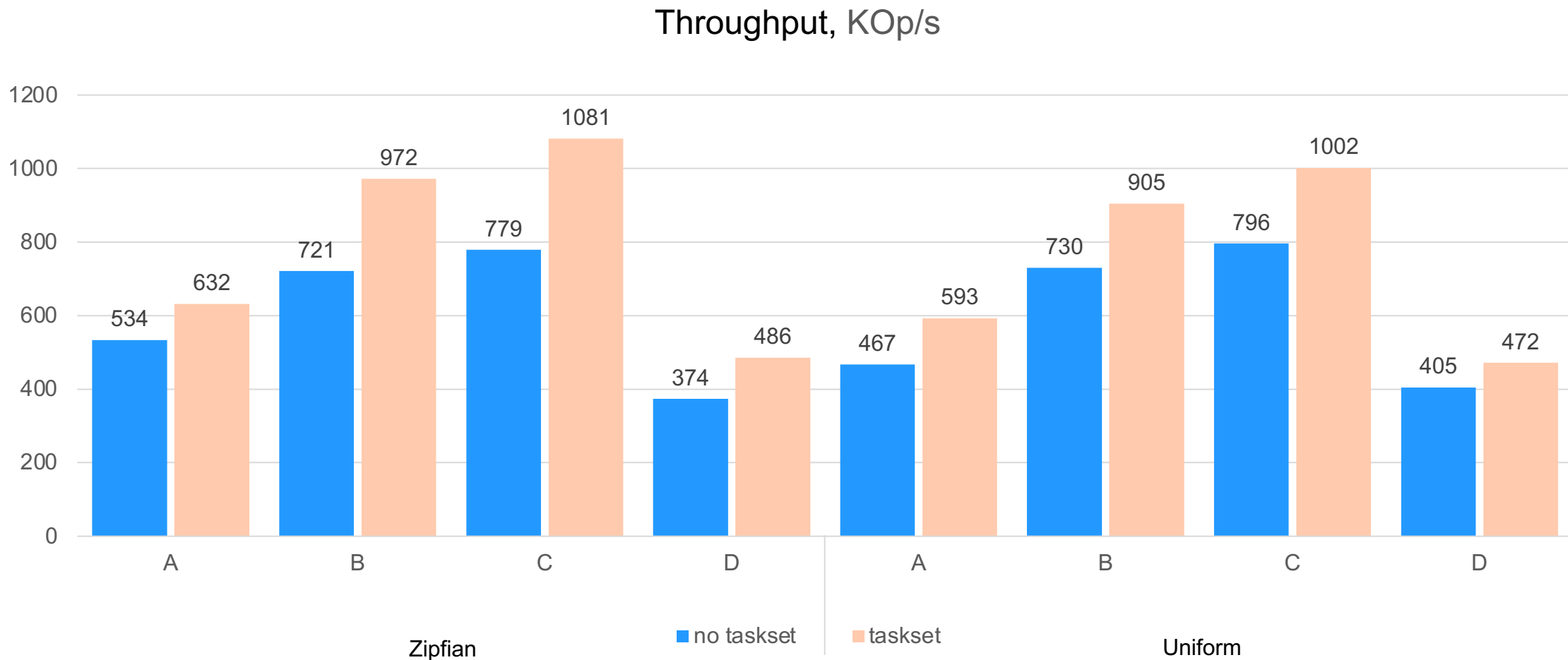
Yes! CPU affinity

- No code changes – same binary
- Applicable to any multithreaded software and any DB
- CPU caches are crucial for performance
- Important metric: instructions per cycle (IPC)
- We want to pin threads to a subset of cores
- Solution: CPU affinity (cpuset/taskset)
- Throughput increases by 20-40% with same CPU consumption (proportionally to IPC)



Yes! CPU affinity

18



Hugepages

- Cache misses are expensive
- TLB misses are even more expensive



Hugepages

- Cache misses are expensive
- TLB misses are even more expensive
- Solution: transparent hugepages
- +3% of free throughput



Actor System

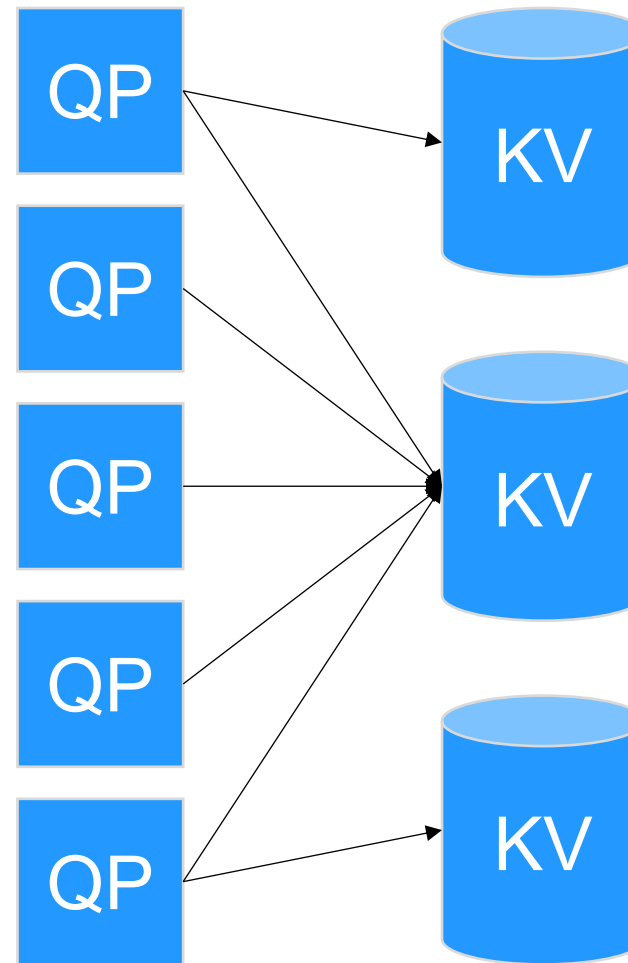
- Is a concurrency model
- Actor – the single threaded entity with own state
- Actors receive messages, send messages and create other actors



Latency bound antipattern

22

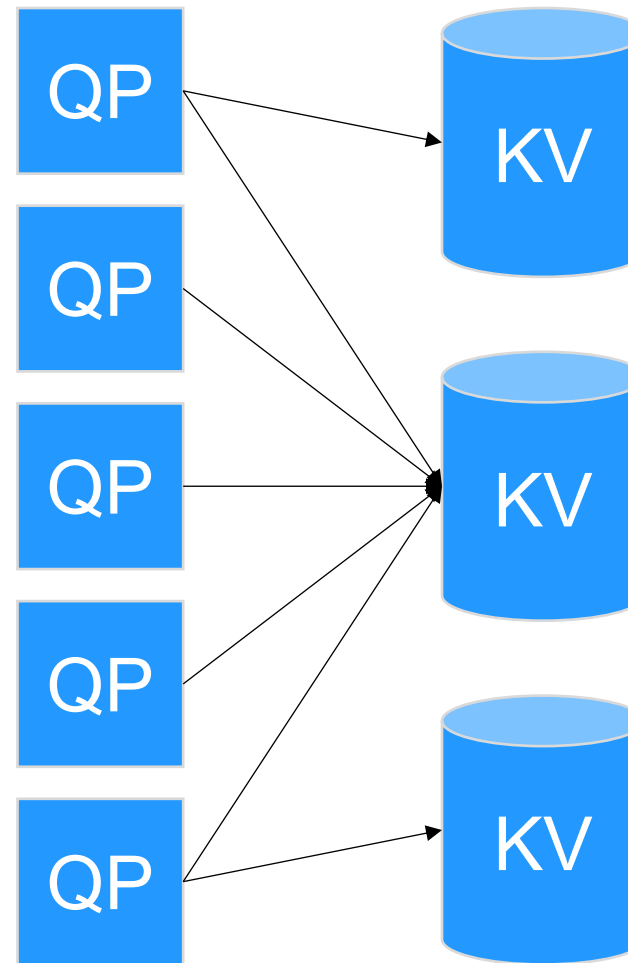
- QP – query processor
- KV – key-value storage
- QP and KV are single threaded
- Queries might be complex
- Push-down is a traditional approach



Latency bound antipattern

23

- Load distribution is uneven
- Pareto principle (the Zipfian distribution)
- Requests to KV are latency bound



Some compiler magic

- Profile-guided optimization (PGO): +20-30% of throughput
- BOLT: post-link binary optimizer: +3-5% of throughput
- Link time optimization (LTO): +1-3% of throughput



Architecture, platform and features

25

- Separate compute and storage layers
- Own storage layer directly on disks ([talk](#))
- Own advanced actor system ([talk](#))
- C++
- MVCC ([talk](#)) and other features
- Beyond Calvin



02

YCSB: key-value performance



Co-organizer

Yandex

Yahoo! Cloud Serving Benchmark (YCSB)

- A popular key-value benchmark
- Created for NoSQL key-value DBs, but still loved by everybody
- Supports almost all modern databases
- It's hard to do distributed transactions well if you can't do key-value workloads well



YCSB workloads

- A (update heavy workload): 50% reads and 50% updates
- B (read mostly workload): 95% reads and 5% updates
- C (read only)
- D (read latest workload): 95% reads, 5% inserts
- F (read-modify-write): 50% reads and 50% read-update operations
- E (short ranges): 95% scans and 5% inserts.



Test setup

- 128 cores: 2x32-cores Intel Xeon Gold 6338 CPU @ 2.00GHz with hyper-threading turned on
- 4xNVMe
- 512 GB RAM
- 50 Gb network
- Transparent hugepages turned on
- Ubuntu 20.04.3 LTS

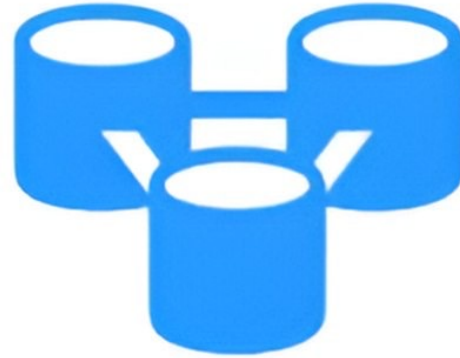


Distributed SQL Databases

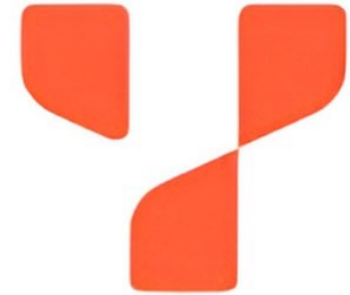
30



VS.



VS.



A bar fight?



Like this?

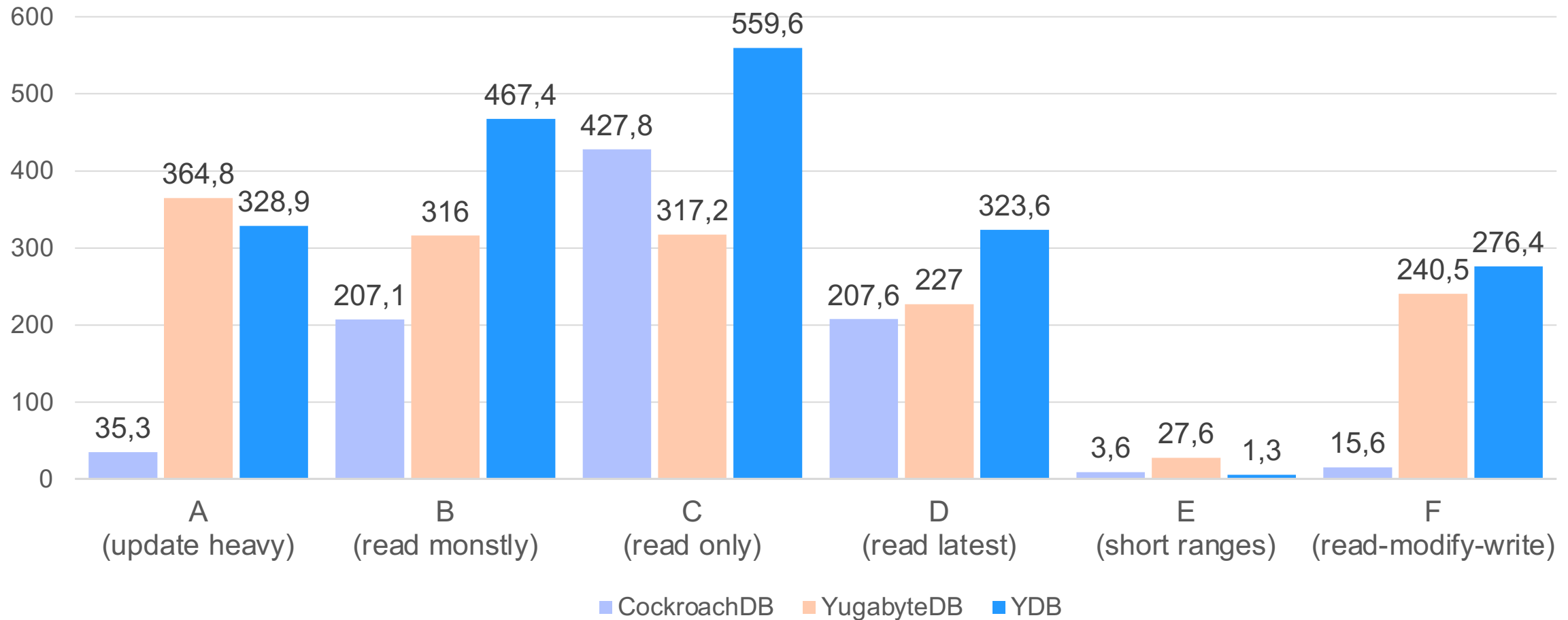




YCSB: 300 GB data

34

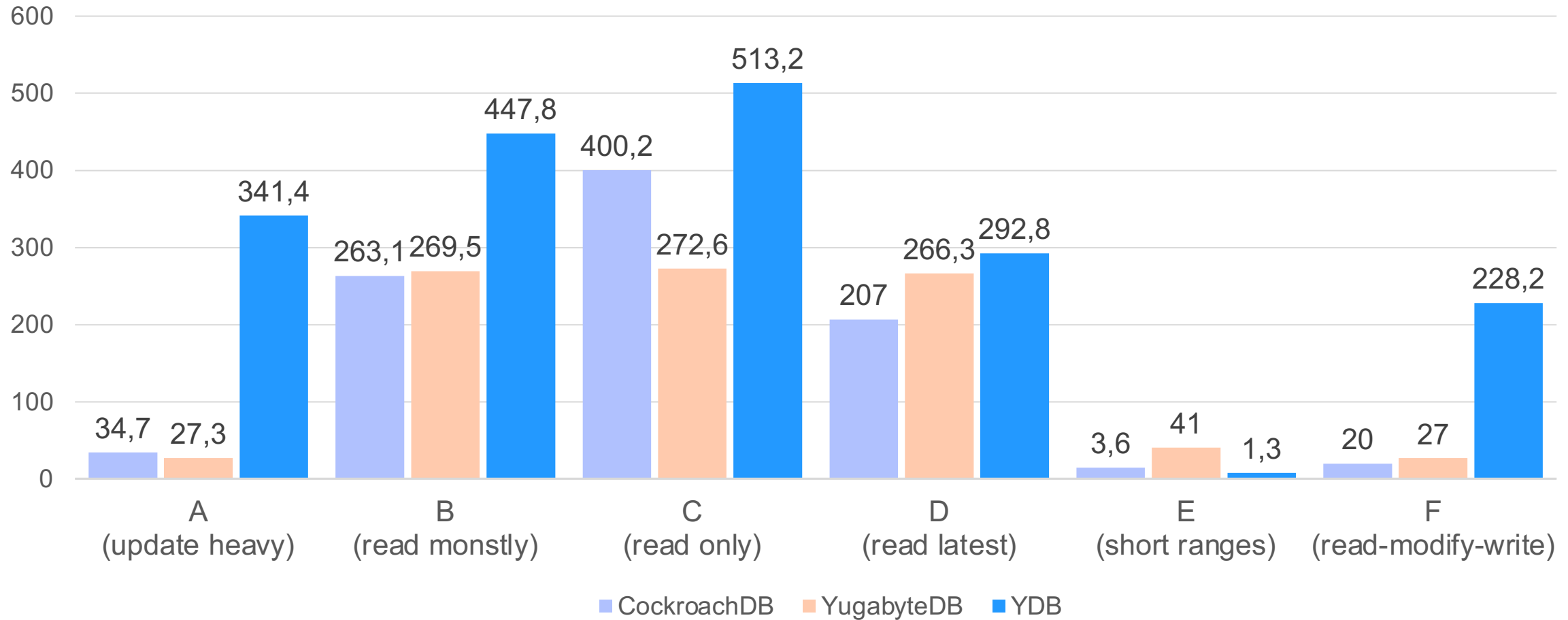
Throughput, KOp/s (higher is better)



YCSB: 2 TB data

35

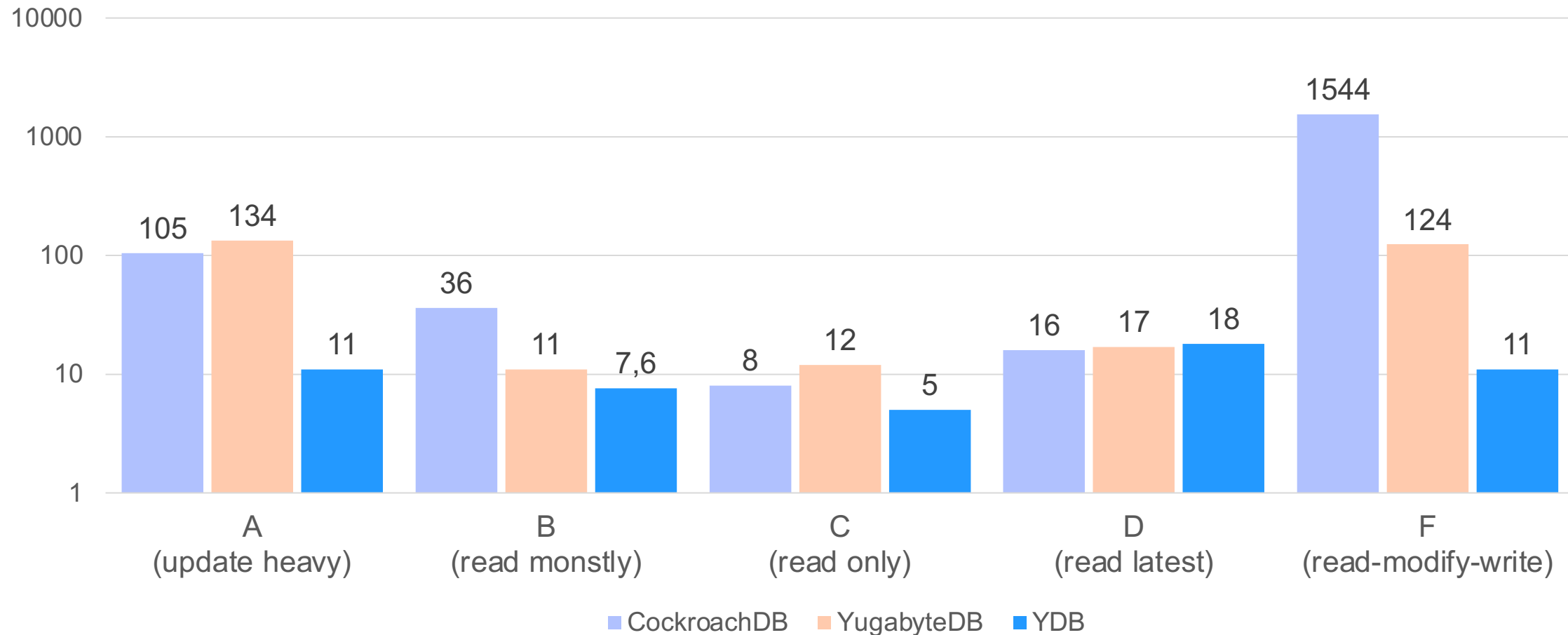
Throughput, KOp/s (higher is better)



YCSB: 2 TB data

36

Latency, ms (lower is better)



A silver bullet?

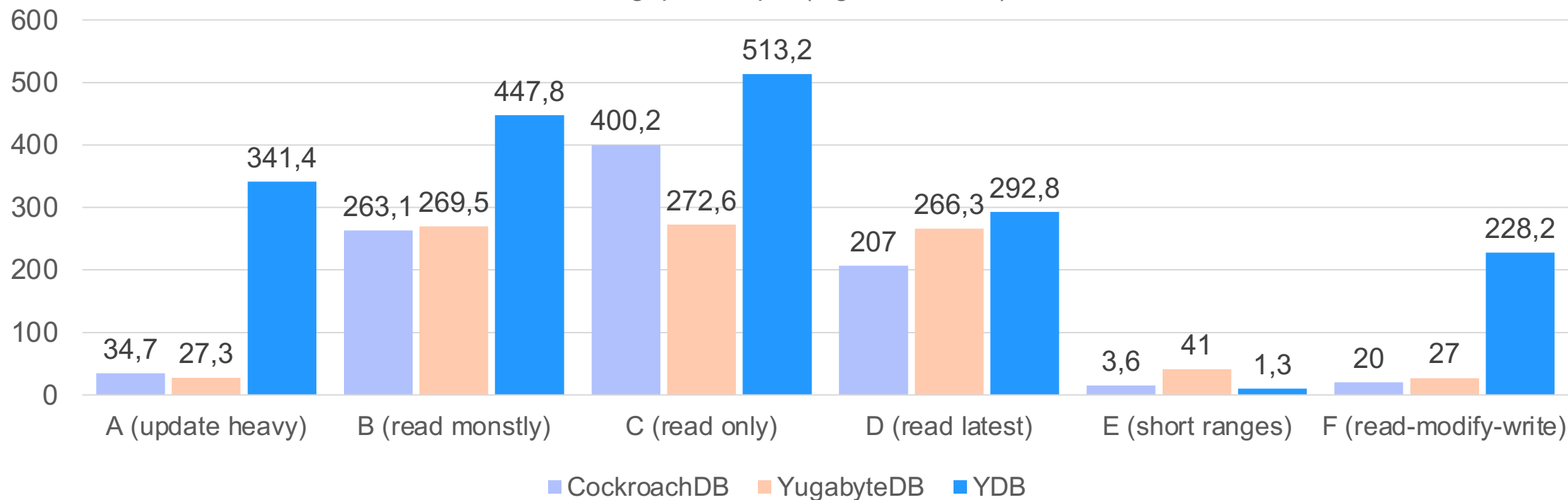


Analyses

38

- UPSERT vs. UPDATE: blind write VS. read-modify transaction
- CockroachDB uses UPDATE in workloads A, B and F

Throughput, KOp/s (higher is better)

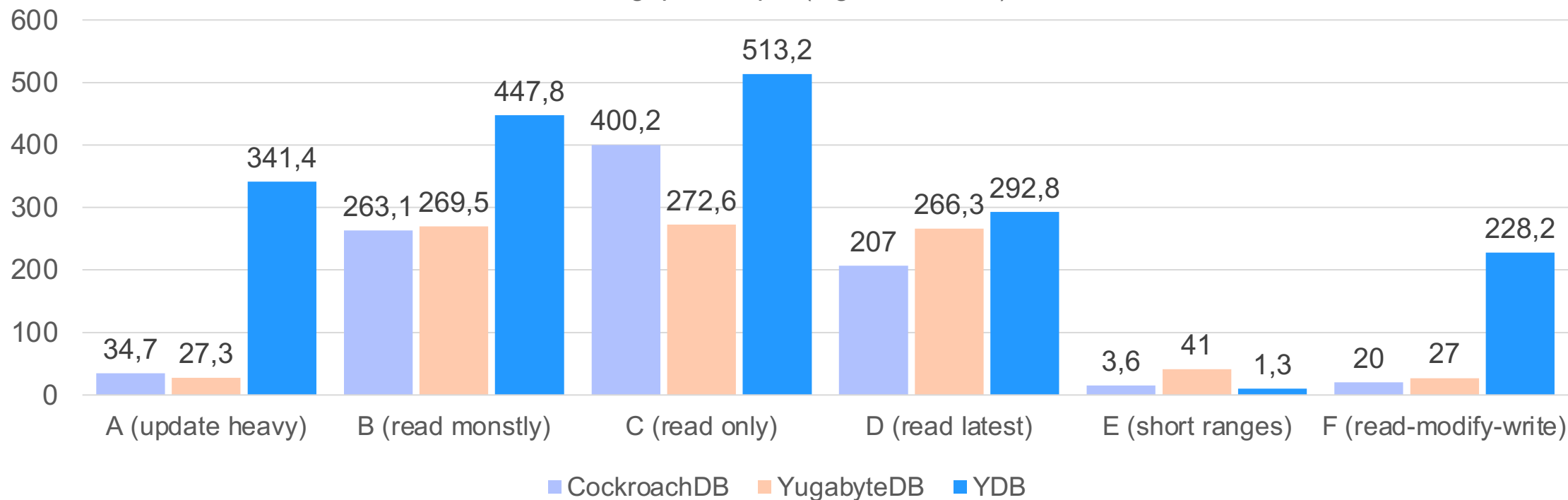


Analyses

39

- YugabyteDB has a bare metal horizontal scaling issue
- YugabyteDB has a vertical scaling issue

Throughput, KOp/s (higher is better)

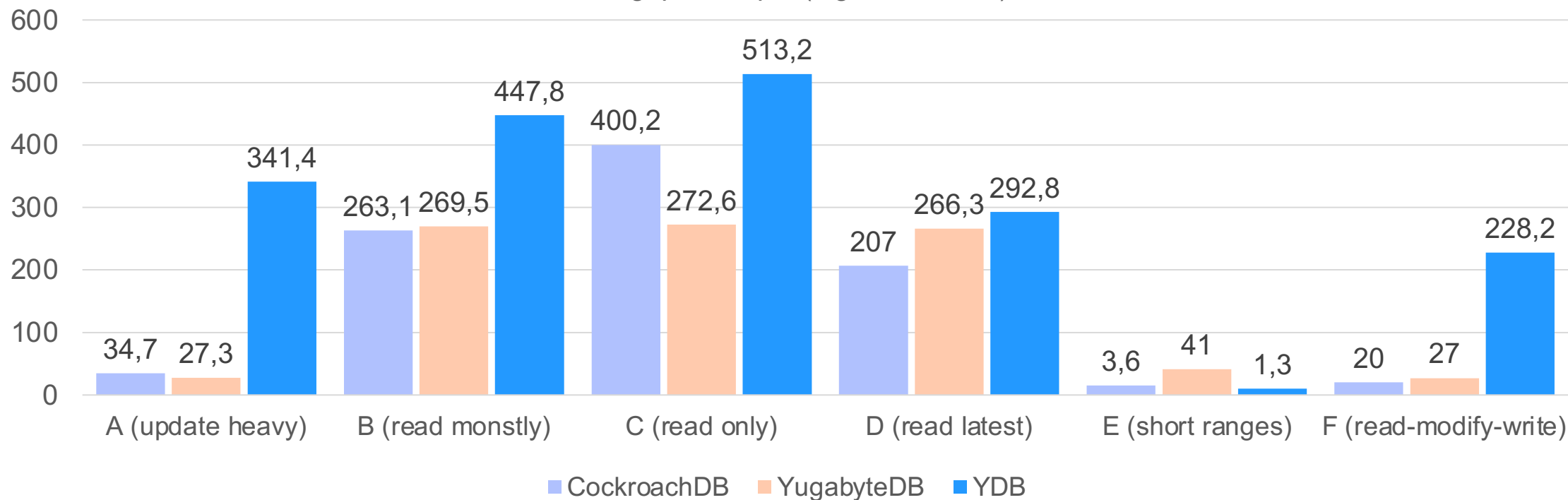


Analyses

40

- Workload C is the same among all DBs
- YDB has an issue with short ranges (workload E)

Throughput, KOp/s (higher is better)



YCSB conclusion

- Be careful when you implement your application
- YDB outperforms others in many YCSB workloads
- In write intensive workloads all 3 DBs exhibited decent performance
- We strongly recommend YDB for read-mostly and read-only workloads

Full YCSB results: <https://bit.ly/3WfMZEq>



03

Distributed SQL vs “Centralized” SQL



Co-organizer

Yandex

Other DBs to compare with?

43



Ask your DB a question

- What will happen if a disk/server/datacenter fails?
- How can I reduce the latency if users are far from their DC?
- What latency will I have when working set exceeds RAM (typically 512 GiB, rarely above 1024 GiB)?



Big data questions to DB

45

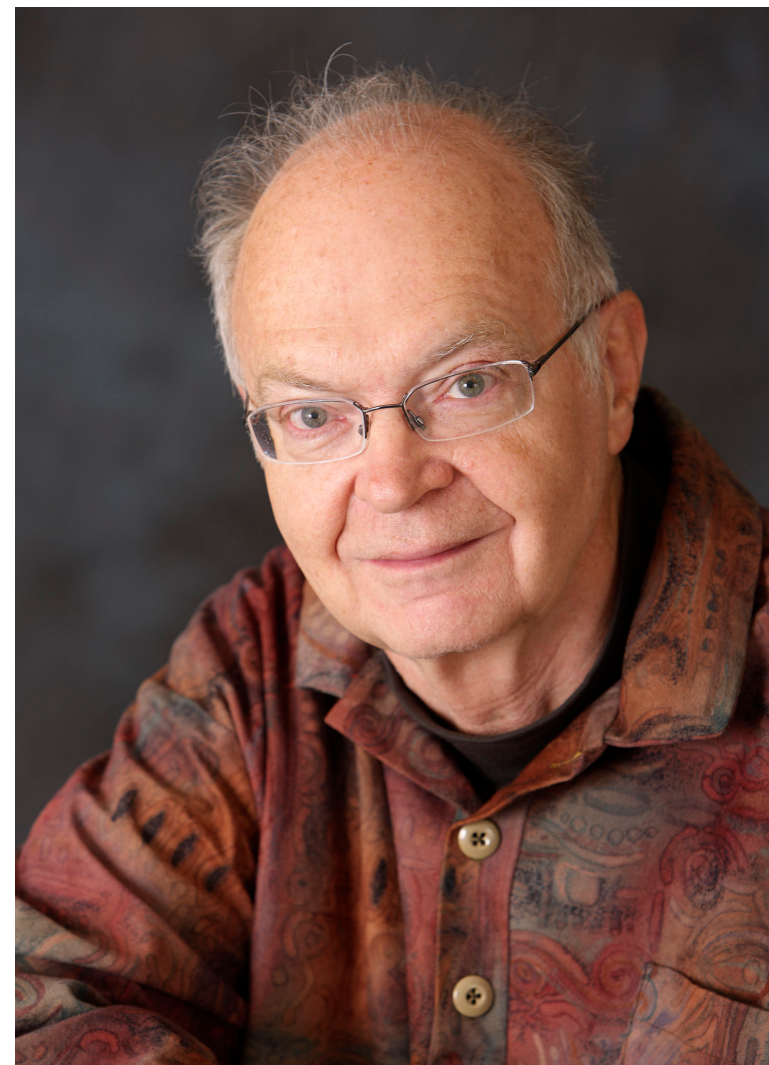
- Can you store a petabyte of data?
- Can you handle more than 100 GB/s of data updates?



“The best theory is inspired by practice. The best practice is inspired by theory.”

Donald Knuth

- There are theoretical reasons behind the advantages of a distributed DB
- Distributed DB theory is based on the best practice



What do we expect from DBs?

47

- ACID – of course!
- Scalability
- Maintainability
- High performance
- Security
- Reliability / fault tolerance / availability



Nothing is reliable

- Software bugs
- Hardware bugs
- Hardware faults
- Unreliable networks

Can a single server DB be reliable and available? No.



Nothing is reliable



Is replication a solution?

50



Replication pitfalls

51

Async replication:

- risk of data loss (not durable)
- stale reads and anomalies

Avoid async replication until you know what you're doing!



Sync replication pitfalls

Sync replication is better, but still there is one single leader:

- only reads can be split among replicas
- load balancer (like HAProxy)
- failover and failback (like Patroni)
- failover requires consensus: etcd, ZooKeeper, etc
- when failed node is back, recovery might take a long time



Vertical and horizontal scaling

53



IBM z16 mainframe



A small part of typical YDB cluster



Sharding

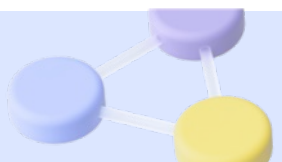
- Horizontal scaling
- Geographical scaling
- Many leader shards (partitions/tablets)



Sharding pitfalls

55

- Consistency – distributed snapshots – complicated!
- You need a coordinator like Citus as well as standby and failover procedure for the coordinator



The choice

Do you really want to maintain a DB with a separate

- load balancer
- consensus cluster
- failover and failback
- coordinator



Conclusion



Co-organizer

Yandex

Conclusions

- There are ways to ruin performance on the application level



Conclusions

- There are ways to ruin performance on the application level
- "Centralized" SQL doesn't match mission critical data
- Distributed DB is the proper solution for availability and scaling challenges



Conclusions

- There are ways to ruin performance on the application level
- "Centralized" SQL doesn't match mission critical data
- Distributed DB is the proper solution for availability and scaling challenges
- There are scenarios, where YDB outperforms other distributed DBs
- You might want to consider YDB when choosing distributed DB for your data



Please leave your feedback



Evgenii Ivanov (tg, twitter, linkedin: @eivanov89)

Senior software developer, Yandex

Slides: bit.ly/3XufHlr



Yandex on the conference



Co-organizer

Yandex