

FAST In-memory Transaction Processing using New Hardware Features

—— *A Perspective from Systems Software* ——

RONG CHEN

Institute of Parallel and Distributed Systems
Shanghai Jiao Tong University, China

Joint work with Haibo, Xingda, Jiaxin, and Yanzhe@IPDADS

Transaction: Key Pillar for Many Systems

amazon.com

426 items/second

Amazon sold 426 items per second during its 'best ever' holiday season

Alibaba.com

\$9.3 billion/day

Alibaba's 'Singles Day' rings up \$9.3 billion in new record

Demand Speedy Distributed Transaction **Over Large Data Volumes**



9.56 million

tickets/day

more than 9.56 million tickets on Dec 19, a new high daily sales, as people rushed to book tickets home before the Spring Festival.

PayPal

11.6 million

payments/day

processing almost 11.6 million payments for our customers per day.

Business Demand: **High Throughput**



**GLOBAL SHOPPING
FESTIVAL 2016**

Peak transactions per second

175,000 new orders

120,000 payment

Business Demand: **Low Latency**



Evolution and Practice: Low-latency
Distributed Applications in Finance

The finance industry has unique demands for low-latency distributed systems

- Read input message from network and parse – 5 microseconds
- Look up client profile – 3.2 milliseconds (3,200 microseconds)
- Compute client quote – 15 microseconds
- Log quote – 20 microseconds
- Serialize quote to a response message – 5 microseconds
- Write to network – 5 microseconds

Conventional Approach [Before 2015]

TPC-C World Record

SPARC SuperCluster*

Perf: **504 K TXs/s**

Cost: **30 Million USD**

Cost / Perf: **60.6** [2010]



IBM Power 780 Server+

Perf: **173 K TXs/s**

Cost: **14 Million USD**

Cost / Perf: **82.6** [2010]



* http://tpc.org/tpcc/results/tpcc_result_detail5.asp?id=110120201

+ http://tpc.org/tpcc/results/tpcc_result_detail5.asp?id=110081702

Conventional Approach [Now]

TPC-C World Record

Alibaba Cloud / OceanBase* [2020]

Perf: **11.8 M TXs/s**, Cost: **2,815 Million CNY**, Cost / Perf: **238.7**



* http://tpc.org/tpcc/results/tpcc_result_detail5.asp?id=120051701

High **COST** for Distributed TX

Many **scalable** systems have **LOW perf.**

- Usually 10s~100s of thousands of TX/second
- High COST* (config. that outperform single thread)
- e.g., HStore, Calvin^{SIGMOD'12}

Emerging speedy TX systems **NOT scale-out**

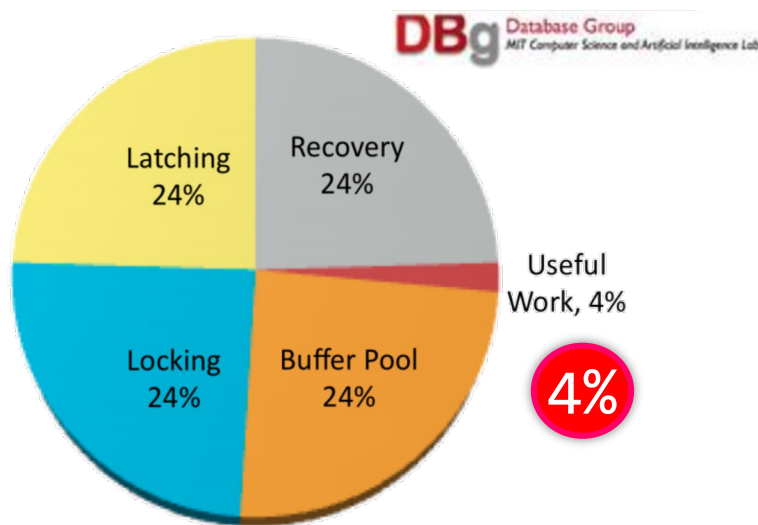
- Achieve over 100s of thousands TX/second
- e.g., Silo @SOSP'13, DBX @EuroSys'14

Dilemma
high perf. vs. scale-out

Why **distributed** TXs are Slow?

Only **4% of wall-clock time** spent on useful data processing, while the rest is occupied with **buffer pools, locking, latching, recovery**.^{*}

— Michael Stonebraker

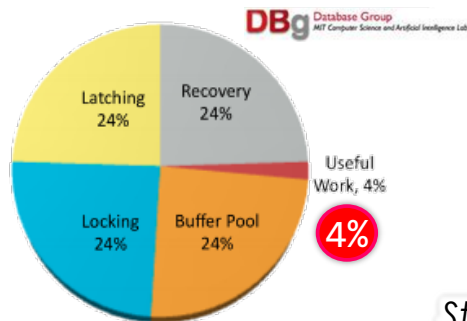


Michael
Stonebraker

2014 Turing Award

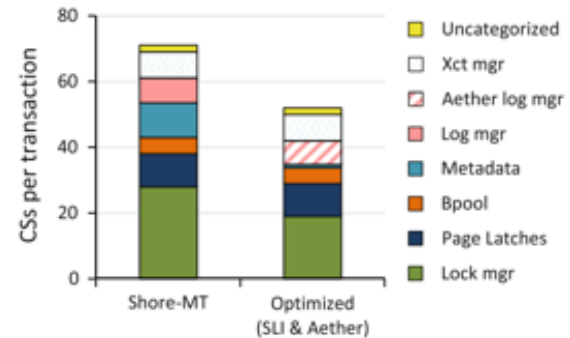
^{*} "The Traditional RDBMS Wisdom is All Wrong"

Why distributed TXs are Slow?

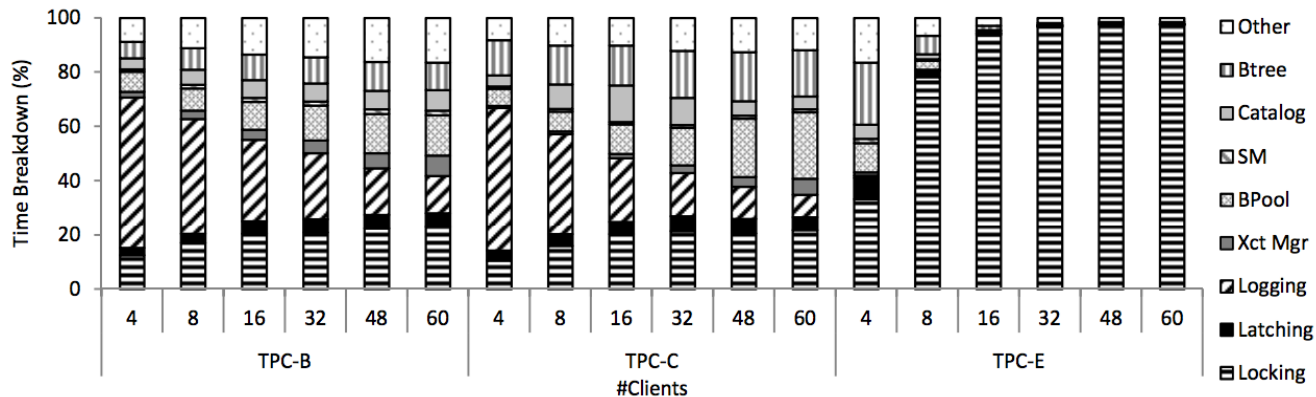


Michael Stonebraker

2014 Turing Award



Eliminating unscalable communication in transaction processing, VLDBJ 2014



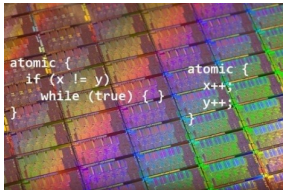
From A to E: Analyzing TPC's OLTP Benchmarks, EDBT/ICDT'13

Our Approach: New Hardware

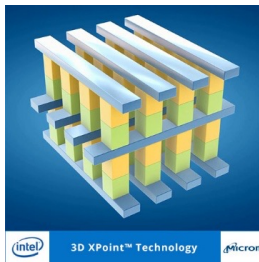
DrTM / DELL Cluster (16 nodes)

Perf: $> 6 \text{ M TXs/s}$, Cost: $< 1 \text{ Million CNY}$, Cost / Perf: < 0.17

HTM

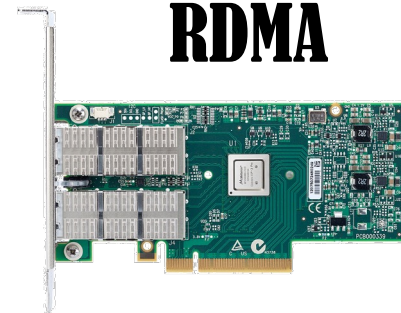


NVM



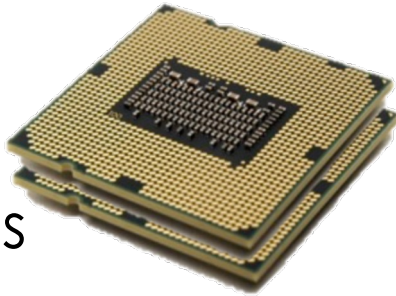
A few
rack-scale
machines

RDMA

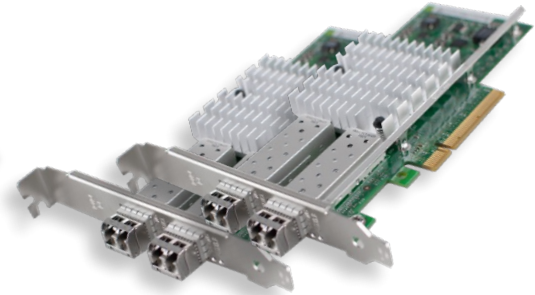


Traditional Server @DC

CPUs:
16-32 cores



1 GB/s
each



NICs: 10 Gbps

50 GB/s



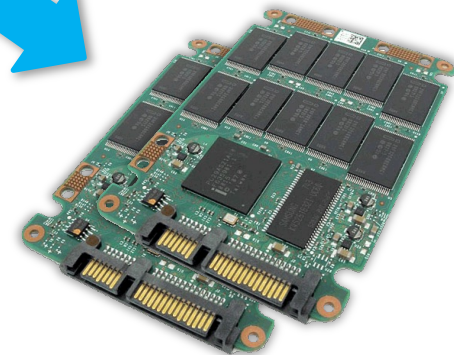
DRAM: 32-64 GB

0.1 GB/s
each



HDDs: 2-40 TB

1 GB/s each

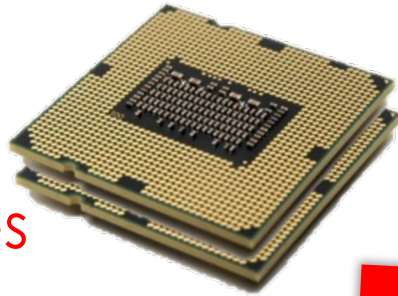


SSDs: 0.25-1 TB

“Not So New” Hardware Features

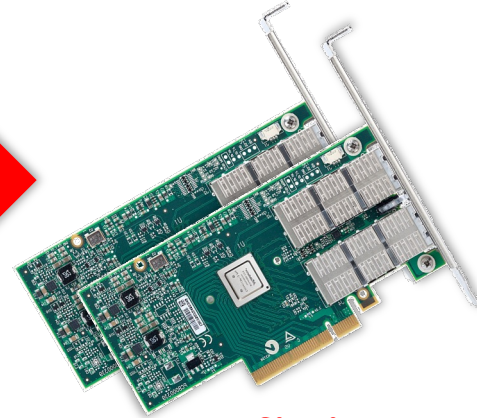
CPU:

16-32 cores
w/ **HTM**



10 GB/s

each



InfiniBand NICs:

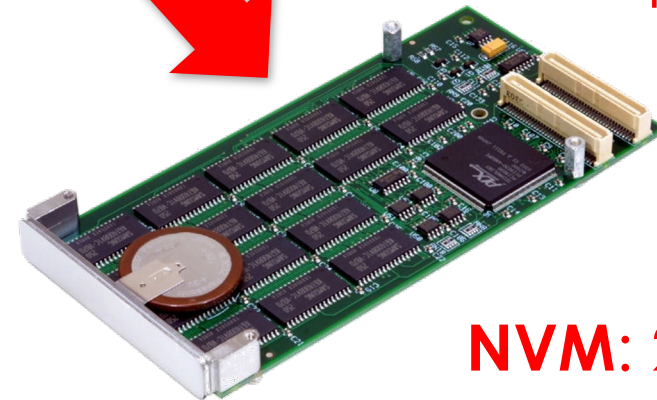
100-200 Gbps
w/ **RDMA**

50 GB/s



DRAM: 128-256 GB

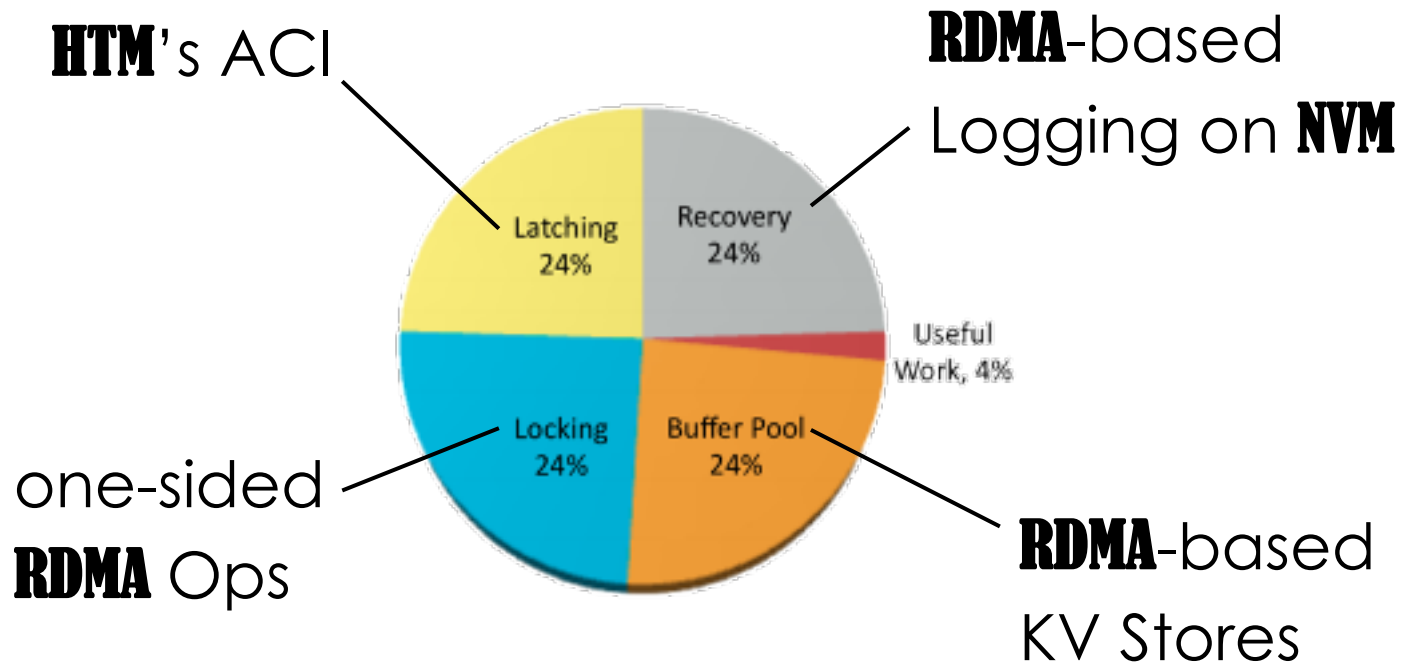
10-50 GB/s



NVM: 256-512 GB
w/ **Persistence**

General Idea

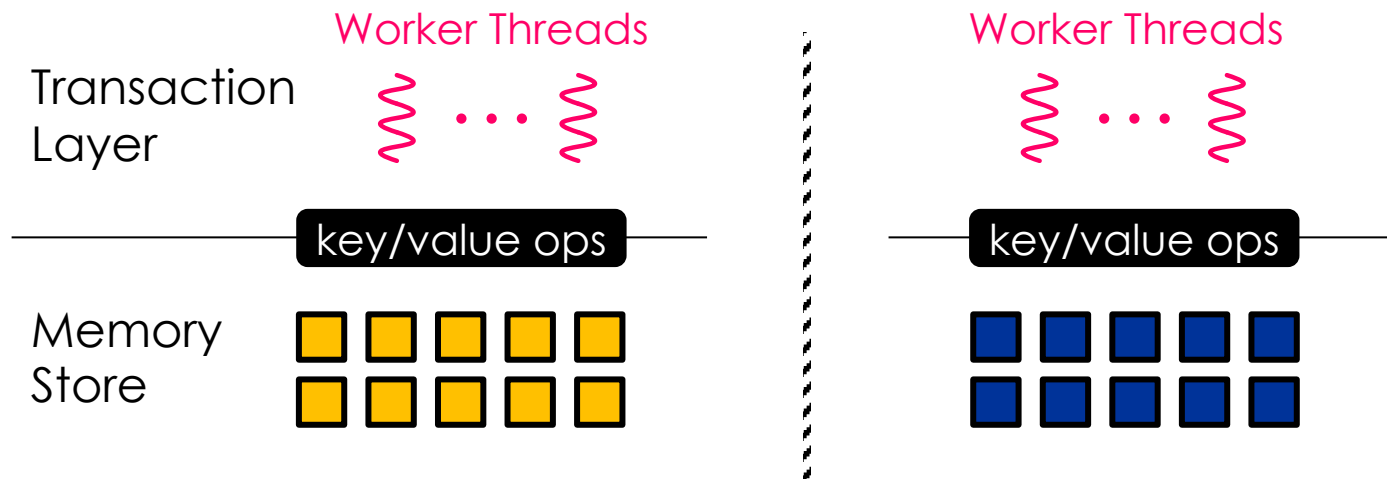
Leveraging **new hardware features** to accelerate distributed transaction processing



System Architecture

Distributed transaction processing

- **Target:** OLTP workloads over large volume of data
- Two independent components on each node
 - ▶ Transaction layer & memory store
- Low COST distributed TX



This Course

101 of New Hardware Features

Hardware



Distributed Transaction Protocol



In-memory Key-Value Store

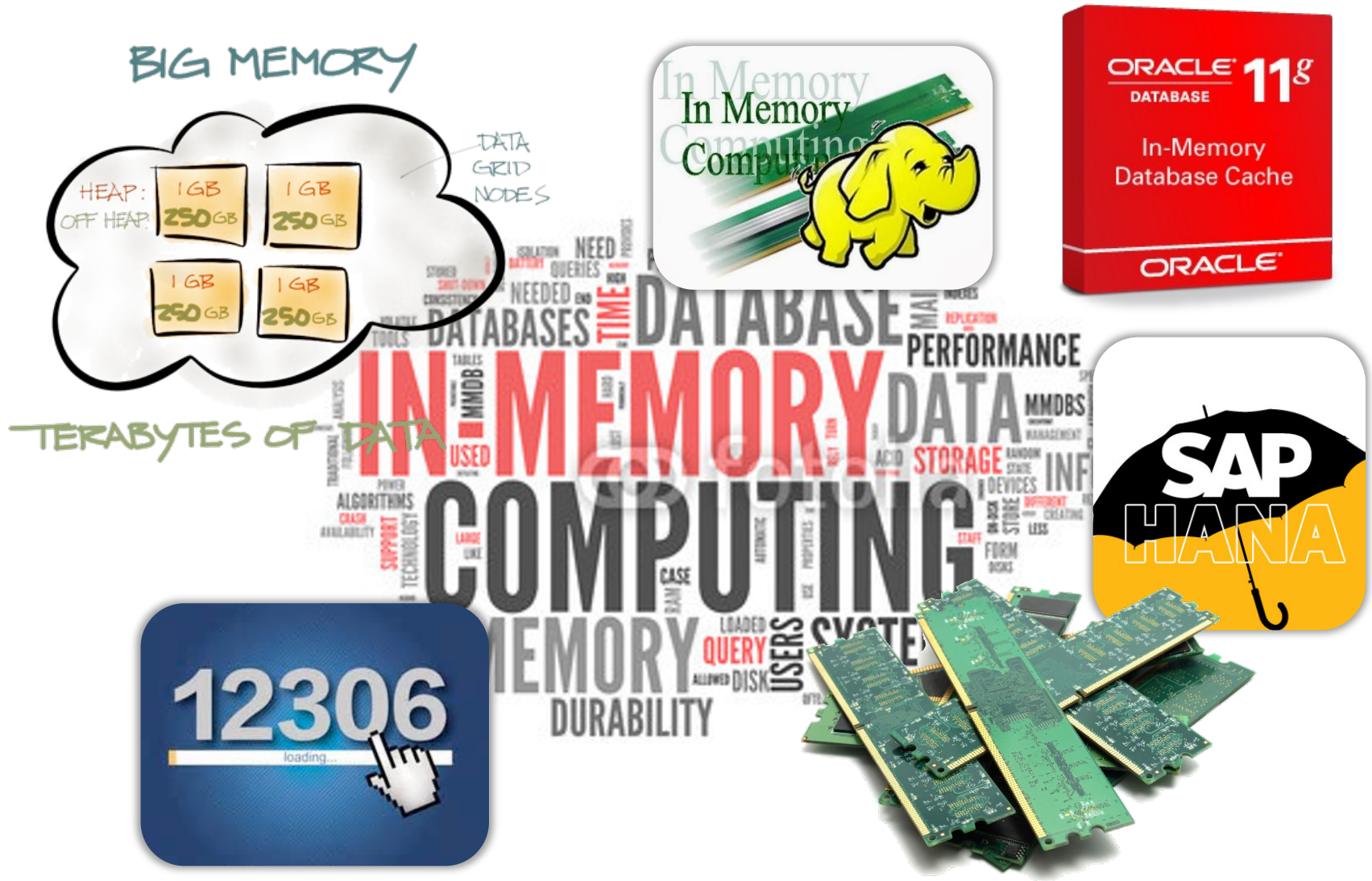
KV Store



High Availability with Replication



In-Memory Computing



HTM: Hardware Transactional Memory

Herlihy
& Moss

1993

Sun
Rock

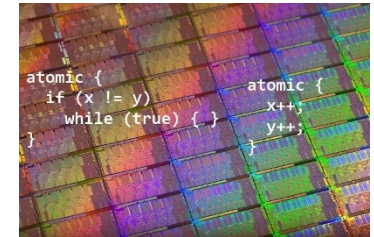
2009

AMD TX
Extension

2010

. . .

2012



Intel Haswell

2013

Massively available

Intel RTM

Restricted Transactional Memory (RTM)

- Commercial hardware TM with limitations

New Instructions:

Xbegin, Xend, Xabort

RTM Usage

```
if _xbegin() == _XBEGIN_STARTED
```

```
    do some critical work
```

```
    _xend()
```

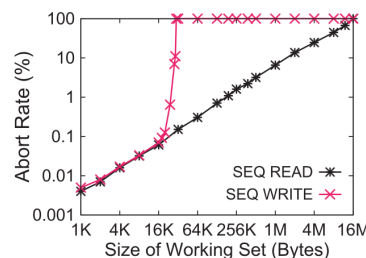
```
else
```

```
    fallback routine
```

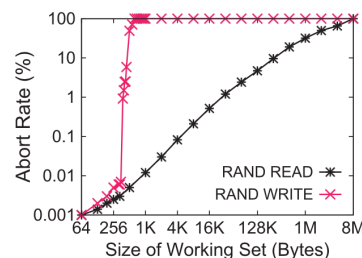
Handle the
abort event

Major Limitations

1. Working set is limited
2. System events abort TX



SEQ Access



RAND Access

Deconstructing **RTM**

RTM prefers read than write

- Asymmetric Read/Write Limits

RTM prefers read before write

- Eviction of cache lines in write set will abort TX

The execution time affects **TX abort**

- Timer interrupt unconditionally abort a TX
(e.g., 4ms on 250HZ kernel)

RDMA: Remote Direct Memory Access

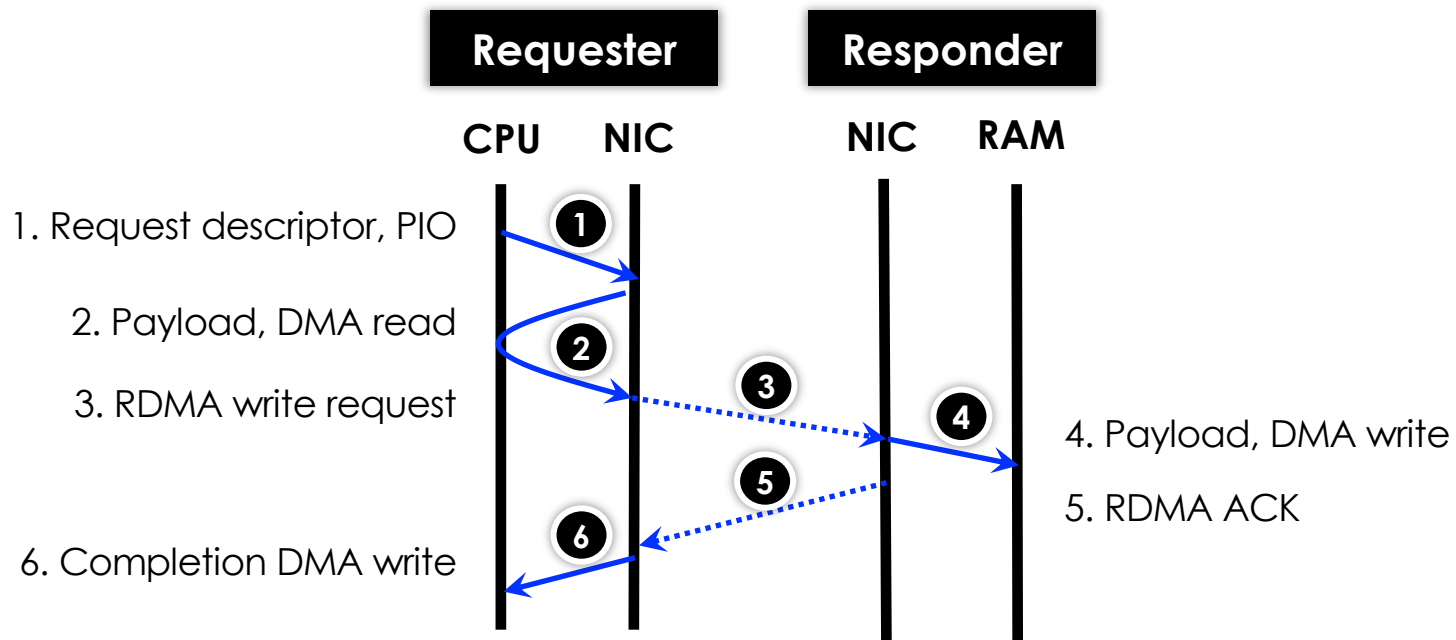
A network feature that allows **direct** access to the memory of a **remote** computer

High speed, low latency & low CPU overhead

- ▶ Interface: IPoB, two-sided **RDMA** (SEND/RECV), and **one-sided RDMA** (READ/WRITE/CAS)
- ▶ Bypasses **OS kernels**: Zero copy
Bypasses **CPU** (one-sided)
- ▶ Round-trip time: **one-sided**/ $\sim 3\mu\text{s}$,
two-sided/ $\sim 7\mu\text{s}$, IPoB/ $\sim 100\mu\text{s}$

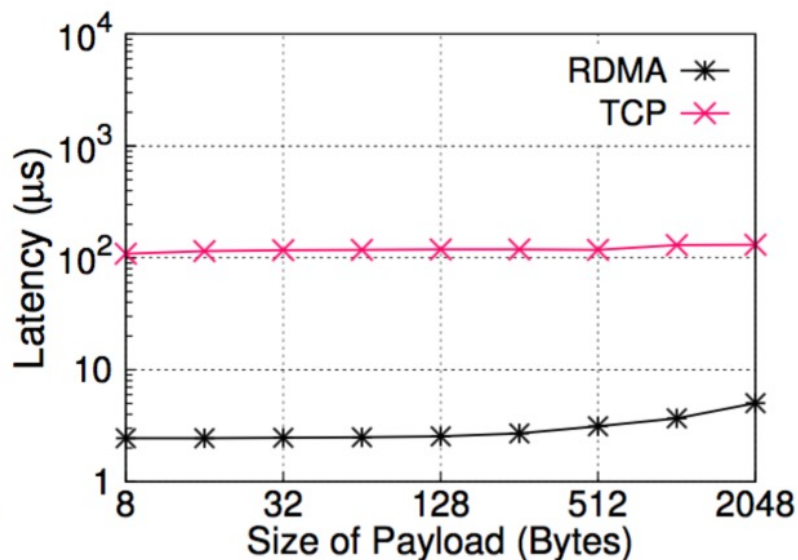
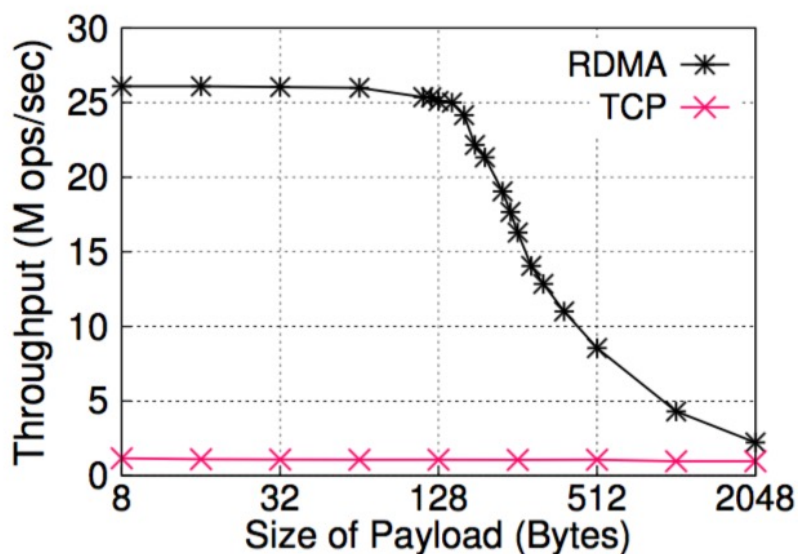
One-sided **RDMA** Primitives

For example: life-cycle of an **RDMA WRITE**



RDMA READ Performance

Performance of Random Read*



Insensitive to payload size: high/near constant throughput and low latency when payload is smaller than a threshold

This Course

101 of New Hardware Features

Hardware



Distributed Transaction Protocol



In-memory Key-Value Store

KV Store



High Availability with Replication



Opportunity of **HTM** and **RDMA**

HTM: Hardware Transaction Memory

a *non-transactional* code will unconditionally abort
a transaction when their accesses conflict

**Strong
Atomicity**

RDMA: Remote Direct Memory Access

Opportunity of **HTM** and **RDMA**

HTM: Hardware Transaction Memory

a *non-transactional* code will unconditionally abort
a transaction when their accesses conflict

**Strong
Atomicity**

RDMA: Remote Direct Memory Access

one-sided RDMA operations are *cache-coherent*
with local accesses

**Strong
Consistency**

Opportunity of **HTM** and **RDMA**

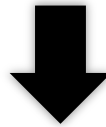
HTM: Hardware Transaction Memory

a *non-transactional* code will unconditionally abort a transaction when their accesses conflict

RDMA: Remote Direct Memory Access

one-sided RDMA operations are *cache-coherent* with local accesses

RDMA ops will abort conflicting **HTM** TX



Basis for **Distributed HTM**

Challenge#1: Restriction of **HTM**

HTM is only a compelling hardware feature for **single machine** platform

- **Distributed** TXs cannot directly benefit from it

Some instructions & system events (e.g. I/O) will **unconditionally** abort **HTM** transactions

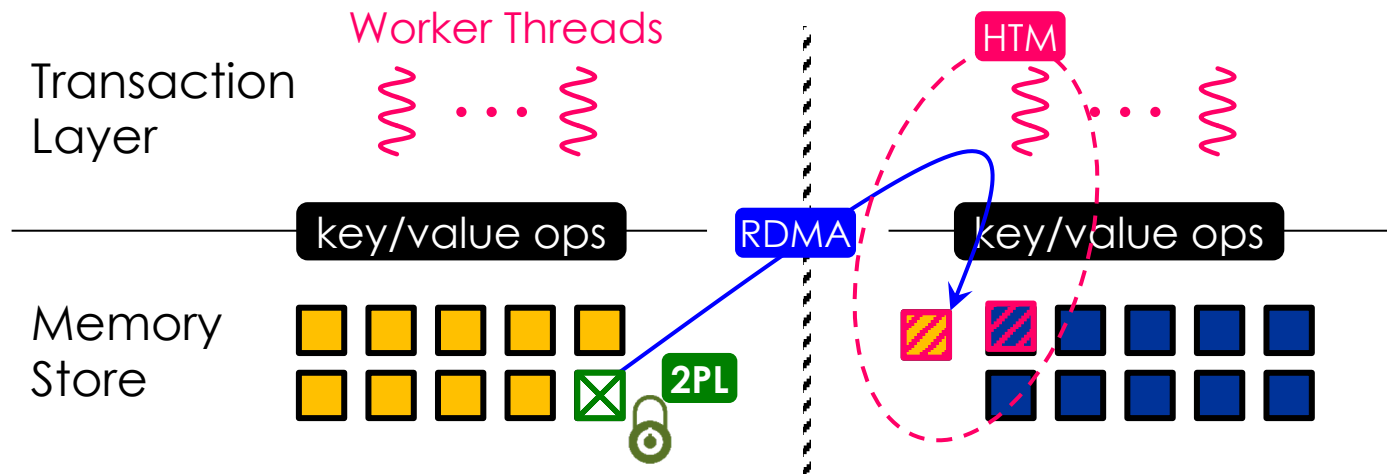
- Like any **RDMA** ops: READ/WRITE, CAS, . . .

How to glue multiple **HTM** transactions together by **RDMA** while preserving serializability?

Combining **HTM** with 2PL

Using **2PL** to accumulate all remote records prior to accesses in an **HTM** transaction

- Transform a **distributed** TX to a **local** one
- **Limitation:** require advanced knowledge of **read/write sets** of transactions*



* similar to prior work (e.g. Sinfonia & Calvin) and the case for typical OLTP workloads

Concurrency Control Protocol

Local TX vs. Local TX: **HTM**

Distributed TX vs. Distributed TX: **2PL**

Local TX vs. Distributed TX: **abort local TX**

Concurrency Control Protocol

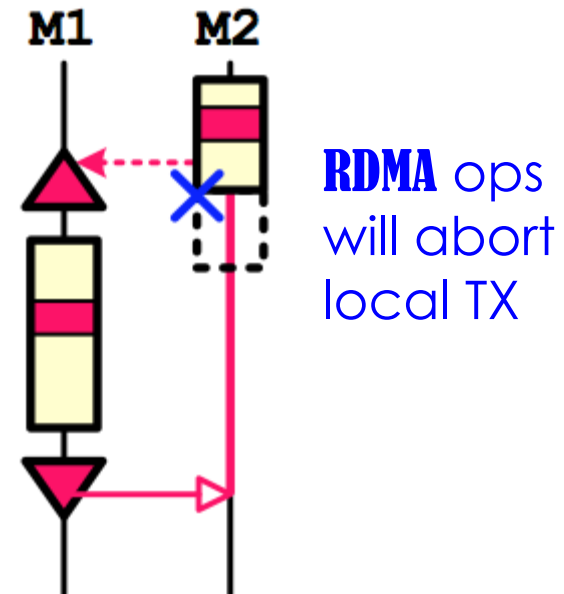
Local TX vs. Local TX: **HTM**

Distributed TX vs. Distributed TX: **2PL**

Local TX vs. Distributed TX: **abort local TX**

Local TX prior to Dist. TX

→ **RDMA** ops abort
local **HTM** tx



Concurrency Control Protocol

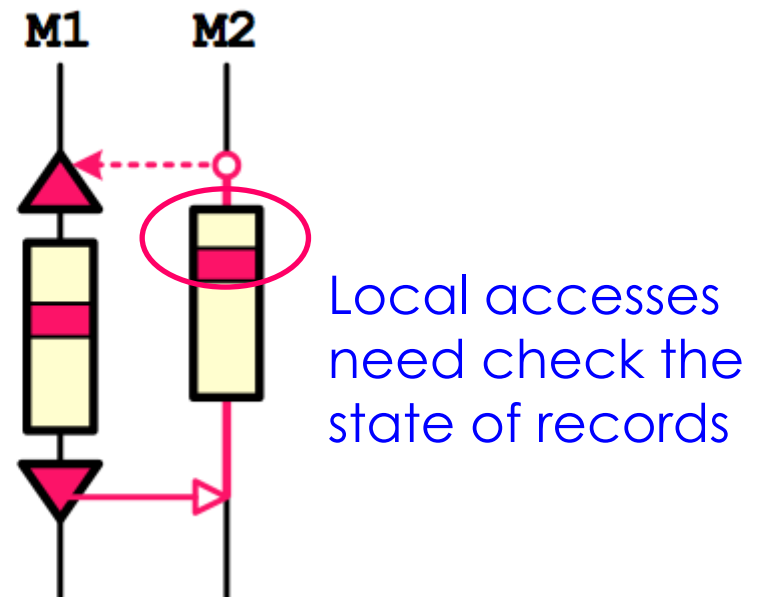
Local TX vs. Local TX: **HTM**

Distributed TX vs. Distributed TX: **2PL**

Local TX vs. Distributed TX: **abort local TX**

Dist. TX prior to local TX

→ Local **HTM** TX checks
(not locks) records



Challenge#2: Limit of **RDMA** Semantics

RDMA provides three communication options

- IPoIB, SEND/RECV and **one-sided RDMA ops**

Good performance (e.g. latency)
and without involving the host CPU

One-sided RDMA has much limited interfaces

- READ, WRITE, CAS and XADD

How to support **exclusive** and **shared** accesses
in 2PL protocol using one-sided **RDMA** ops?

RDMA-based Lock

RDMA CAS: atomic compare-and-swap

- Similar to the semantic of normal CAS (i.e., local CAS)

1. exclusive lock

- ▶ Spinlock: use **RDMA** CAS to **acquire** & **release**

2. shared lock

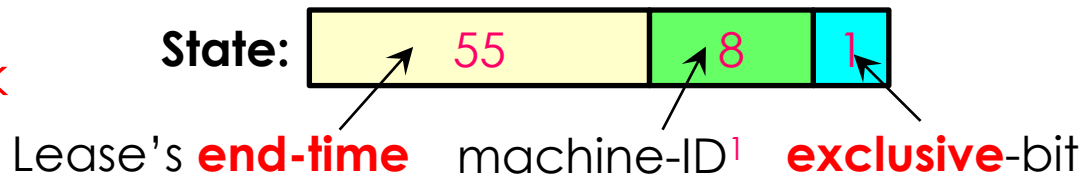
- ▶ **Lease-based** protocol

Shared (Read) Lock

Lease-based protocol

- Grant **read right** to the lock holder in a **time period**
- No need to **explicit** release or invalidate the lock

exclusive &
shared lock



`000...0002` **unlocked**
`000...yy12` **exclusive locked**
`xxx...0002` **shared locked**

State is atomically compare and swap using RDMA CAS

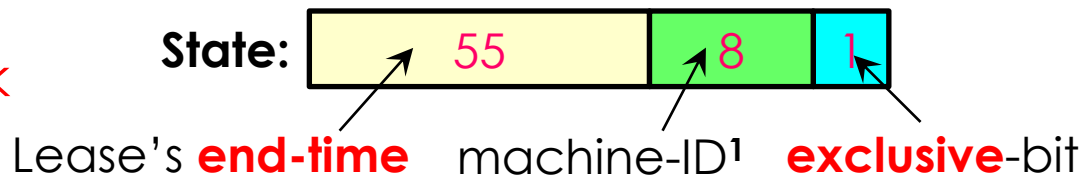
¹ Machine ID is only used by recovery

Shared (Read) Lock

Lease-based protocol

- Grant **read right** to the lock holder in a **time period**
- No need to **explicit** release or invalidate the lock
- Synchronized time is provided by PTP²

exclusive &
shared lock



000...000₂ **unlocked**
000...yy1₂ **exclusive locked**
xxx...000₂ **shared locked**

DELTA is used to tolerate the time bias among machines

→ **EXPIRED:** if **now** > end-time + **DELTA**

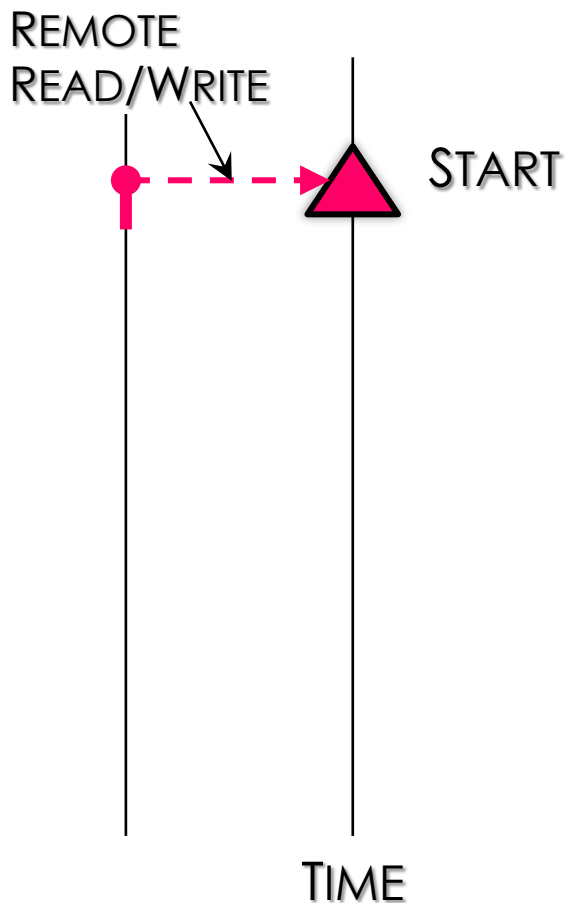
INVALID: if **now** < end-time - **DELTA**

¹ Machine ID is only used by recovery

² PTP: precision time protocol, <http://sourceforge.net/p/ptpd/wiki/Home/>

Transaction Execution Flow

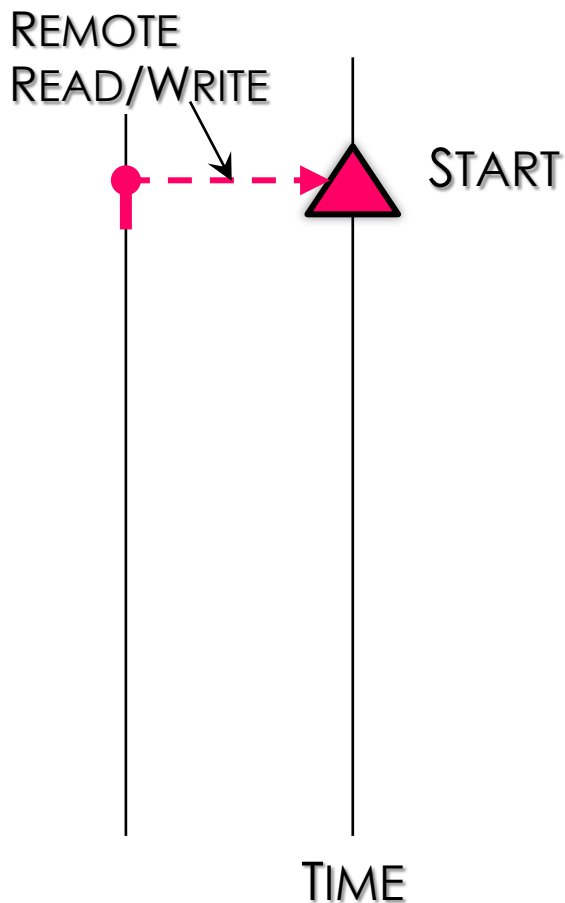
Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



```
START (remote_writeset, remote_readset)  
  foreach key in remote_writeset  
    value = Exclusive_lock_fetch(key)  
    cache[key] = value  
  foreach key in remote_readset  
    value = Shared_lease_fetch(key)  
    cache[key] = value  
  
XBEGIN()
```

Transaction Execution Flow

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**

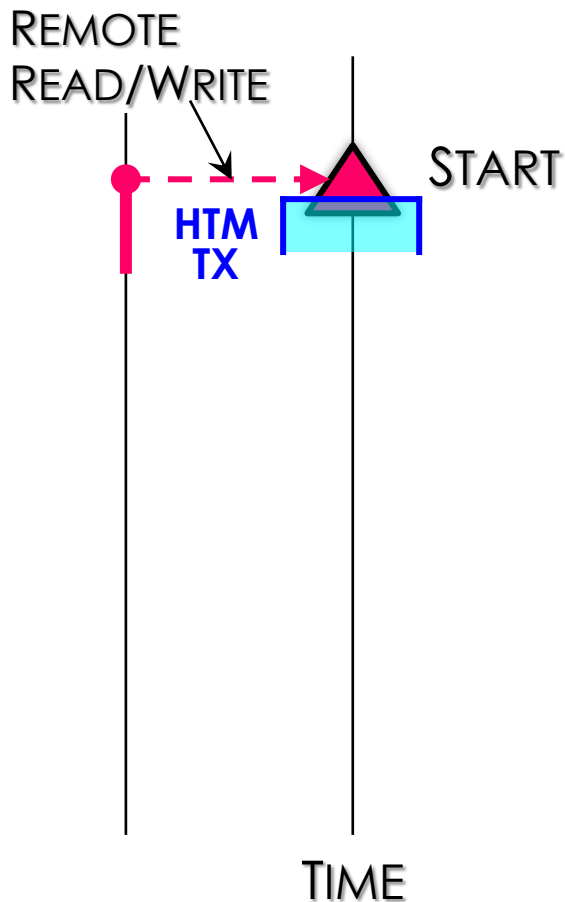


```
START(remote_writeset, remote_readset)
  foreach key in remote_writeset
    value = Exclusive_lock_fetch(key)
    cache[key] = value
  foreach key in remote_readset
    value = Shared_lease_fetch(key)
    cache[key] = value

XBEGIN()
```

Transaction Execution Flow

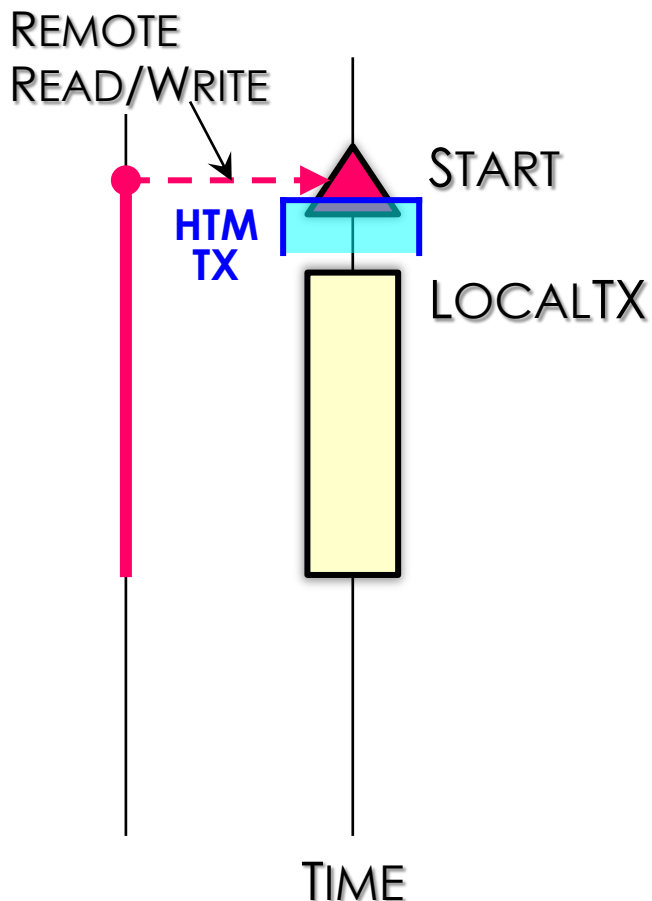
Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



```
START(remote_writeset, remote_readset)
  foreach key in remote_writeset
    value = Exclusive_lock_fetch(key)
    cache[key] = value
  foreach key in remote_readset
    value = Shared_lease_fetch(key)
    cache[key] = value
XBEGIN()
```

Transactional Read & Write

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



READ(key)

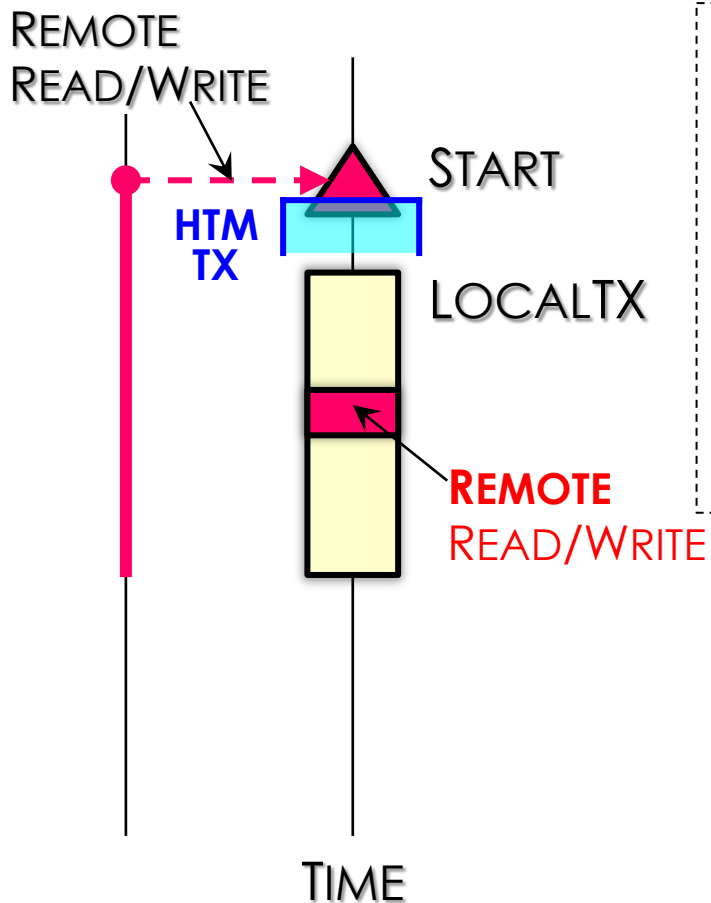
```
if key.is_remote() == true
    return cache[key]
else return LOCAL_READ(key)
```

WRITE(key, value)

```
if key.is_remote() == true
    cache[key] = value
else LOCAL_WRITE(key, value)
```

Transactional Read & Write

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**

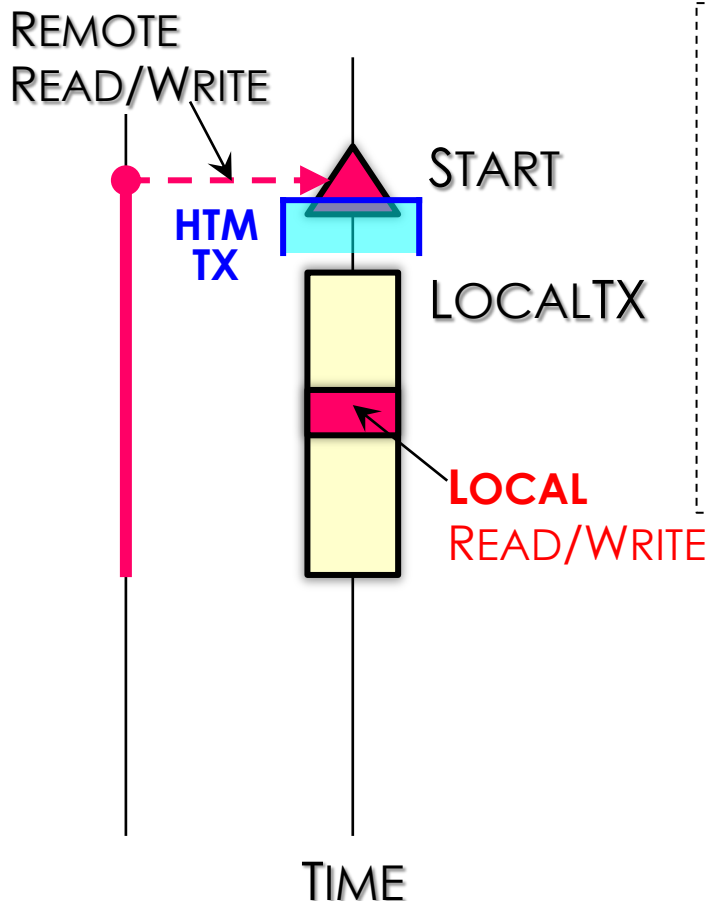


```
READ(key)
    if key.is_remote() == true
        return cache[key]
    else return LOCAL_READ(key)

WRITE(key, value)
    if key.is_remote() == true
        cache[key] = value
    else LOCAL_WRITE(key, value)
```

Transactional Read & Write

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**

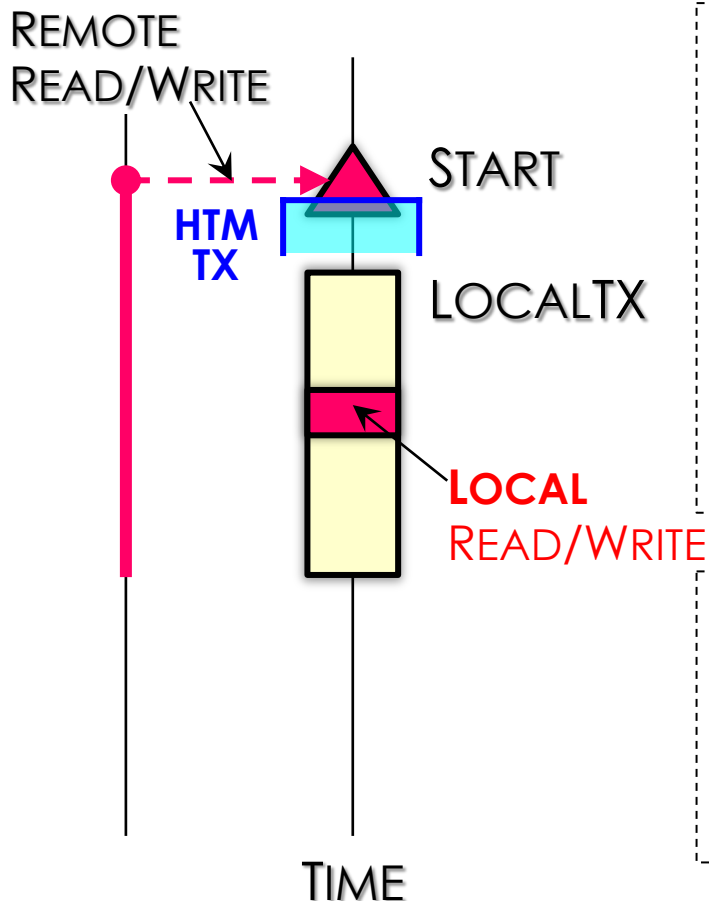


```
READ (key)
  if key.is_remote() == true
    return cache[key]
  else return LOCAL_READ(key)

WRITE (key, value)
  if key.is_remote() == true
    cache[key] = value
  else LOCAL_WRITE(key, value)
```

Transactional Read & Write

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



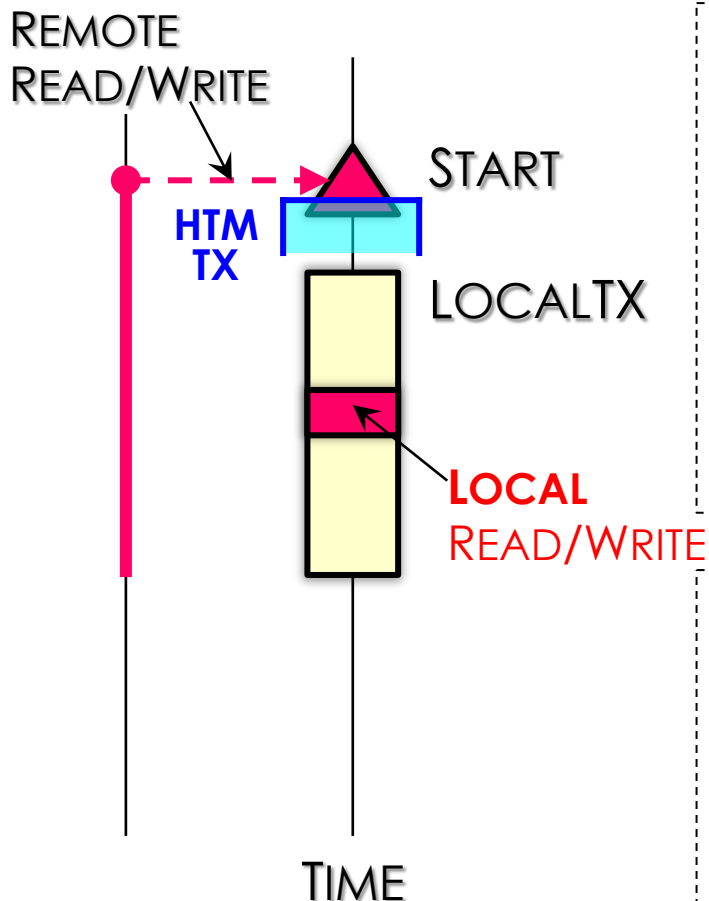
```
READ(key)
    if key.is_remote() == true
        return cache[key]
    else return LOCAL_READ(key)
```

```
WRITE(key, value)
    if key.is_remote() == true
        cache[key] = value
    else LOCAL_WRITE(key, value)
```

```
LOCAL_READ(key)
    if states[key].w_lock == W_LOCKED
        ABORT()
    else
        return values[key]
```

Transactional Read & Write

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



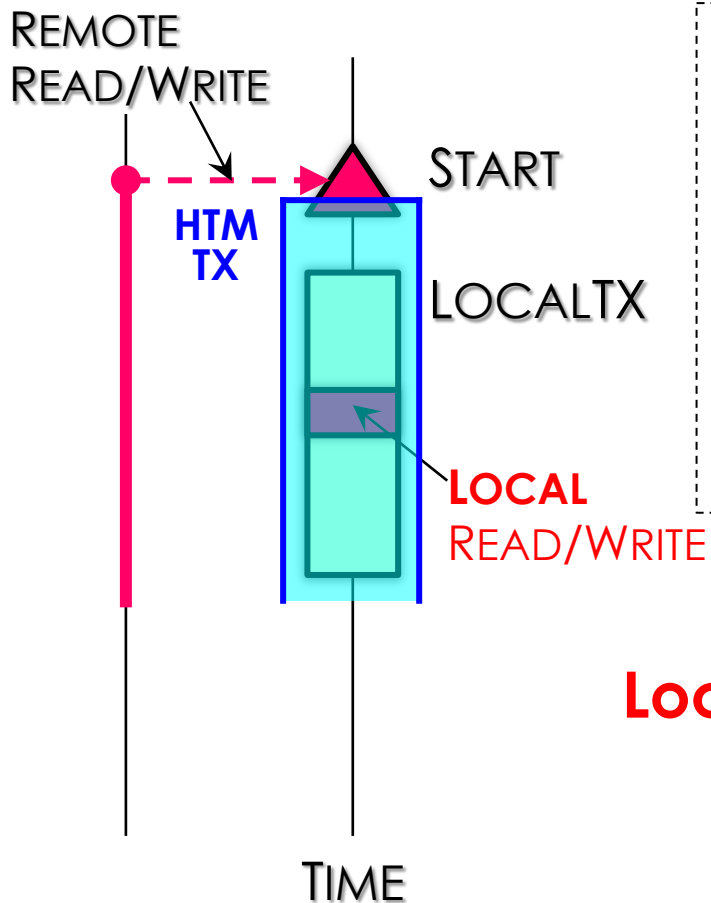
```
READ(key)
    if key.is_remote() == true
        return cache[key]
    else return LOCAL_READ(key)
```

```
WRITE(key, value)
    if key.is_remote() == true
        cache[key] = value
    else LOCAL_WRITE(key, value)
```

```
LOCAL_WRITE(key, value)
    if states[key].w_lock == W_LOCKED
        ABORT()
    else if EXPIRED(END_TIME(states[key]))
        values[key] = value
    else ABORT()
```

Transactional Read & Write

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



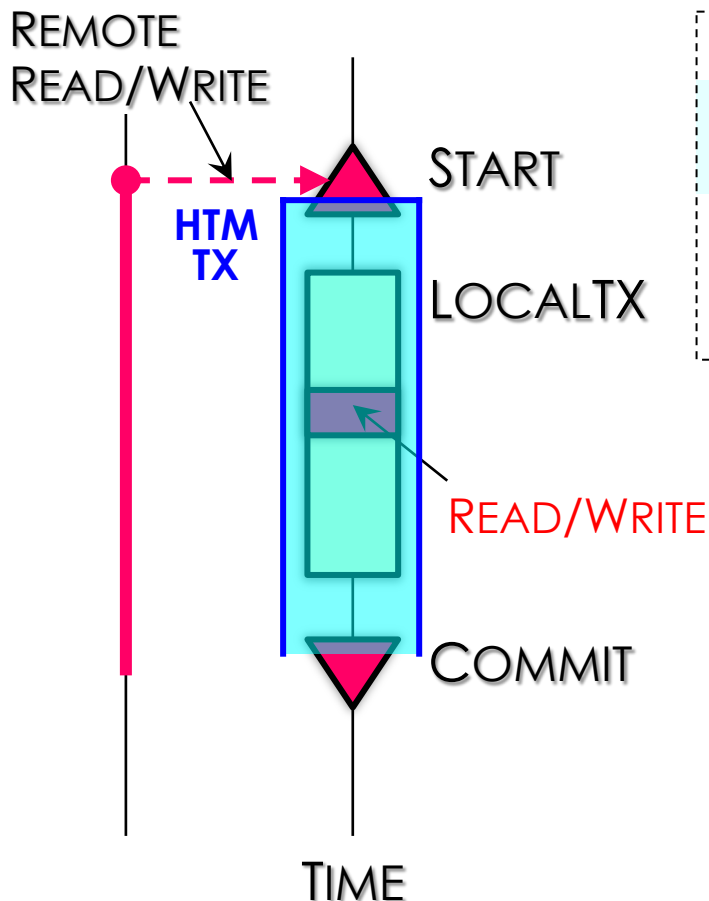
```
READ(key)
  if key.is_remote() == true
    return cache[key]
  else return LOCAL_READ(key)

WRITE(key, value)
  if key.is_remote() == true
    cache[key] = value
  else LOCAL_WRITE(key, value)
```

Local conflicts are detected by **HTM**

Transaction Execution Flow

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



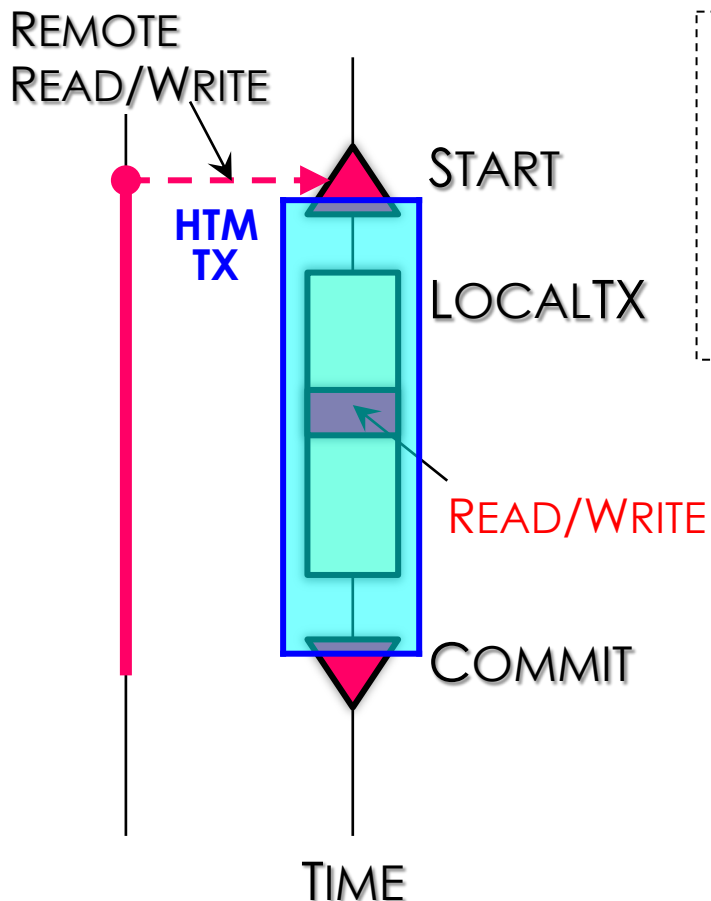
```
COMMIT(remote_writeset, remote_readset)
  if !VALID(end_time)
    ABORT()
  XEND()
  foreach key in remote_writeset
    RELEASE_WRITE_BACK(key, cache[key])
```

2PL: all **shared locks** must be released in **shrinking phase**

- Insert validation to all leases **just before HTM** commit

Transaction Execution Flow

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**

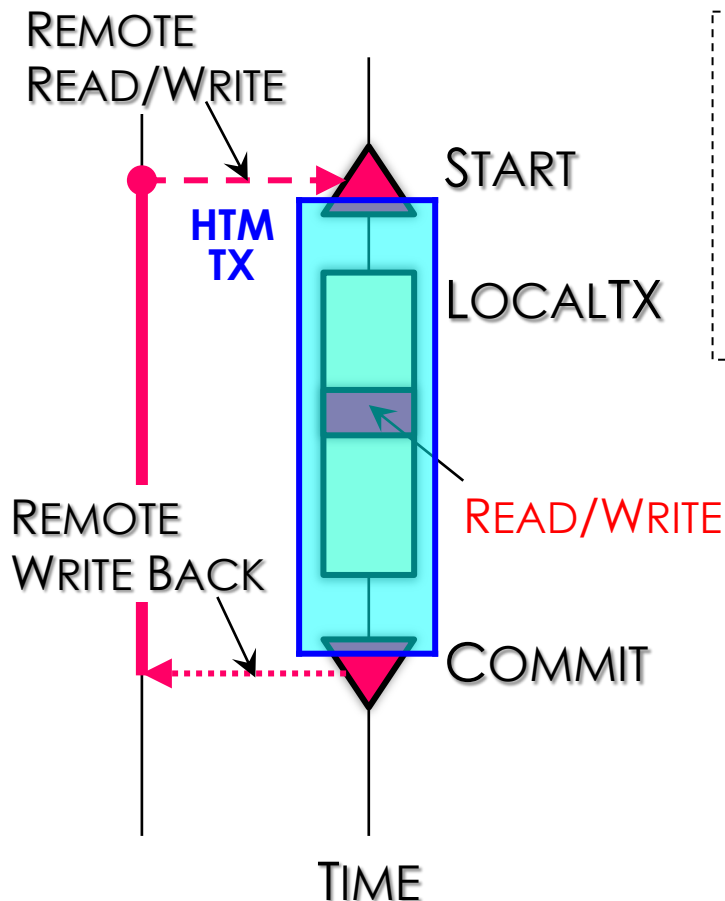


```
COMMIT(remote_writeset, remote_readset)
  if !VALID(end_time)
    ABORT()
  XEND()
  foreach key in remote_writeset
    RELEASE_WRITE_BACK(key, cache[key])
```

Commit **local** updates by **HTM**

Transaction Execution Flow

Transaction Protocol: **START** + **LOCALTX** + **COMMIT**



```
COMMIT(remote_writeset, remote_readset)
  if !VALID(end_time)
    ABORT()
  XEND()
  foreach key in remote_writeset
    RELEASE_WRITE_BACK(key, cache[key])
```

2PL & HTM → Serializability

Commit **local** updates by **HTM**
Commit **remote** updates by **RDMA**

Read-only Transaction

Read-only TX usually has a very large read set

- Likely abort an HTM transaction due to **capacity**

Two-round schema (read & confirm)

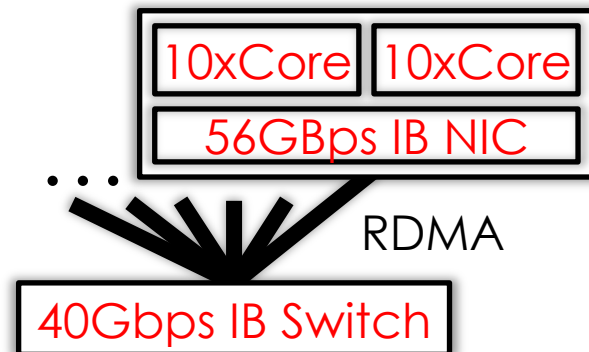
- Similar to ROCOCO @OSDI'14
- Using lease & uniform end-time to avoid twice read

```
START_RO(readset)
L: end_time = now + duration
  foreach key in readset
    end_time = MIN(end_time, Set_lease(key, end_time))
  if !VALID(end_time) goto L //round2: confirm

READ_RO(key)
  return read_cache[key]
```

Evaluation

Baseline: Latest *Calvin* (Mar. 2015)



Platforms: A small-scale 6-machine cluster

- Each: two 10-cores, **RTM-enabled** Intel Xeon E5-2650 (disabled HT), 64GB DRAM, Mellanox ConnectX-3 MCX353A 56Gbps InfiniBand NIC w/ **RDMA**¹

Benchmarks²

- TPC-C
- SmallBank

TPC-C	NEW	PAY	DLY	OS	SL
Ratio	45%	43%	4%	4%	4%
Type	d+rw	d+rw	l+rw	l+ro	l+ro

SmallBank	SP	AMG	BAL	DC	WC	TS
Ratio	25%	15%	15%	15%	15%	15%
Type	d+rw	d+rw	l+ro	l+rw	l+rw	l+rw

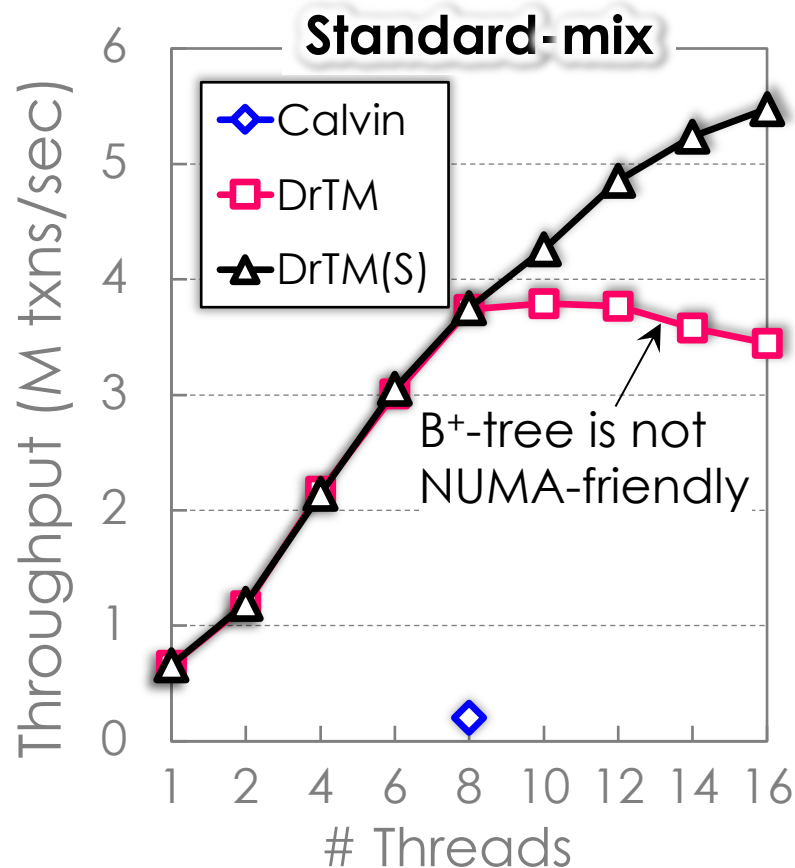
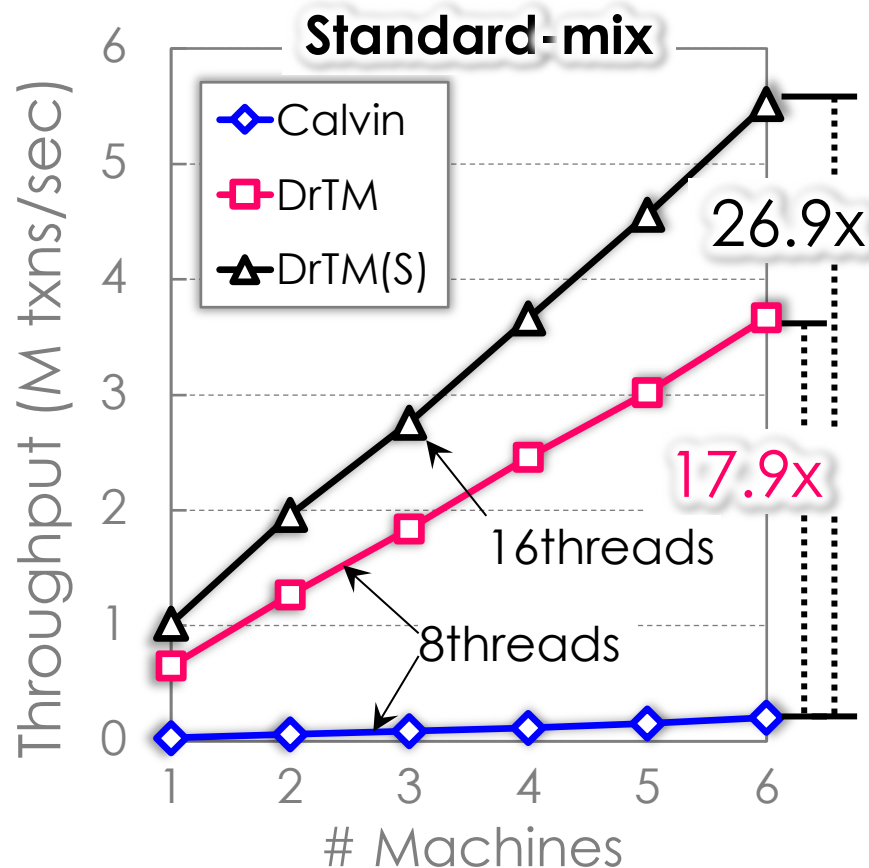
¹ All machines run Ubuntu 14.04 with Mellanox OFED v3.0-2.0.1 stack.

² **d** and **l** stand for distributed and local. **rw** and **ro** stand for read-write and read-only.

Performance on TPC-C

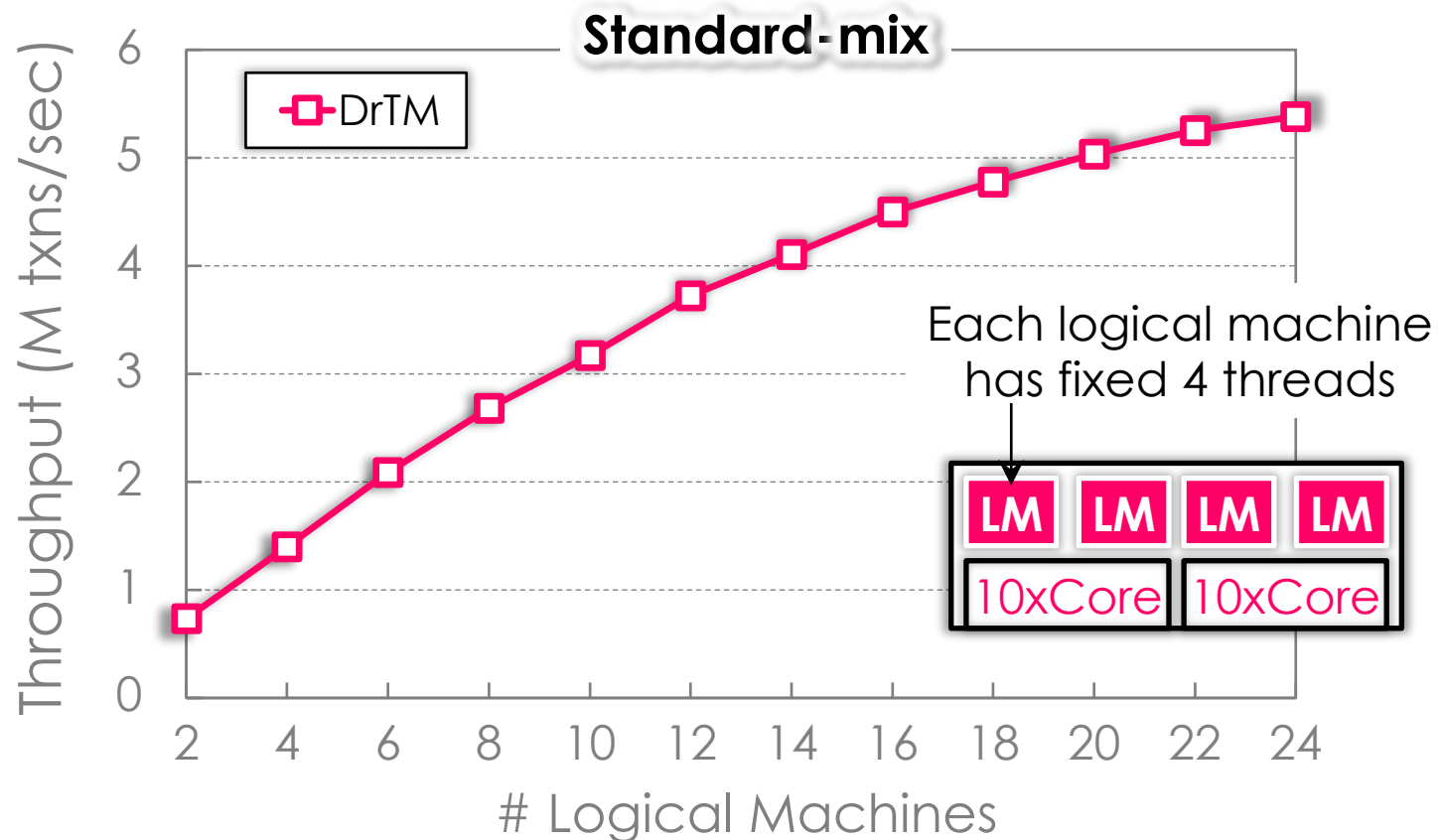
New-order TX
≈ Standard-mix x45%

DrTM(S): run a separate **logical** node on each socket



Scalability on TPC-C

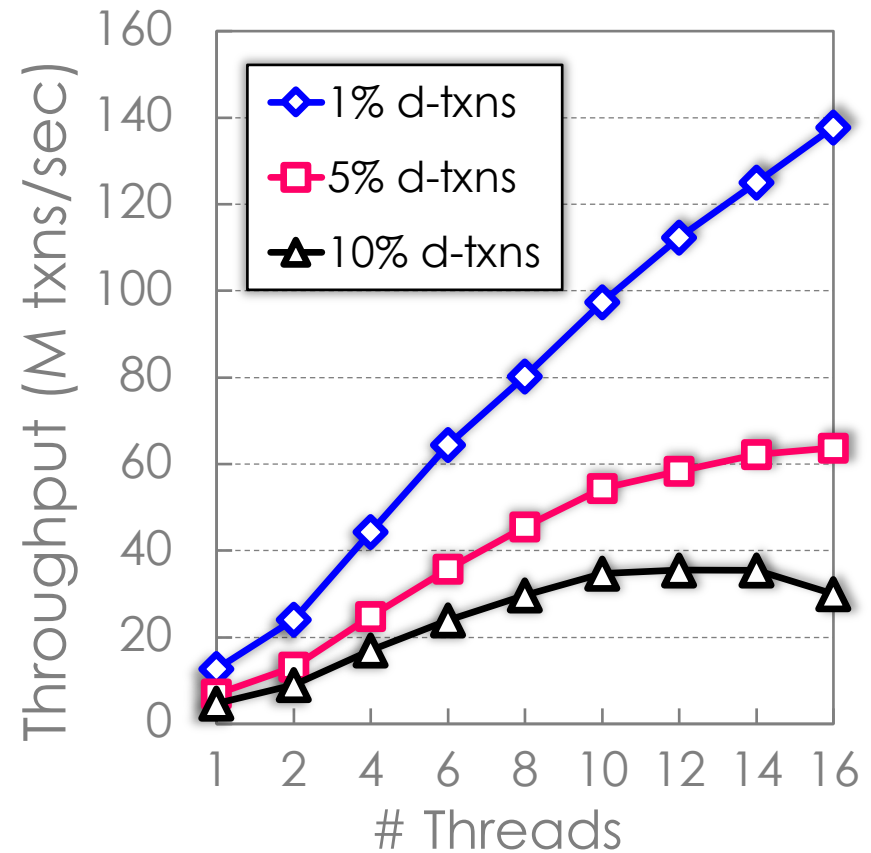
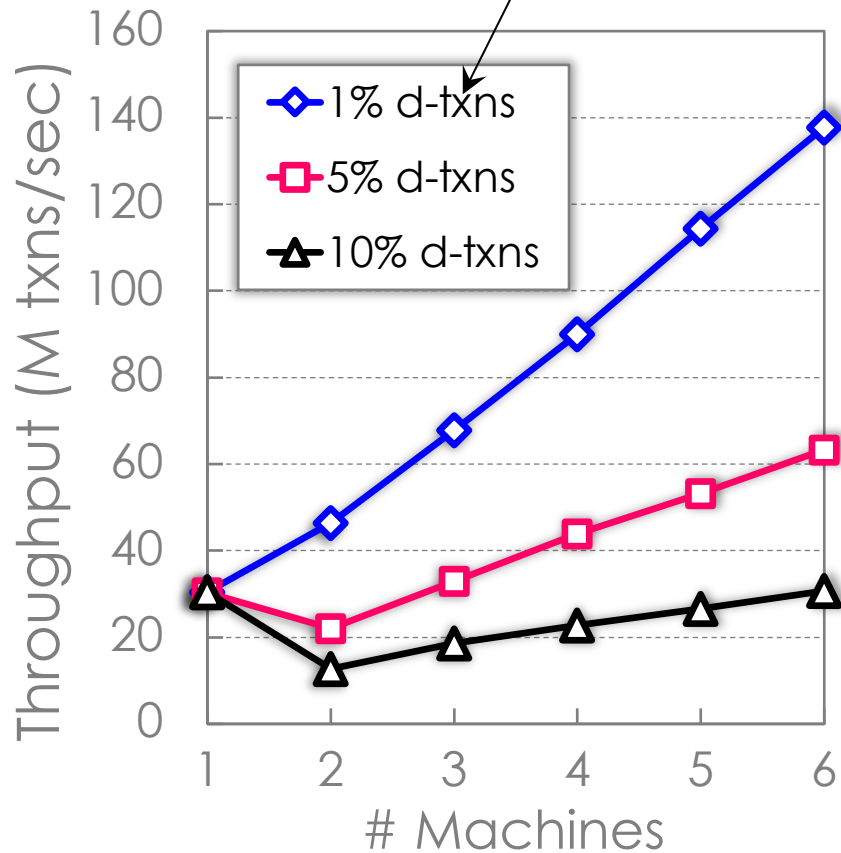
New-order TX
≈ Standard-mix x45%



NOTE: the **interaction** btw. two logical nodes sharing the same machine still uses our **RDMA-based 2PL protocol**

Performance on SmallBank

The probability of distributed transactions



Extended Reading

Transaction chopping: reduce **HTM** working set

Fine-grained RTM's fallback handler

Atomicity Issues: **RDMA** CAS vs. Local CAS

Horizontal scaling across socket: logical node

Reference

1. *Fast In-memory Transaction Processing using RDMA and HTM*. SOSP 2015
2. *Fast In-memory Transaction Processing using RDMA and HTM*. ACM TOCS.

This Course

101 of New Hardware Features

Hardware



Distributed Transaction Protocol



In-memory Key-Value Store

KV Store



High Availability with Replication



Memory Store

Separating **ordered** and **unordered** store

- Ordered store: B⁺ tree from DBX @EuroSys'14

No inevitable **remote accesses** to ordered stores in our OLTP workloads (i.e. TPC-C & SmallBank)

- Unordered store: **RDMA/HTM-based** hash table

Architecture

- **Symmetric**: each node is both a server and a client
- Most memory accesses are local with **HTM**

Existing Approach

Prior systems (e.g. Pilaf @ATC'13 and FaRM @NSDI'14)

- Complicated INSERT: hard to leverage **HTM**
- Only leverage one-sided **RDMA** to read
- No **RDMA**-based caching mechanism

Content-based caching (e.g. replication) is hard to perform strong-consistent read and write locally, especially using **RDMA**

	Pilaf	FaRM
Hashing	Cuckoo	Hopscotch
Race Detection	Checksum	Versioning
Remote Read	One-sided RDMA	
Remote Write	Messaging	
Caching	No	

RDMA & HTM provides a new design space

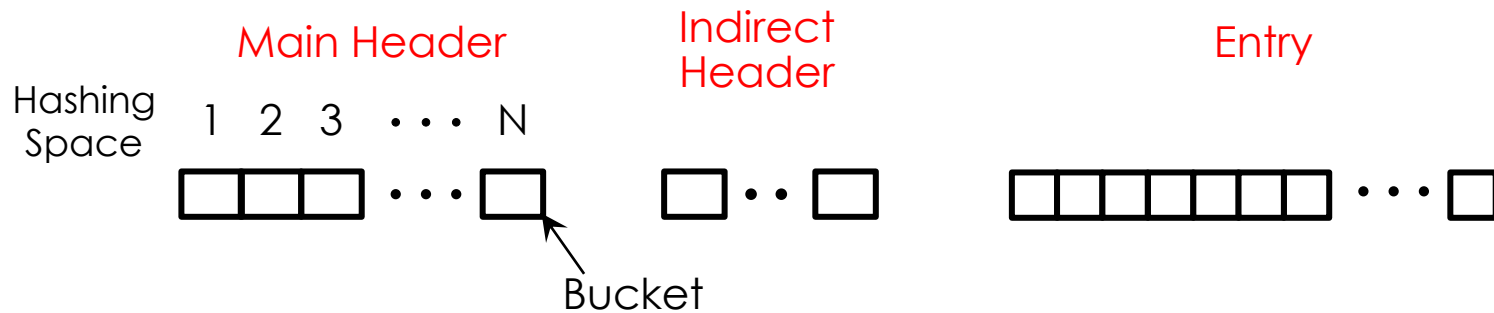
General Design

- ✓ Simple hash structure to fully leverage **HTM**
- ✓ **Decouple** race detection from memory store
 - Rely on transaction layer (**HTM** & Locking)
 - Use **one-sided RDMA** ops for remote read & **write**
- ✓ **Location-based** and fully transparent cache

	Pilaf		FaRM	DrTM	
Hashing	Cuckoo	Hopscotch	Chaining	Simple	
Race Detection	Checksum	Versioning	L:HTM / D: Lock		
Remote Read	One-sided RDMA		One-sided RMDA	Efficient	
Remote Write	Messaging				
Caching	No		Yes		

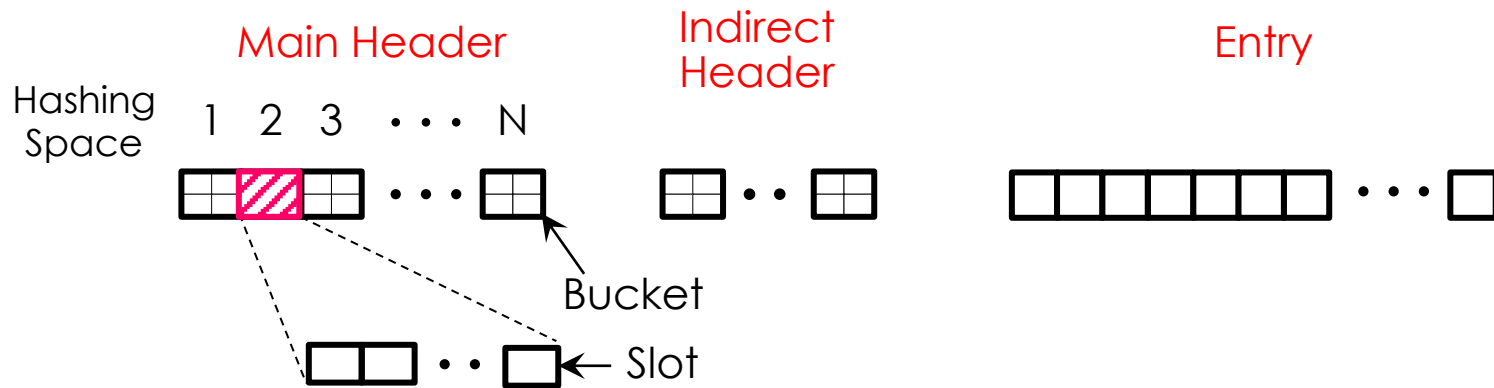
Cluster Chaining

- ✓ Decoupled memory region: **index** & **data**



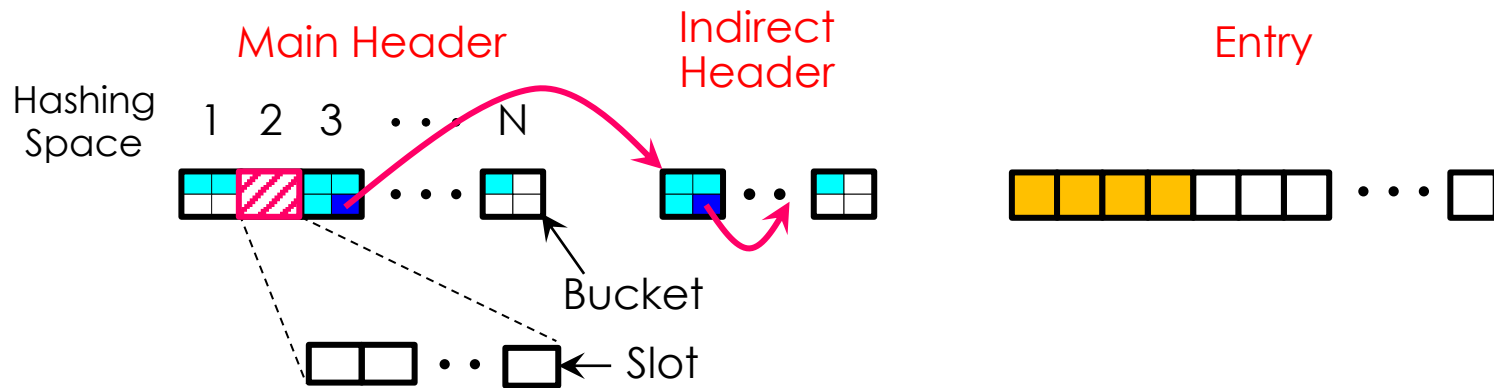
Cluster Chaining

- ✓ Decoupled memory region: **index** & **data**
- ✓ Similar to traditional **chaining** HT with **associativity**



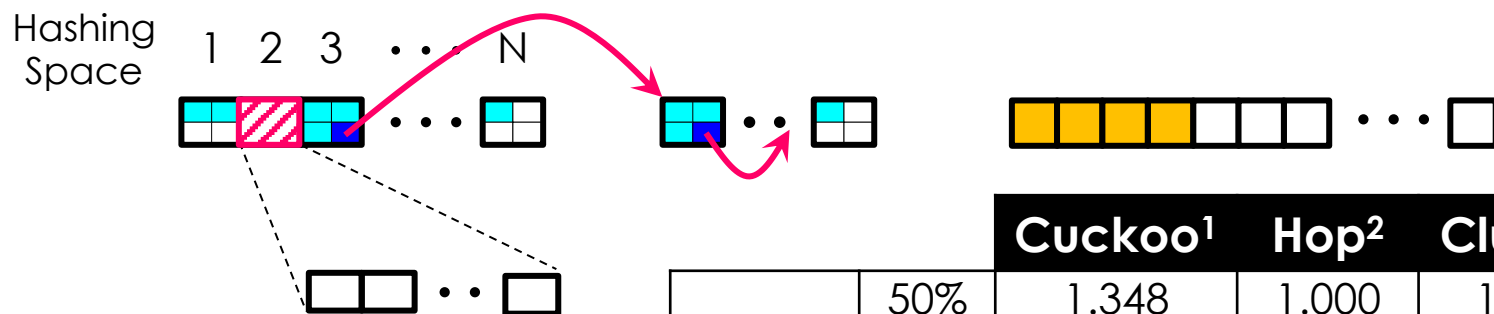
Cluster Chaining

- ✓ Decoupled memory region: **index** & **data**
- ✓ Similar to traditional **chaining** HT with **associativity**
- ✓ **Shared** indirect headers: high **space** efficiency



Cluster Chaining

- ✓ Decoupled memory region: **index** & **data**
- ✓ Similar to traditional **chaining** HT with **associativity**
- ✓ **Shared** indirect headers: high **space** efficiency



		Cuckoo ¹	Hop ²	Cluster ³
Uniform	50%	1.348	1.000	1.008
	75%	1.652	1.011	1.052
	90%	1.956	1.044	1.100
Zipf $\theta=0.99$	50%	1.304	1.000	1.004
	75%	1.712	1.020	1.039
	90%	1.924	1.040	1.091

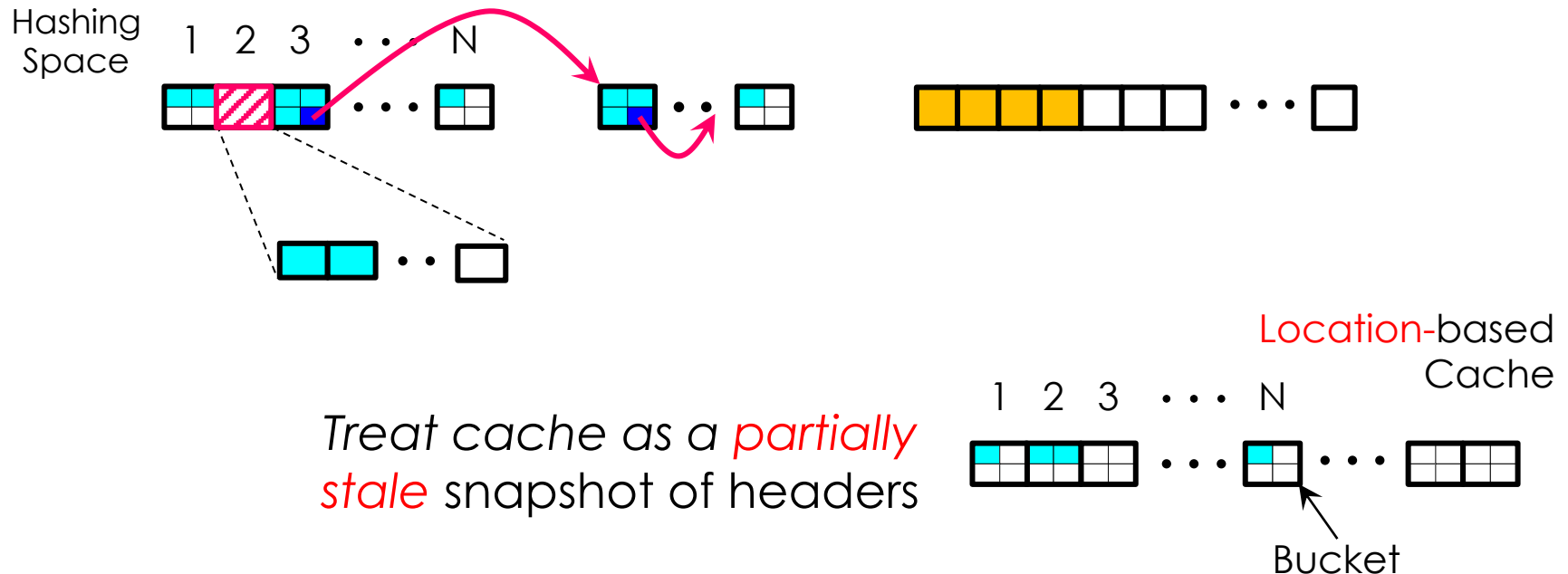
¹ Hopscotch hashing in FaRM configures the neighborhood with 8 ($H=8$).

² Cuckoo hashing in Pilaf uses 3 orthogonal hash functions and each bucket contains 1 slot.

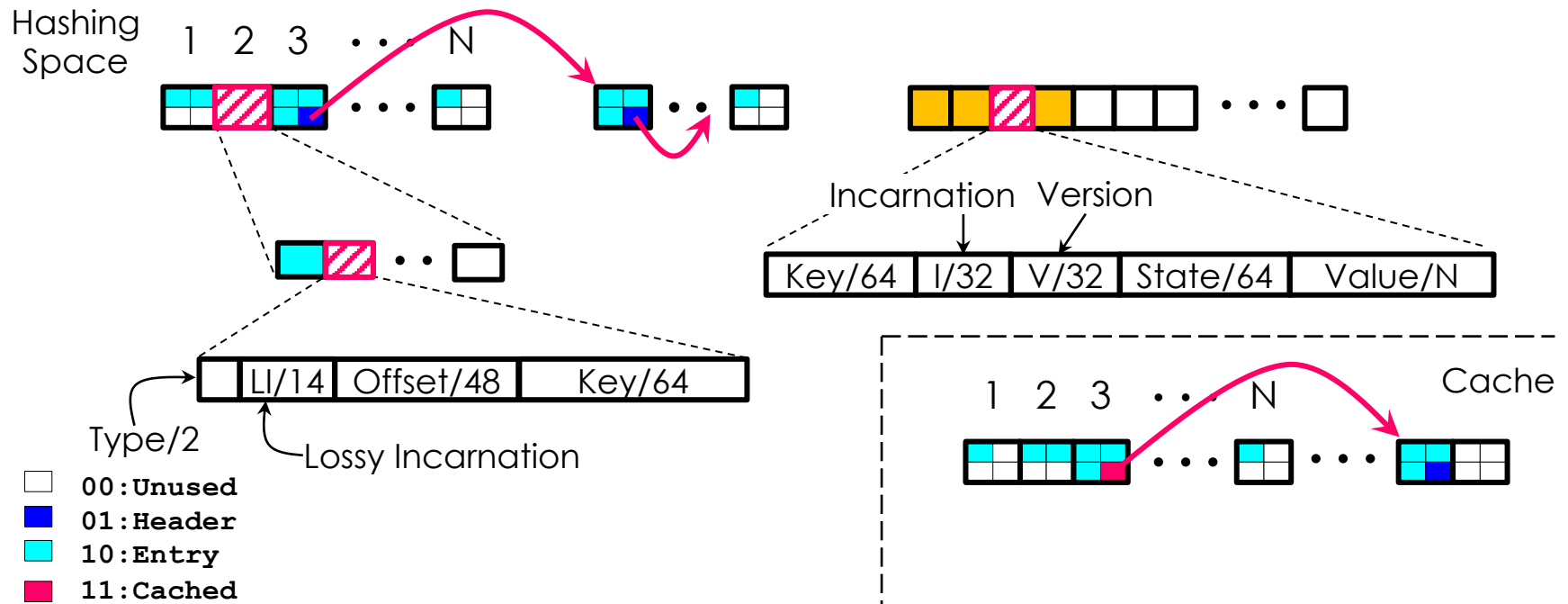
³ Cluster hashing in DrTM configures the associativity with 8.

Location-based Caching

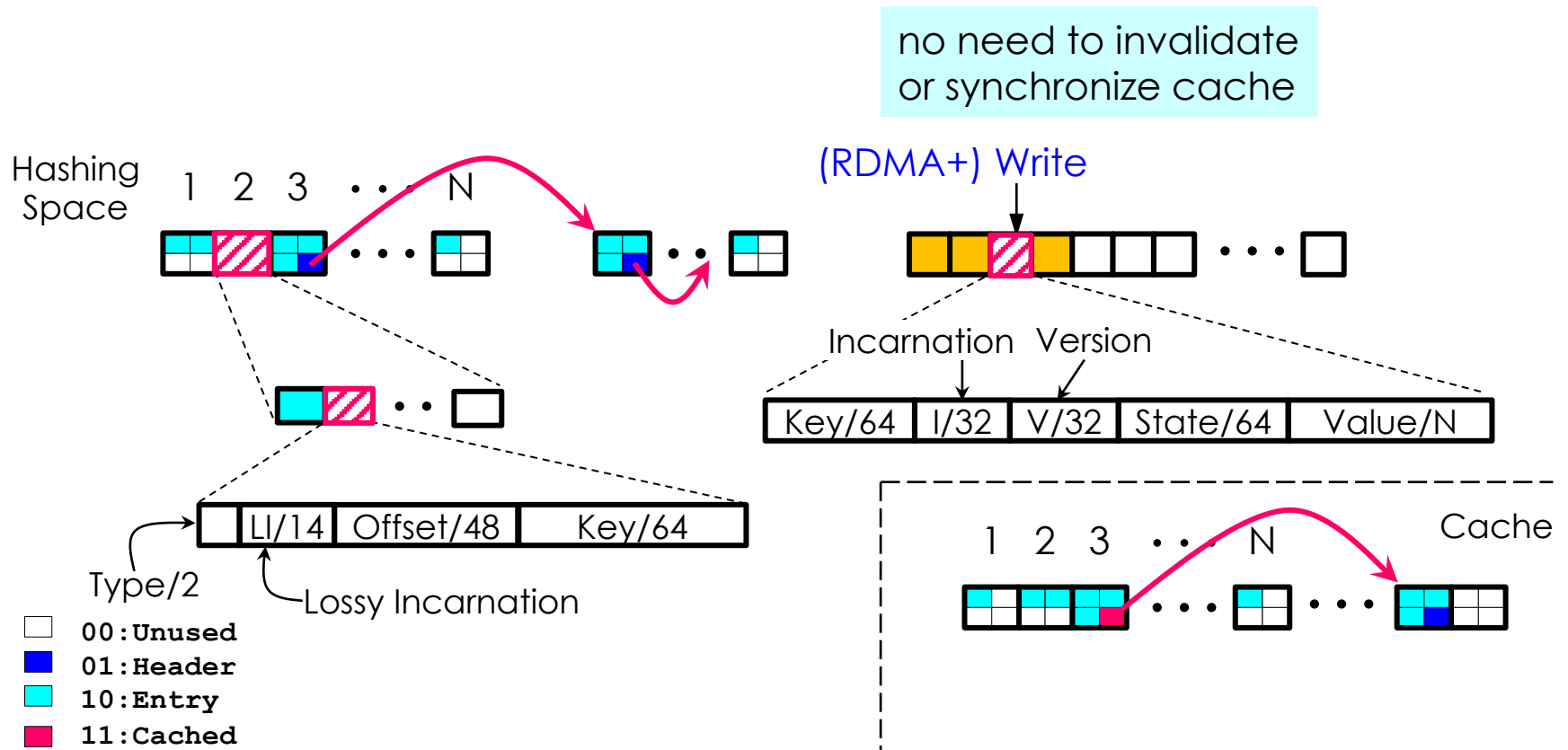
- ✓ **RDMA-based**: focus on minimizing the lookup cost



- ✓ **RDMA-based**: focus on minimizing the lookup cost
- ✓ Retain the full **transparency** to the host
 - All metadata used by concurrency control mechanisms are encoded in the key-value **entry**

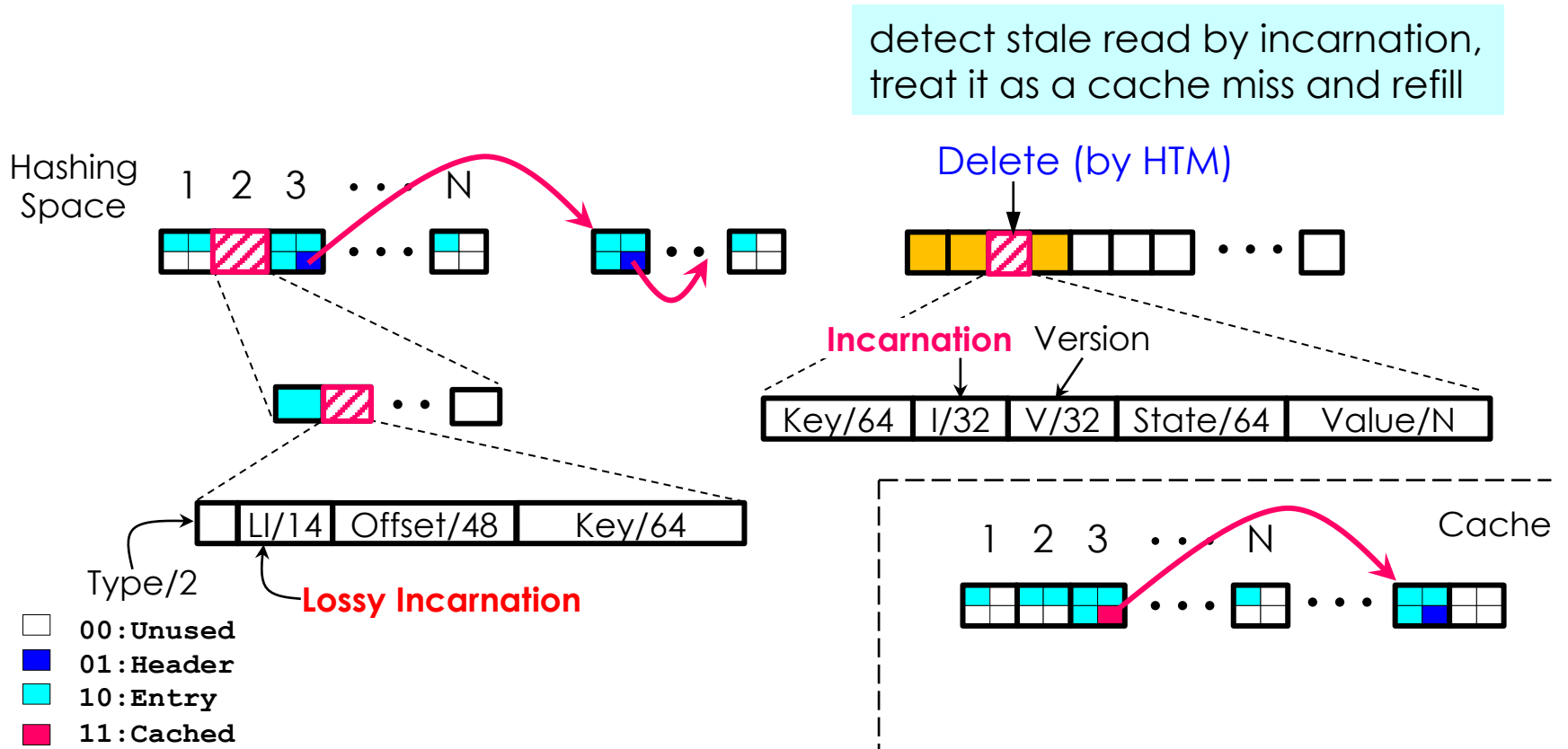


- ✓ **RDMA-based**: focus on minimizing the lookup cost
- ✓ Retain the full **transparency** to the host



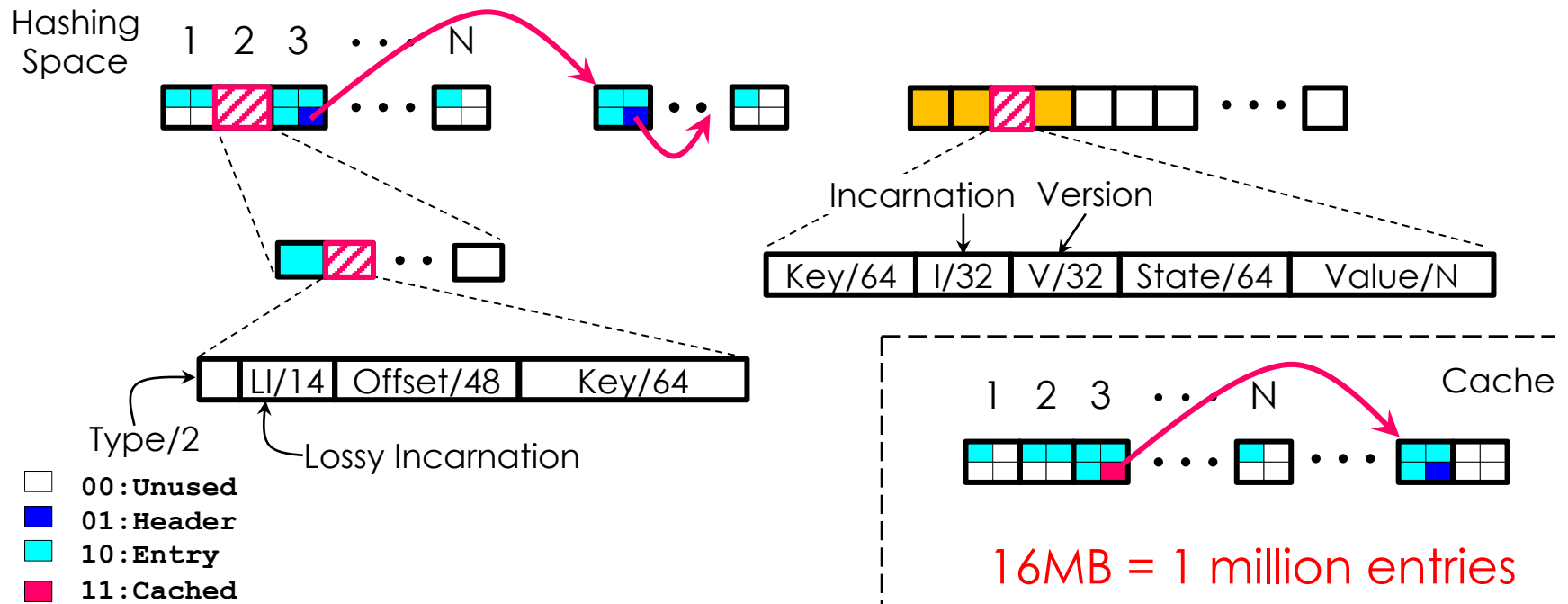
Abstract

- ✓ **RDMA-based**: focus on minimizing the lookup cost
- ✓ Retain the full **transparency** to the host



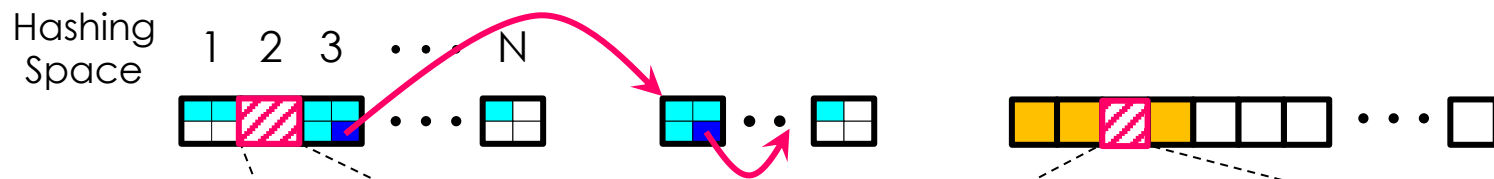
Location-based Caching

- ✓ **RDMA-based**: focus on minimizing the lookup cost
- ✓ Retain the full **transparency** to the host
- ✓ The size of cache for location is **small**



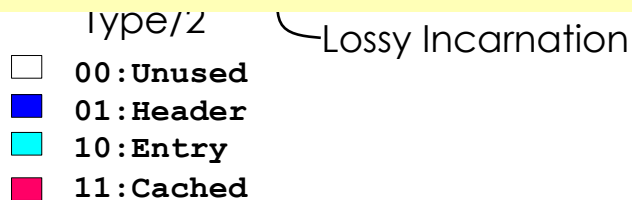
Location-based Caching

- ✓ **RDMA-based**: focus on minimizing the lookup cost
- ✓ Retain the full **transparency** to the host
- ✓ The size of cache for location is **small**
- ✓ All client threads can directly **share** the cache



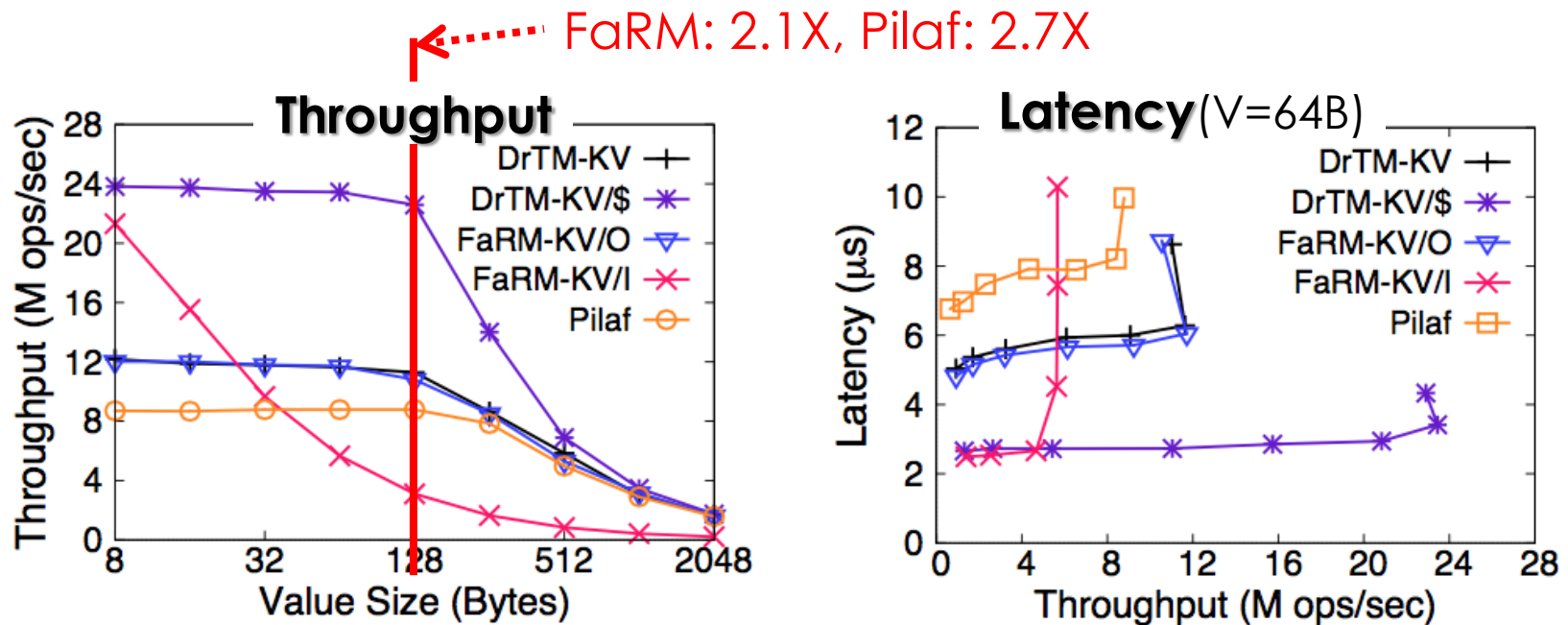
The average lookup cost = **0.178**

20 million kv-pairs (>40 GB), **20MB** cache (from **empty**),
8 client threads, skewed workload (**Zipf** $\theta=0.99$)



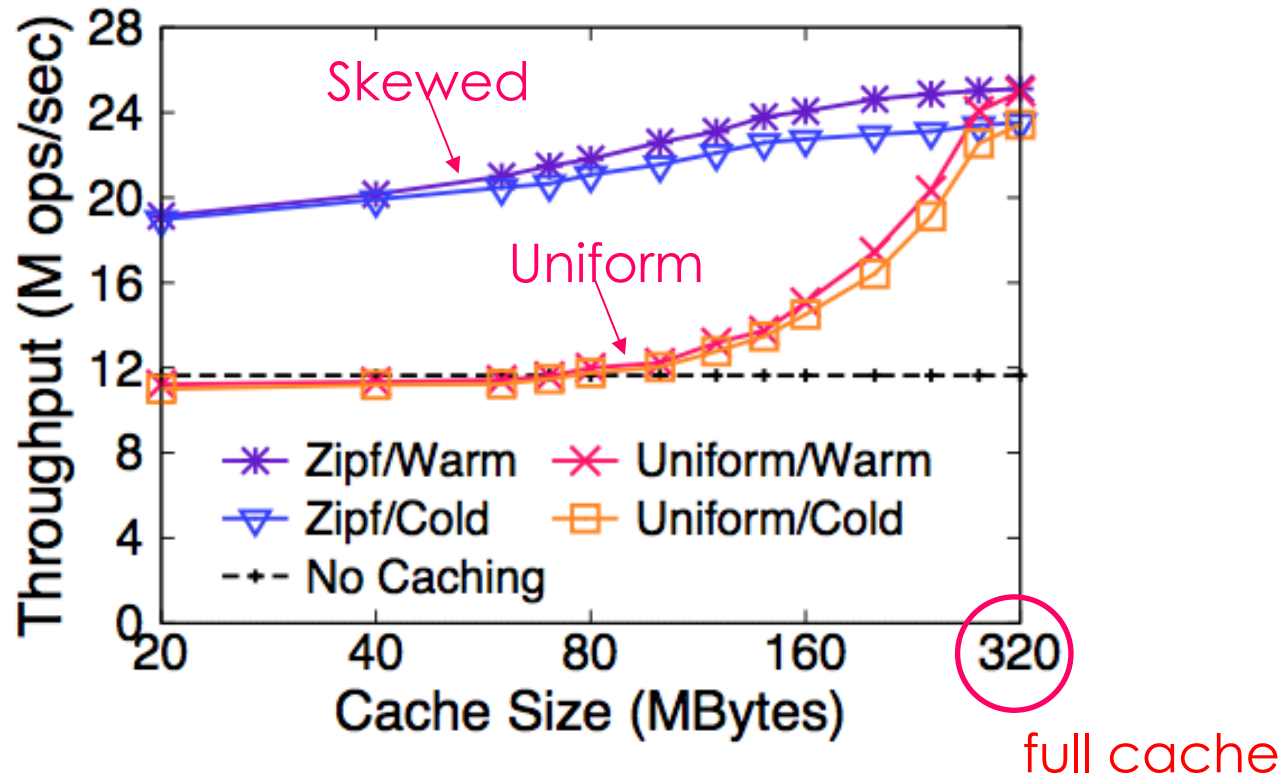
Read Performance

- ✓ DrTM-KV w/o caching provides a comparable perf.
- ✓ DrTM-KV w/ caching (DrTM-KV/\$) can achieve both lowest latency ($3.4\ \mu\text{s}$) and highest thpt ($23.4\ \text{Mops/s}$)



Setting: 1 server and 5 clients (up to 8 threads), 20 million k/v pairs
peak throughput of random RDMA READ $\approx 26\ \text{Mops/sec}$

Location-based Cache



Setting: 1 server and 5 clients (up to 8 threads), 20 million k/v pairs
Traditional replacement policy (i.e., LRU)

Extended Reading

RDMA-based KV store for range query

RDMA-based index caching for B⁺Tree

Supporting variable-length key

Scale-out ordered store and durability

Reference

1. *FaRM: fast remote memory*. NSDI 2014
2. *Fast In-memory Transaction Processing using RDMA and HTM*. SOSP 2015
3. *Balancing CPU and Network in the Cell Distributed B-Tree Store*. Usenix ATC 2016.
4. *Fast RDMA-based Ordered Key-Value Store using Remote Learned Cache*. OSDI 2020.

Conclusion

NEW hardware features open NEW opportunities



DrTM: The first design and impl. of combining **HTM**, **RDMA** and **NVM** to boost in-memory transaction processing with imperative functionality

Achieving **orders-of-magnitude** higher throughput and lower latency than prior traditional designs

Thanks



Questions

Home: <http://ipads.se.sjtu.edu.cn/drtm>

GitHub: <https://github.com/SJTU-IPADS>

IPADS, SJTU





Backup