

CS 610: Cache Coherence and False Sharing

Swarnendu Biswas

Department of Computer Science and Engineering,
Indian Institute of Technology Kanpur

Sem 2026-27-I



Types of Parallelism

Instruction-level Parallelism

Overlap instructions within a single thread of execution (e.g., pipelining, superscalar issue, and out-of-order execution)

Data-level Parallelism

Execute an instruction in parallel on multiple data values (e.g., vector instructions)

```
for (int i = 0; i < N; i++) {  
    c[i] = a[i] + b[i];  
}
```

Thread-level Parallelism

Concurrently execute multiple threads

Shared Memory Multiprocessor Architecture

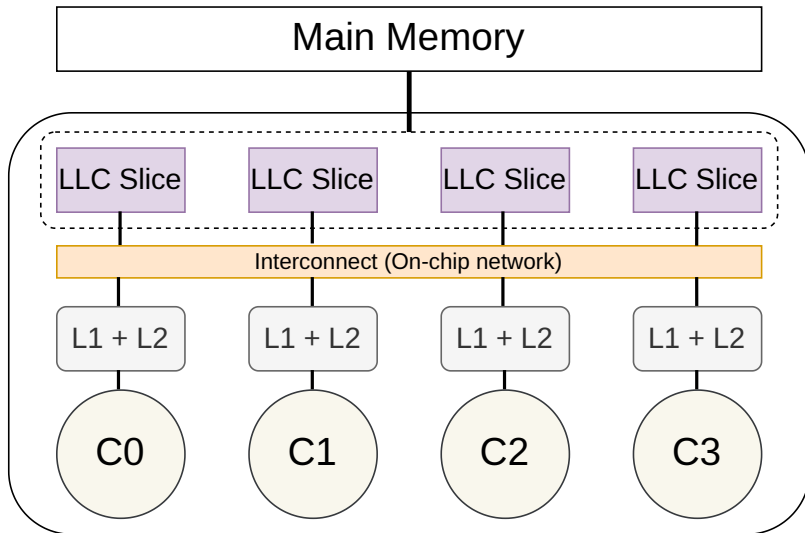
Single address space shared by multiple cores

- + Exploits TLP by having a number of cores
- + Can share data efficiently, communication is implicit through memory instructions (i.e., loads and stores)
- Cost for accessing shared memory can be uniform or non-uniform across cores

Processors privately cache data to improve performance

Reduces average data access time and saves interconnect bandwidth

Block Diagram of a SMP

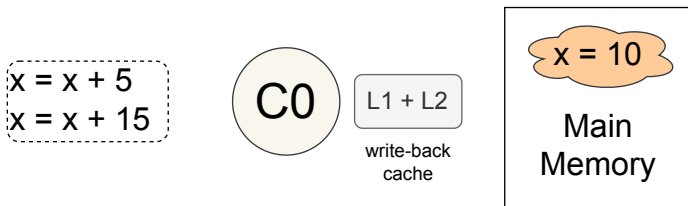


Data Coherence

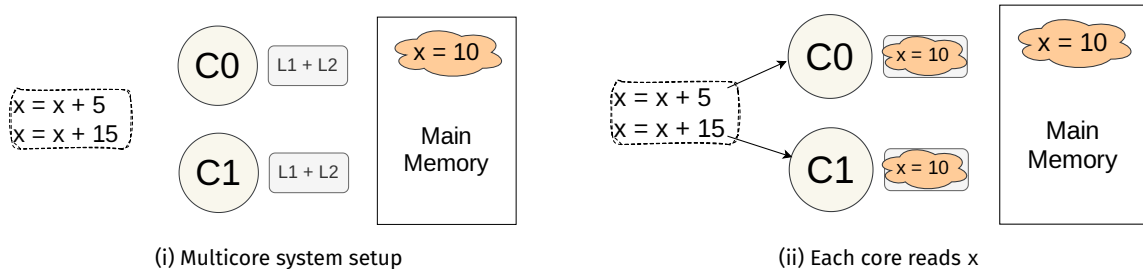
⚠ Private caches create data coherence problem

- Copies of a variable can be present in multiple caches
- Private copies of shared data must be **coherent**, i.e., all copies must have the same value (okay if the requirement holds eventually)

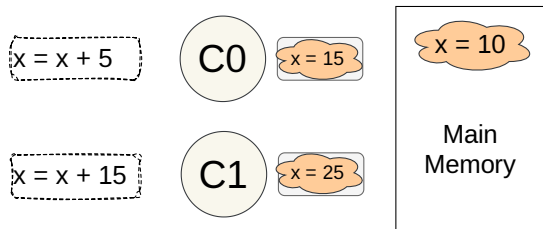
Consider the following sequence of operations on a single core system with write-back caches but without coherence



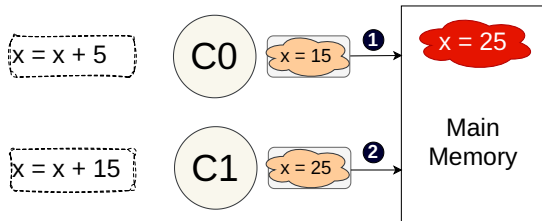
Coherence Challenge with Multicores



Coherence Challenge with Multicores



(iii) Each core updates x in its private cache



(iv) Cores write back x , a store is lost depending on the order of write backs

Can Write-through Caches Avoid the Coherence Problem?

Assume 3 cores with write-through caches

- (i) Core C₀ reads x from memory, caches it, and gets the value 10
- (ii) Core C₁ reads x from memory, caches it, and gets the value 10
- (iii) C₁ writes x=20, and updates its cached and memory values
- (iv) C₀ reads x from its cache and gets the value 10
- (v) C₂ reads x from memory, caches it, and gets the value 20
- (vi) C₂ writes x=30, and updates its cached and memory value

Sources of Errors in the Previous Setups

✖ Write-back cache

- Stores are not visible to memory immediately
- Order of write backs are important
- **Lesson learned:** do not allow more than one copy of a cache line in dirty state

⚠ Write-through cache

- The value in memory may be correct if the writes are correctly ordered
- A store proceeded when there is already a cached copy
- **Lesson learned:** must invalidate all cached copies before allowing a store to proceed

Understanding Coherence

A memory system is coherent if the following hold:

- (i) A read from a location X by a core C after a write by C to X always returns the value written by C, provided there are no writes of X by another processor between the two accesses by C.
- (ii) A read from a location X by a core C after a write to X by another core returns the written value if the read and write are sufficiently separated in time and no other writes to X occur between the two accesses.
- (iii) Writes to the same location are serialized. That is, two writes to the same location by any two cores are seen in the same order by all processors.

Correctness Requirement

✓ For sequential programs, there is only one correct output

A read from a memory location must return the “latest” value written to it

⚠ For parallel programs, there can be multiple correct outputs

- Defining “latest” precisely is crucial
- Assume that the latest value of a location is the latest value “committed” by any thread/process

Cache Coherence Protocol

Multicore processors implement a cache coherence protocol to keep private caches in sync

💡 A “cache coherence protocol” is a set of actions that ensure that a load to address A returns the “last committed” value to A

- Essentially, makes one core's write to A visible to other cores by propagating the write to other caches
- Aims to make the presence of private caches functionally invisible
- Coherence protocols usually operate at the granularity of whole cache blocks (e.g., 64 bytes)

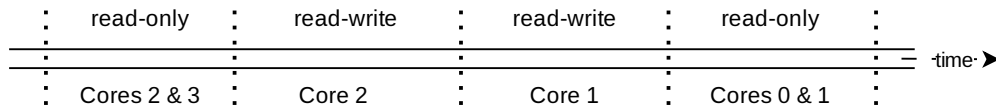
Cache Coherence Protocol Invariants

1. Enforces the Single-Writer-Multiple-Reader (SWMR) invariant

For any given memory location, at any given moment in time, there is either a single core that may write it (including read) or some number of cores that may read it

2. Data values must be propagated correctly (data invariant)

The value of a memory location at the start of a read-only time period is the same as the value of the location at the end of its last read-write time period



Alternate Definitions of Coherence

Definition 2

A coherent system must appear to execute all threads' loads and stores to a **single** memory location in a total order that respects the program order of each thread

Definition 3

A coherent system satisfies two invariants:

write propagation every store is eventually made visible to all cores

write serialization writes to the same memory location are serialized (i.e., observed in the same order by all cores)

Memory Consistency Model

Correctness of Shared-Memory Programs

“To write correct and efficient shared memory programs, programmers need a precise notion of how memory behaves with respect to read and write operations from multiple processors”

Busy-Wait Paradigm

```
1 Object X = null;  
2 boolean done = false;
```

Thread 1

```
1 X = new Object();  
2 done = true;
```

Thread 2

```
1 while (!done) {}  
2 X.compute();
```

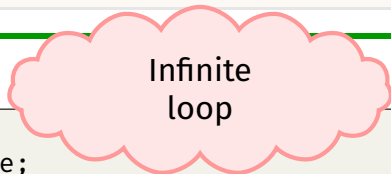
Possible Errors

Thread 1

```
1 X = new Object();  
2  
3  
4 done = true;
```

Thread 2

```
1  
2 tmp = done;  
3 while (!temp) {}  
4
```



Thread 1

```
1 done = true;  
2  
3  
4 X = new Object();
```

Thread 2

```
1  
2 while (!done) {}  
3 X.compute();  
4
```



Reordering of Accesses by Hardware

Accesses are to **different** addresses

- Store-store ● Non-FIFO write buffer (first store misses in the cache while the second hits or the second store can coalesce with an earlier store)
- Load-load ● Cache hits, dynamic scheduling, execute out of order
- Load-store ● Cache hits, out-of-order core
- Store-load ● FIFO write buffer with bypassing, out-of-order core

Reordering of Accesses by Hardware

Accesses are to **different** addresses

- Store-store ● Non-FIFO write buffer (first store misses in the cache while the second hits or the second store can coalesce with an earlier store)
- Load-load ● C ● Correct in a single-threaded context
- Load-store ● C ● Non-trivial in a multithreaded context
- Store-load ● FIFO write buffer with bypassing, out-of-order core

What values can a load return?

Return the “last” write

- Uniprocessor: program order defines the “last” write
- Multiprocessor: operations from different cores/threads are not related by program order

Can both r1 and r2 be set to zero?

```
1 x = 0;  
2 y = 0;
```

Core 1

```
1 S1: x = new Object();  
2 L1: r1 = y;
```

Core 2

```
1 S2: y = new Object();  
2 L2: r2 = x;
```

Memory Consistency Model

Set of rules that govern how systems process memory requests from multiple processors

- Determines the **order** in which memory operations appear to execute
- Specifies allowed behaviors of multithreaded programs executing with shared memory
 - ▶ Both at the hardware-level and at the programming-language-level
 - ▶ There can be multiple correct behaviors

Importance of memory consistency models

- + Determines what optimizations are correct
- + Contract between the programmer and the hardware
- + Influences ease of programming and program performance
- + Impacts program portability

Issues with Memory Consistency

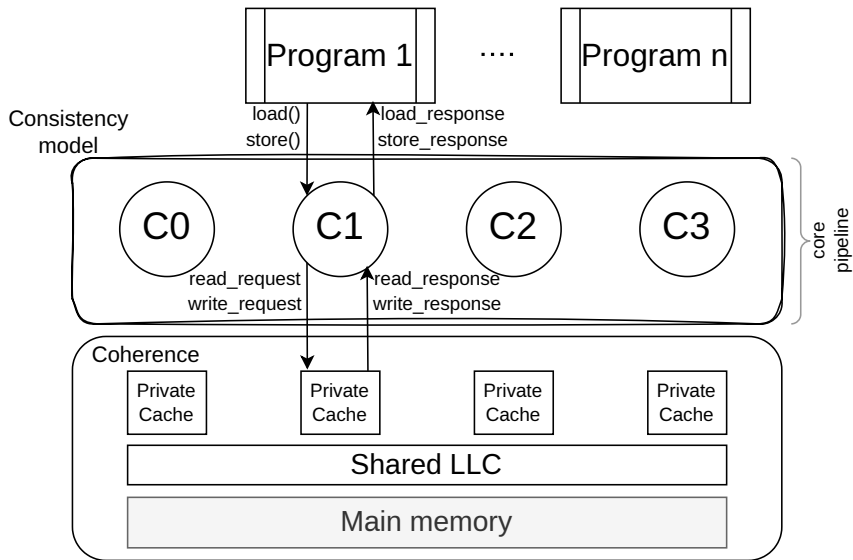
Visibility

When are the effects of one thread (e.g., updating a memory location) visible to another?

Ordering

When can operations of any given thread appear out of order to another thread?

Memory Consistency vs Cache Coherence



Memory Consistency vs Cache Coherence

Memory Consistency

- Defines shared memory behavior
- Related to all shared-memory locations
- Policy on when new value is propagated to other cores
- Memory consistency implementations can use cache coherence as a “black box”

Cache Coherence

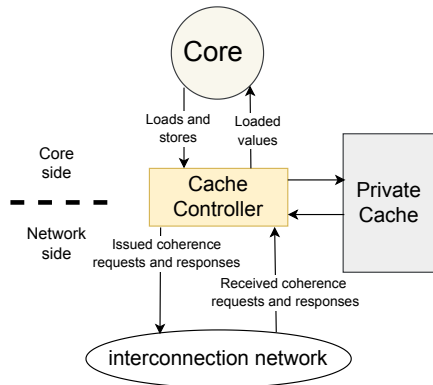
- Does not define shared memory behavior
- Specific to a single shared-memory location
- Propagates a new value to other cached copies
- Invalidation-based or update-based

Cache Coherence Protocols

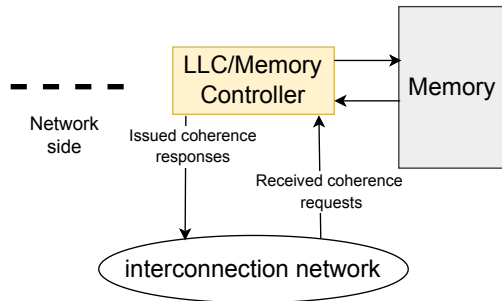
Implementing Coherence Protocols

Protocols are implemented as finite state machines called coherence controllers

A protocol formalizes the interactions between the different coherence controllers



Cache controller



Memory controller

Important Characteristics of a Cache Block

Coherence protocols are implemented by associating states with each cache block

Validity A valid block has the most up-to-date values for the block. The block may be read. It can be written if it is also exclusive.

Dirtyness A cache block is dirty if its value is the most up-to-date, and the value differs from the value in the LLC/memory.

Exclusivity A cache block is exclusive if it is the only privately cached copy of that block except perhaps in the shared LLC.

Ownership A cache or memory controller is the owner of a block if it is responsible for responding to coherence requests for that block. In most protocols, there is exactly one owner of a given block at all times.

Stable States

M, S, and I are commonly-used states

- Modified (M)** The block is valid, exclusive, owned, and potentially dirty. The cache has the only valid copy of the block, the cache must respond to requests for the block, and the copy of the block at the LLC/memory is potentially stale.
- Shared (S)** The block is valid but not exclusive, not dirty, and not owned. The cache has a read-only copy of the block. There may be multiple processors caching a line in S state.
- Invalid (I)** The cache either does not contain the block (not present) or it contains a potentially stale copy that it may not read or write.

Different protocol extensions add additional states (e.g., E, O, and F) to optimize for certain sharing patterns

Common Coherence Transactions

Transaction	Goal of Requestor
GetS	Obtain block in Shared (read-only) state
GetM	Obtain block in Modified (read-write) state
Upg	Upgrade block state from read-only (Shared or Owned) to read-write (Modified); Upg (unlike GetM) does not require data to be sent to requestor
PutS	Evict block in Shared state
PutE	Evict block in Exclusive state
PutO	Evict block in Owned state
PutM	Evict block in Modified state

Communication between Core and Cache Controller

Event	Response from Cache Controller
Load	If cache hit, respond with data from cache; else initiate GetS transaction
Store	If cache hit in state E or M, write data into cache; else initiate GetM or Upg transaction
Atomic RMW	If cache hit in state E or M, atomically execute RMW semantics; else initiate GetM or Upg transaction
Instruction fetch	If I-cache hit, respond with instruction from cache; else initiate GetS transaction
Read-only prefetch	If cache hit, ignore; else (optionally) initiate GetS transaction
Read-write prefetch	If cache hit in state M, ignore; else (optionally) initiate GetM or Upg transaction
Replacement	Depending on state of block, initiate PutS, PutE, PutO, or PutM transaction

Types of Coherence Protocols

💡 Protocols differ in when and how writes are propagated

- The writes can be propagated synchronously or asynchronously
- Synchronous propagation means a write is made visible to other cores *before* returning

❗ Two main axes to classify synchronous protocols

- (i) Invalidation-based protocol and Update-based protocol
- (ii) Snoopy protocol and Directory protocol

Invalidation-Based Protocols

Invalidate all cached copies before allowing a store to proceed

Need to know the location of cached copies

Solution 1 : Broadcast that a core is going to do a store and sharers invalidate themselves

Solution 2 : Keep track of the sharers and invalidate them when needed

- + Only store misses go on bus and subsequent stores to the same line are cache hits
- Sharers will miss next time they try to access the line

Update all cached copies with the new value of the store

- + Sharers continue to hit in the cache, do not need to initiate and wait for a GetS transaction to complete
- Prevalence of spatial and temporal locality can lead to unnecessary updates, leading to increased bandwidth requirements
- Complicates implementing many memory consistency models

Snoopy Protocol

Cache controller initiates a request for a block by broadcasting a request message to all other coherence controllers

- Each cache controller snoops (i.e., continuously monitors) the shared medium (e.g., bus or switch) for write activity concerned with its cached data addresses
- Assumes a global bus structure where communication can be seen by all
- Relies on the interconnection network to deliver the broadcast messages in a consistent order to all cores

How do you prevent simultaneous writes from different controllers?

Snoopy Protocol

Invalidate on a write

- Core that wants to write to an address grabs a bus cycle and broadcasts a “write invalidate” message
- All snooping caches invalidate their private copy of the appropriate cache line
- Core writes to its cached copy
- Any future read in other cores will now miss in cache and refetch new data

Update on a write

- Core that wants to write to an address grabs a bus cycle and broadcasts new data as it updates its own copy
- All snooping caches update their copy

Cache controller initiates a request for a block by unicasting it to the block's home memory controller

- Memory controller maintains a directory that holds state about each block in the LLC/memory (e.g., coherence state, the current owner ID, and a bitvector for the list of current sharers)
- If the LLC/memory is the owner, the memory controller completes the transaction by sending a data response to the requestor
- If a cache controller is the owner, the memory controller forwards the request to the owner cache
- When the owner cache receives the forwarded request, it completes the transaction by sending a data response to the requestor

Snoopy vs Directory Protocol

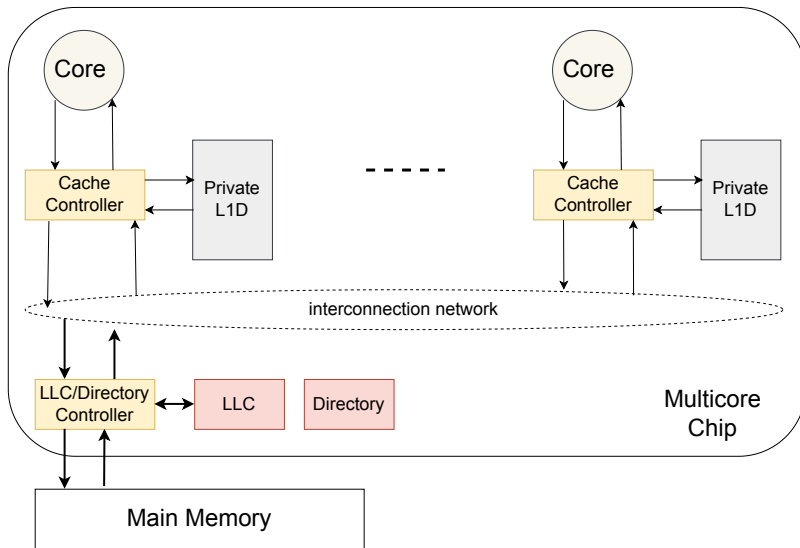
Snoopy Protocol

- Does not scale to large core counts because of broadcast messages
- Requires some ordering guarantees on messages which limits network optimizations

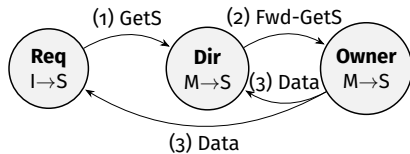
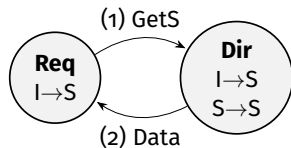
Directory Protocol

- + Scalable because messages are unicast
- + The directory can be distributed to improve scalability
- Few transactions take more cycles when the home is not the owner
- Memory requirement increases with core count

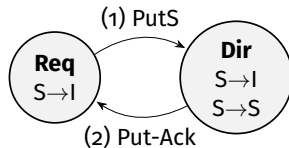
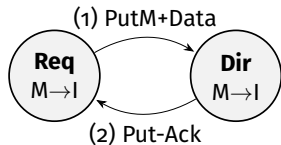
Directory System Model



MSI Protocol

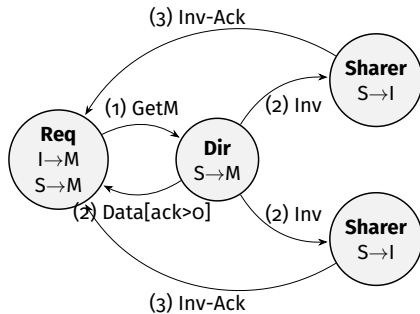
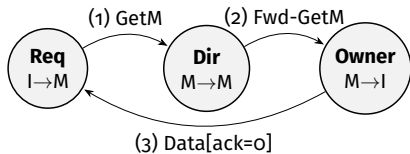
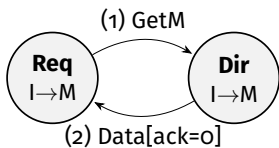


Transitions from I to S



Transitions from M or S to I

MSI Protocol



Transitions from I or S to M

Usefulness of E State

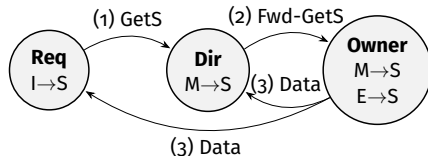
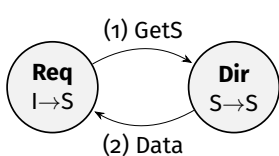
💡 Cores often read data before updating it

- Oftentimes, there is only one sharer in the system (also applicable for single-threaded programs)
- With MSI, the core will issue two coherence transactions: GetS followed by an Upg

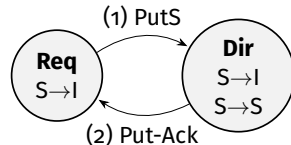
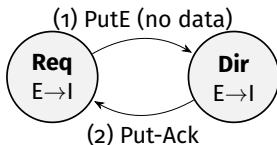
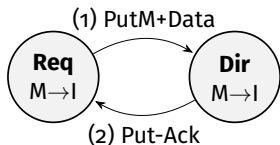
👍 Optimization with E state

- The core that issues GetS for a block gets it in E state if there are no existing sharers
- E state indicates the cache line is clean and is the only cached copy
- The core may then silently upgrade the block from E to M without issuing another coherence request
- We will assume E is an ownership state, which implies evictions cannot be silent

MESI Protocol

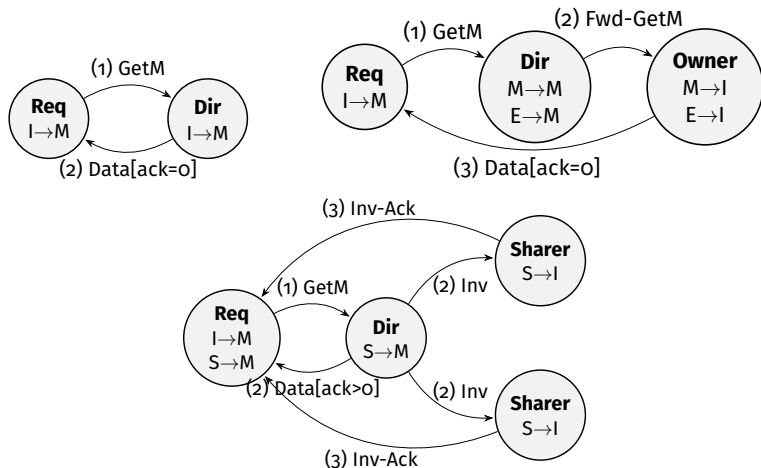


Transitions from I to S



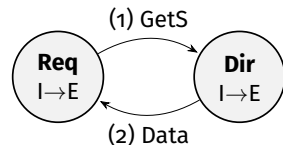
Transitions from M or E or S to I

MESI Protocol



If the only sharer is the requestor, then no Inv messages are sent and the Data message from the Dir to Req has an Ack count of zero.

Transitions from I or S to M



Transitions from I to E

Adding an Owned (O) State

💡 **Suppose a cache has a block in state M or E and receives a GetS**

The cache changes the block state from M or E to S and sends the data to **both** the requestor and the memory controller

- Why is it necessary to send the data to the memory controller?

Adding an Owned (O) State

Owned state indicates that the block is valid, dirty, and shared, and the cache is the owner

- The owner cache does **not** have permission to modify the block, and is responsible for eventually updating memory

👍 Advantages of MOESI protocol (used in AMD Opteron)

- + Eliminates the extra data message to update the LLC/memory when a cache receives a GetS request in the M or E state
- + Eliminates potentially unnecessary writes to the LLC if the block is written again before being written back to the LLC
- + Allows subsequent requests to be satisfied by the private owner cache instead of the slower LLC/memory

Cache Contention

Types of Cache Contention

Cache line contention arises from two types of read-write data sharing: true sharing and false sharing

True Sharing

- Same location is accessed by multiple cores
- Can be fixed only by means of algorithmic changes

False Sharing

- Two unrelated locations lie on the same cache line and are accessed by multiple cores
- Can be fixed by modifying the data layout (manual or automated)

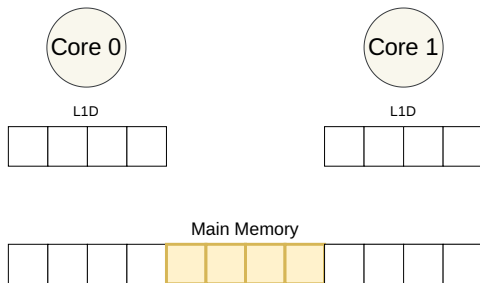
False Sharing

False sharing is a performance problem in cache-coherent systems

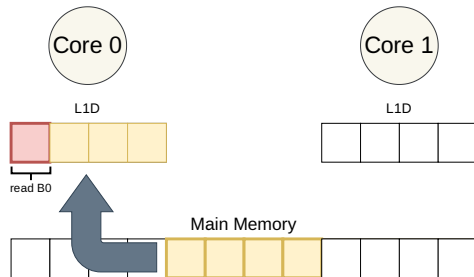
- Cores contend on cache blocks instead of data
- Can arise when threads access global or heap memory



Understanding False Sharing

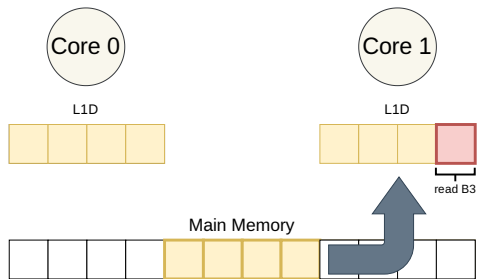


(i) Multicore system setup

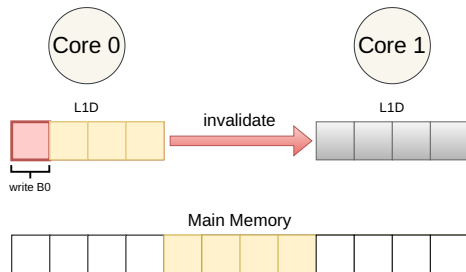


(ii) Core 0 reads block offset B0

Understanding False Sharing

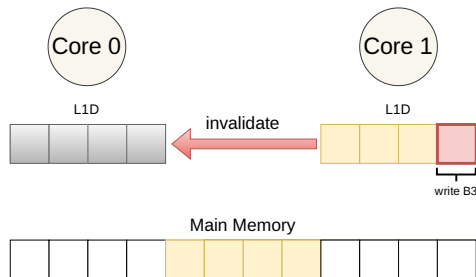


(iii) Core 1 reads block offset B3

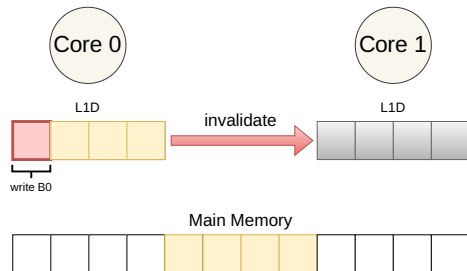


(iv) Core 0 writes block offset B0

Understanding False Sharing



(v) Core 1 writes block offset B3



(ii) Core 0 writes block offset B0

Cache misses resulting from data sharing across cores are called **coherence** misses (e.g., the write to B3 by Core 1 in (v))

Impact of False Sharing

```
1  int array[100];  
2  
3  void *func(void *param) {  
4      int index = *((int*)param);  
5      for (int i = 0; i < 100000000; i++)  
6          array[index]+=1;  
7  }  
8  
9  int main() {  
10     int first_elem = 0;  
11     int bad_elem = 1;  
12     int good_elem = 99;  
13     pthread_t thread_1;  
14     pthread_t thread_2;  
  
15     clock_gettime(CLOCK_REALTIME, ...);  
16     func((void*)&first_elem);  
17     func((void*)&bad_elem);  
18     clock_gettime(CLOCK_REALTIME, ...);
```

```
22     clock_gettime(CLOCK_REALTIME, ...);  
23     pthread_create(&thread_1, NULL, func,  
24                   (void*)&first_elem);  
25     pthread_create(&thread_2, NULL, func,  
26                   (void*)&bad_elem);  
27     pthread_join(thread_1, NULL);  
28     pthread_join(thread_2, NULL);  
29     clock_gettime(CLOCK_REALTIME, ...);  
  
30     clock_gettime(CLOCK_REALTIME, ...);  
31     pthread_create(&thread_1, NULL, func,  
32                   (void*)&first_elem);  
33     pthread_create(&thread_2, NULL, func,  
34                   (void*)&good_elem);  
35     pthread_join(thread_1, NULL);  
36     pthread_join(thread_2, NULL);  
37     clock_gettime(CLOCK_REALTIME, ...);  
38 }
```

Impact of False Sharing

```
1  int array[100];

3  void *func(void *param) {
    int index = *((int*)param);

5      swarnendu@vindhya:~/falseSharing$ ./a.out
    array[first_element]: 300000000      array[bad_element]: 200000000      array[good_element]: 100000000

7  }

9  Time take with false sharing      : 330.773397 ms
    Time taken without false sharing : 173.216272 ms
    Time taken in Sequential computing : 325.908369 ms
11  swarnendu@vindhya:~/falseSharing$ ./a.out
    array[first_element]: 300000000      array[bad_element]: 200000000      array[good_element]: 100000000

13

15  Time take with false sharing      : 326.360157 ms
    Time taken without false sharing : 176.690517 ms
    Time taken in Sequential computing : 324.763425 ms
17  func((void*)&first_elem);
    func((void*)&bad_elem);
19  clock_gettime(CLOCK_REALTIME, ...);

22  clock_gettime(CLOCK_REALTIME, ...);
    pthread_create(&thread_1, NULL, func,
                   (void*)&first_elem);
24  pthread_create(&thread_2, NULL, func,

38 }
```

Introducing False Sharing is Easy

```
1 // Global variables accessed by
2 // different threads me and you
3 me = 1;
4 you = 2;
```

```
1 // Class/struct fields can lie on
2 // the same line
3 class X {
4     int me;
5     float you;
6 };
```

```
1 // Heap objects can lie on the
2 // same line
3 me = new Foo();
4 you = new Bar();
```

```
1 // Array accesses by different
2 // threads me and you
3 array[me] = 12;
4 array[you] = 13;
```


False Sharing in Real Applications

False sharing problems were reported in Linux kernel, JVM, and Boost library

The screenshot shows a blog post from mikaelronstrom.blogspot.com dated Tuesday, April 10, 2012. The main title is "MySQL team increases scalability by >50% for Sysbench OLTP RO in MySQL 5.6 labs release april 2012". The text discusses a performance breakthrough in MySQL 5.6, mentioning a 6x improvement in performance with specific scalability fixes. It also mentions that the solutions to scalability are no longer only related to protected regions and their related mutexes and read-write locks or atomic variables. The post is categorized under "INSPIRATIONAL MESSAGES OF THE WEEK" and "Sense and Sensibility and Experiments". There is a "FOLLOWERS" section at the bottom showing a grid of user avatars.

Seeing through hardware counters: a journey to threefold performance increase



Netflix Technology Blog · Follow

Published in Netflix TechBlog · 10 min read · Nov 10, 2022



1.7K



18



By [Vadim Filanovsky](#) and [Harshad Sane](#)

In one of our previous blogposts, [A Microscope on Microservices](#) we outlined three broad domains of observability (or “levels of magnification,” as we referred to them) — Fleet-wide, Microservice and Instance. We described the tools and techniques we use to gain insight within each domain. There is, however, a class of problems that requires an even stronger level of magnification going deeper down the stack to introspect CPU microarchitecture. In this blogpost we describe one such problem and

False Sharing Mitigation Techniques

- Compiler optimizations (cache block padding)
 - Inflates memory requirement, can complicate address computations
- Run-time solutions (e.g., use hardware performance counters to detect false sharing)
 - Can miss false sharing instances
- Sub-block coherence or false-sharing-aware coherence protocols
- Cache-conscious programming

Fixing False Sharing can be Non-trivial

mikaelronstrom.blogspot.com/2012/04/

TUESDAY, APRIL 10, 2012

MySQL team increases scalability by >50% for Sysbench OLTP RO in MySQL 5.6 labs release april 2012

A MySQL team focused on performance recently met in an internal meeting to discuss and work on MySQL scalability issues. We had gathered specialists on InnoDB and all its aspects of performance including scalability, adaptive flushing and other aspects of InnoDB, we had also participants from MySQL support to help us understand what our customers need and a number of generic specialists on computer performance and in particular performance of the MySQL software.

The fruit of this meeting can be seen in the MySQL 5.6 labs release april 2012 released today. We have a new very interesting solution to the adaptive flushing problem. We also made a significant breakthrough in MySQL scalability. On one of our lab machines we were able to increase performance of the Sysbench OLTP RO test case by more than 50% by working together to find the issues and then quickly coming up with the solution to the issues. Actually in one particular test case we were able to improve MySQL performance by 6x with these scalability fixes.

Fixing False Sharing can be Non-trivial

Problem is often embedded inside the source code

False sharing is sensitive to

- Application behavior (e.g., mapping of threads to cores)
- Compiler toolchain (e.g., data layout optimizations and memory allocator)
 - ▶ GCC **unintentionally eliminates** false sharing in Phoenix linear_regression benchmark at certain optimization levels, while LLVM does not do so at any optimization level
- Execution environment (e.g., object placements on the cache line, hardware platform with different cache line sizes)

Detect False Sharing with perf c2c

Idea is to check whether loads/stores frequently hit in a remote cache line that is in M state

Input  with false sharing

Compile with `gcc -O0 -g fs-perf-c2c.c -pthread`

Using perf c2c

```
# May need to update /proc/sys/kernel/perf_event_paranoid to -1
# sudo sh -c 'echo 1 >/proc/sys/kernel/perf_event_paranoid'
perf c2c record -F 30000 -u -- ./a.out
perf c2c report -NN -i perf.data --stdio > ./perf-report.out
# Check the analysis report
```

Code Examples

-  fs-time.cpp
-  fs-perf-c2c.c
-  Makefile

References

-  V. Nagarajan et al. A Primer on Memory Consistency and Cache Coherence. Chapters 1,2,6–8, 2nd edition, Morgan and Claypool.
-  J. Hennessy and D. Patterson. Computer Architecture: A Quantitative Approach. Sections 5.2, 5.4, 6th edition, Morgan Kaufmann.
-  A. Gupta et al. Parallel Computer Architecture: A Hardware/Software Approach. Morgan Kaufmann.
-  Mainak Chaudhuri. Cache Coherence. Computer Architecture Summer School, IIT Kanpur, 2018.