

CS 610: Shared-Memory Synchronization

Swarnendu Biswas

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

Sem 2026-27-I



Is the Code Thread-Safe?

```
1  class Set {  
2      final Vector elems = new Vector();  
3  
4      void add(Object x) { // Free of data races  
5          if (!elems.contains(x))  
6              elems.add(x);  
7      }  
8  }  
9  
10 class Vector {  
11     synchronized void add(Object o) { ... }  
12     synchronized boolean remove(Object o) { ... }  
13     synchronized boolean contains(Object o) { ... }  
14 }
```

What is the Desired Property?

```
1  class Set {  
2      final Vector elems = new Vector();  
3  
4      void add(Object x) { // Free of data races  
5          if (!elems.contains(x))  
6              elems.add(x);  
7      }  
8  }  
9  
10 class Vector {  
11     synchronized void add(Object o) { ... }  
12     synchronized boolean remove(Object o) { ... }  
13     synchronized boolean contains(Object o) { ... }  
14 }
```

Synchronization Patterns

Mutual exclusion – updates need to be serialized

```
bool lock = false;
```

```
1 lock_acquire():  
2   while TAS(&lock)  
3     // spin
```

```
1 lock_release():  
2   lock = false;  
3
```

Conditional synchronization – events need to occur in a specified order

```
1 while !condition  
2   // spin
```

Other forms – e.g., synchronize across threads or control the number of simultaneous accesses to a shared resource

Desired Synchronization Properties

Mutual exclusion (safety property)

- Critical sections on the same lock from different threads do not overlap

Deadlock freedom (liveness property)

- If some threads attempt to acquire the lock, then some thread should be able to acquire the lock
- Individual threads may be infinitely delayed

Starvation freedom (liveness property)

- Every thread that acquires a lock eventually releases it
- A lock acquire request must eventually succeed within bounded steps
- Implies deadlock freedom

Classic Mutual Exclusion Algorithms

LockOne: What could go wrong?

```
1  class LockOne implements Lock {  
2      private boolean[] flag = new boolean[2];  
3      public void lock() {  
4          int i = ThreadID.get();  
5          flag[i] = true;  
6          j = 1-i;  
7          while (flag[j]) {}  
8      }  
9      public void unlock() {  
10         int i = ThreadID.get();  
11         flag[i] = false;  
12     }  
13 }
```

- LockOne satisfies mutual exclusion
- LockOne fails deadlock-freedom, concurrent execution can deadlock

LockTwo: What could go wrong?

```
1  class LockTwo implements Lock {  
2      private int victim;  
3      public void lock() {  
4          int i = ThreadID.get();  
5          victim = i;  
6          while (victim == i) {}  
7      }  
8  
9      public void unlock() {}  
10 }
```

- LockTwo satisfies mutual exclusion
- LockTwo fails deadlock-freedom, sequential execution deadlocks

Peterson's Algorithm

```
1  class PetersonLock {
2      private boolean[] flag = new boolean[2];
3      private int victim;
4
5      public void lock() {
6          int i = ThreadID.get();
7          int j = 1-i;
8          flag[i] = true;
9          victim = i;
10         while (flag[j] && victim == i) {}
11     }
12
13     public void unlock() {
14         int i = ThreadID.get();
15         flag[i] = false;
16     }
17 }
```

Peterson's Algorithm

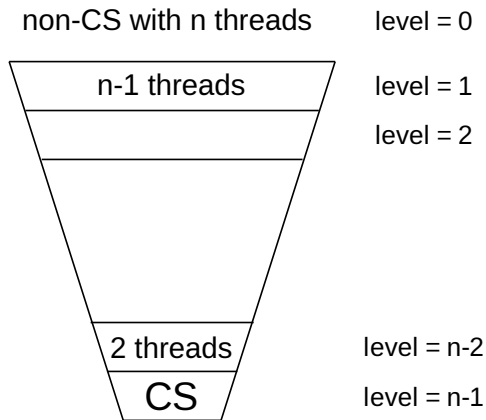
```
1 class PetersonLock {
2     private boolean[] flag = new boolean[2];
3     private int victim;
4
5     public void lock() {
6         int i = ThreadID.get();
7
8         // Does this algorithm satisfy mutual exclusion under
9         // sequential consistency?
10
11     }
12
13     public void unlock() {
14         int i = ThreadID.get();
15         flag[i] = false;
16     }
17 }
```

- Does this algorithm satisfy mutual exclusion under sequential consistency?
- What if we do not have sequential consistency?

Filter Lock for n Threads

Filter lock is a generalization of Peterson's lock to $n > 2$ threads

- There are $n - 1$ waiting rooms called “levels”
- At least one thread trying to enter a level succeeds
- One thread gets blocked at each level if many threads try to enter



Filter Lock

```
1 class FilterLock {
2     int[] level;
3     int[] victim;
4     public FilterLock() {
5         level = new int[n];
6         victim = new int[n];
7         for (int i=0; i<n; i++)
8             level[i] = 0;
9     }
10    public void unlock() {
11        int me = ThreadID.get();
12        level[me]= 0;
13    }
14 }
```

```
15 public void lock() {
16     int me = ThreadID.get();
17     // Attempt to enter level i
18     for (int i=1; i<n; i++) {
19         // visit level i
20         level[me] = i;
21         victim[i] = me;
22         // spin while conflict exists
23         while (( $\exists k \neq me$ )
24             level[k] >= i && victim[i] ==
25                 me) {}
26     }
27 }
```

Fairness

Starvation freedom is good, but maybe threads should not wait too much

- For example, it would be great if we could order threads by the order in which they performed the first step of the `lock()` method

Bounded Waiting

- Divide the `lock()` method into two parts
Doorway interval (DA) expresses intent to synchronize, finishes in finite steps
Waiting interval (WA) wait for turn to synchronize, may take unbounded steps
- A lock is first-come first-served if $D_A^j \rightarrow D_B^k$, then $CS_A^j \rightarrow CS_B^k$

r-bounded waiting

For threads A and B, if $D_A^j \rightarrow D_B^k$, then $CS_A^j \rightarrow CS_B^{k+r}$

Lamport's Bakery Algorithm

```
1  class Bakery implements Lock {
2      boolean[] choosing;
3      Label[] lbl;
4      public Bakery(int n) {
5          choosing = new boolean[n];
6          lbl = new Label[n];
7          for (int i = 0; i < n; i++) {
8              choosing[i] = false;
9              lbl[i] = 0;
10         }
11     }
12
13     public void unlock() {
14         choosing[ThreadID.get()] = false;
15     }
```

Lamport's Bakery Algorithm

```
15  public void lock() {  
16      int i = ThreadID.get();  
17      choosing[i] = true; // Getting a label  
18      lbl[i] = max(lbl[0], ..., lbl[n-1]) + 1;  
19      while (( $\exists$  k != i) choosing[k] && (lbl[k], k) << (lbl  
        [i], i)) {}  
20  }  
21  }
```

$(\text{lbl}[i], i) << (\text{lbl}[j], j)$ iff

- $\text{lbl}[i] < \text{lbl}[j]$, or
- $\text{lbl}[i] = \text{lbl}[j]$ and $i < j$

Lamport's Bakery Algorithm

```
15 public void lock() {  
16     int i = ThreadID.get();  
17     choosing[i] = true; // Getting a label  
18     lbl[i] = max(lbl[0], ..., lbl[n-1]) + 1;  
19     while (( $\exists$  k  $\neq$  i) choosing[k] && (lbl[k], k) << (lbl
```

- Need to compare own label with all other threads' labels irrespective of their intent to enter the critical section
- Cost of locking increases with the number of threads

$(\text{lbl}[i], i) << (\text{lbl}[j], j)$ iff

- $\text{lbl}[i] < \text{lbl}[j]$, or
- $\text{lbl}[i] = \text{lbl}[j]$ and $i < j$

Lamport's Fast Lock

- Programs with highly contended locks are likely to not scale
- **Insight:** Ideally spin locks should be free of contention in well-designed systems, so optimize for the common case
- Idea:
 - ▶ Use two lock fields `fast_check` and `slow_check`
 - ▶ Acquire: Thread `t` writes its ID to `fast_check` and `slow_check` and checks for intervening writes

Lamport's Fast Lock

```
1  class LFL implements Lock {  
2      // Two checkpoints to guarantee mutual exclusion  
3      private int fast_check, slow_check;  
4      boolean[] trying;  
5      LFL() {  
6          slow_check = 1;  
7          for (int i = 0; i < n; i++)  
8              trying[i] = false;  
9      }  
10     public void unlock() {  
11         slow_check = 1;  
12         trying[ThreadID.get()] = false;  
13     }  
14 }
```

Lamport's Fast Lock

```
14 public void lock() {
15     int self = ThreadID.get();
16 start:
17     trying[self] = true;
18     fast_check = self;
19     if (slow_check !=  $\perp$ ) {
20         // Someone else is in the CS
21         trying[self] = false;
22         while (slow_check !=  $\perp$ ) {}
23         goto start; // Retry
24     }
25     slow_check = self;
```

```
25 // Ensure atomicity
26 if (fast_check != self) {
27     trying[self] = false;
28     for (i  $\in$  T) {
29         while (trying[i] == true) {}
30     }
31     if (slow_check != self) {
32         while (slow_check !=  $\perp$ ) {}
33         goto start;
34     }
35 }
36 }
```

Evaluating Performance of a Lock

- Acquisition latency** Lock acquire should be cheap in the absence of contention
- Space overhead** Maintaining lock metadata should not impose high memory overhead
- Fairness** Processors should enter the CS in the order of lock requests
- Bus traffic** Worst case lock acquire traffic should be low
- Scalability** Latency and traffic should scale slowly with the number of processors

Practicality of Classical Mutual Exclusion Algorithms

A write (i.e., regular memory store) by a thread to a memory location can be overwritten without any other thread seeing the first write

Need to read and write n distinct memory locations where n is the maximum number of concurrent threads

- n is a lower bound on the number of required locations
- Motivates the need for read-write operations with stronger guarantees

Atomic Hardware Instructions

Hardware Locks

- Locks can be completely supported by hardware
- Ideas:
 - (i) Have a set of lock lines on the bus, processor wanting the lock asserts the line, others wait, priority circuit used for arbitrating
 - (ii) Special lock registers, processors wanting the lock acquire ownership of the registers

What could be some problems?

Limitations with Hardware Locks

- Waiting logic is critical for lock performance
 - ▶ A thread can (i) busy wait (i.e., spin), (ii) block, or (iii) use a hybrid strategy (e.g., busy wait for some time and then block)
- Hardware locks are not popularly used
 - Limited in number due resource constraints
 - Inflexible in implementing wait strategies

We continue to rely on software locks

Can optionally make use of hardware instructions for better performance

Common Atomic Primitives

Modern architectures provide many atomic read-modify-write (RMW) instructions for synchronization

- For example, test-and-set, fetch-and-add, compare-and-swap, and load-linked/store-conditional

TAS

X86, SPARC

```
1 bool TAS(word* loc):  
2     atomic {  
3         tmp := *loc;  
4         *loc := true; // set  
5     }  
6     return tmp;
```

swap

X86, SPARC

```
1 word swap(word* a, word b):  
2     atomic {  
3         tmp := *a;  
4         *a := b;  
5     }  
6     return tmp;
```

Spin Lock with TAS on X86_64

```
static inline uint8_t tas(volatile uint8_t *p) {
    uint8_t old = 1;
    asm volatile(
        "xchg %0, %1" // atomic xchg with mem
        : "+q"(old), "+m"(*p) // old is both input/output; *p is read/write
        : // no extra inputs
        : "memory" // compiler barrier for reordering
    );
    return old;
}

static volatile uint8_t lock = 0;

static inline void acquire() {
    while (tas(&lock)) { } // Add delay to reduce contention
}

static inline void release() {
    asm volatile("" ::: "memory");
    lock = 0;
}
```

Spin Lock with TAS on X86_64

```
#include <atomic>

std::atomic_flag lock = ATOMIC_FLAG_INIT;

// Using sequential consistency as the default ordering

// Spin until the lock is acquired
while (std::atomic_flag_test_and_set(&lock)) {}

std::atomic_flag_clear(&lock);
```

Common Atomic Primitives

fetch_and_inc

uncommon

```
1 int FAI(int* loc):  
2     atomic {  
3         tmp := *loc;  
4         *loc := tmp+1;  
5     }  
6     return tmp;
```

fetch_and_add

uncommon

```
1 int FAA(int* loc, int n):  
2     atomic {  
3         tmp := *loc;  
4         *loc := tmp+n;  
5     }  
6     return tmp;
```

C++ 11 onward provides `std::atomic<T>::fetch_add()`

Common Atomic Primitives

fetch_and_inc

uncommon

```
1 int FAI(int* loc):  
2     atomic {  
3         tmp := *loc;  
4         *loc  
5     }  
6     return
```

fetch_and_add

uncommon

```
1 int FAA(int* loc, int n):  
2     atomic {  
3         tmp := *loc;
```

How can we implement a mutual exclusion lock with FAI?

C++ 11 onward provides `std::atomic<T>::fetch_add()`

Compare-and-Swap (CAS) Primitive

Compare-and-Swap (CAS) compares the contents of a memory location with a given value and, only if they are the same, updates the contents of that memory location to a new given value

```
1 bool CAS(word* loc, word oldval, word newval) {  
2     atomic { // Code block will execute atomically  
3         res := (*loc == oldval);  
4         if (res)  
5             *loc := newval;  
6     }  
7     return res;  
8 }
```

Compare-and-Swap (CAS) Primitive

- CAS is implemented as the compare-and-exchange (CMPXCHG) instruction in x86 architectures
 - ▶ On a multiprocessor, the LOCK prefix must be used
- CAS is a popular synchronization primitive for implementing both lock-based and nonblocking concurrent data structures

```
1      xor %ecx, %ecx ; ecx=0
2      inc %ecx ; ecx=1
3  RETRY: xor %eax, %eax ; eax=0
4          lock cmpxchg %ecx, &lk
5          jnz RETRY
6          ret
7
```

```
1  void spinLock(lock* lk) {
2      // flg attribute is set when
3      // the lock is acquired
4      while (CAS(&lk->flg, 0, 1) == 1) {
5          // Keep spinning
6      }
7  }
```

Compare-and-Swap (CAS) Primitive

- CAS is implemented as the compare-and-exchange (CMPXCHG) instruction in x86 architectures
 - ▶ On a multiprocessor, the LOCK prefix must be used
- CAS is a popular synchronization primitive for implementing both lock-based and nonblocking concurrent data structures

```
1  xor %ecx, %ecx ; ecx=0
2  inc %ecx ; ecx=1
3  RETRY: xor %eax, %eax ; eax=0
4  lock cmpxchg %ecx, &lk
5  jnz RETRY
6  ret
7
```

```
1  void spinLock(lock* lk) {
2      // flg attribute is set when
3      // the lock is acquired
4      while (CAS(&lk->flg, 0, 1) == 1) {
5          // Keep spinning
6      }
```

How can you implement
fetch_and_xyz() with CAS?

Spin Lock with CAS on X86_64

```
static inline bool cas32(volatile uint32_t *addr, uint32_t oldVal, uint32_t newVal) {
    unsigned char result;
    asm volatile(
        "lock cmpxchgl %3, %0 \n sete %1"
        : "+m"(*addr), "=q"(result), "+a"(oldVal) // loads oldVal into EAX
        : "r"(newVal)
        : "memory");
    return result;
}

uint32_t lock = 0;

void acquire() {
    while (!cas32(&lock, 0, 1)) { } // add delay to reduce contention
}

static inline void release() {
    asm volatile("" ::: "memory"); // compiler barrier
    lock = 0;
}
```

Load Linked (LL)/Store Conditional (SC) Instructions

LL/SC

POWER, MIPS, ARM

```
1 word LL(word* a):  
2     atomic {  
3         remember a;  
4         return *a;  
5     }  
6  
7 bool SC(word* a, word w):  
8     atomic {  
9         res := (a is remembered, and has not been evicted  
10              since LL)  
11         if (res)  
12             *a = w;  
13         return res;  
14     }
```

Load Linked (LL)/Store Conditional (SC) Instructions

LL/SC

POWER, MIPS, ARM

```
1 word LL(word* a):  
2     atomic {  
3         remember a;  
4         return *a;  
5     }  
6  
7 bool SC(word* a, word w):  
8     atomic {  
9         res := (a is remembered and has not been evicted)  
10  
11         if (res == false) {  
12             fetch_and_func() with LL/SC?  
13         }  
14         *a = w;  
15         return res;  
16     }
```

How can you implement
fetch_and_func() with LL/SC?

ABA Problem

Nonblocking Algorithms

Blocking algorithms (e.g., lock-based concurrent data structures)

Failure or delay of any one thread can delay other threads

- Use of locks can lead to deadlocks, livelocks, and priority inversion
- Blocked threads do not do useful work, problematic for high-priority or real-time applications
- Getting the right degree of concurrency and correctness with locks is challenging

Failure or delay of one thread cannot delay other threads in nonblocking algorithms

- Provides different progress guarantees: wait-freedom and lock-freedom
- Use RMW instructions like CAS or LL/SC for mutual exclusion
- Eliminate locks altogether

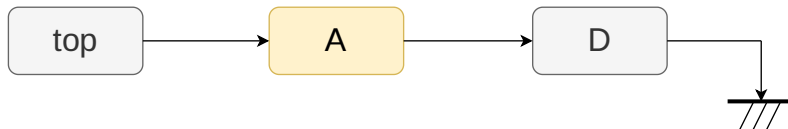
Lock-free Stack Data Structure

push

```
1 void push(node** top, node* new):  
2   node* old  
3   repeat  
4     old := *top  
5     new->next := old  
6   until CAS(top, old, new)  
7  
8
```

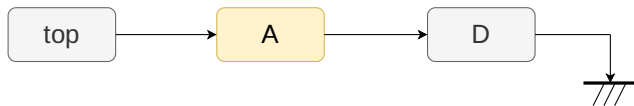
pop

```
1 node* pop(node** top):  
2   node* old, new  
3   repeat  
4     old := *top  
5     if old = null return null  
6     new := old->next  
7   until CAS(top, old, new)  
8   return old
```



Concurrent Modifications to a Lock-free Stack

Thread 1 is executing pop(A)

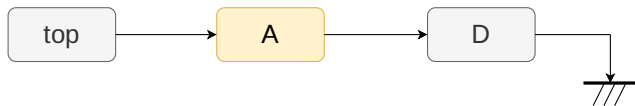


Thread 1 sees top points to A, but gets delayed while executing pop(A)

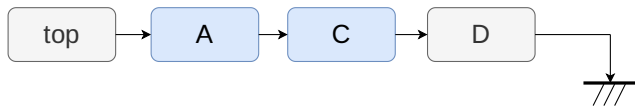
Assume that deleted nodes can be reused

Concurrent Modifications to a Lock-free Stack

1 Thread 1 is executing pop(A)

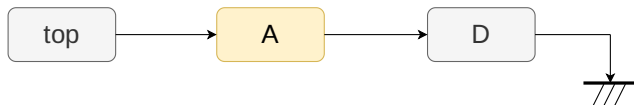


2 Other threads execute pop(A), push(C), and push(A)

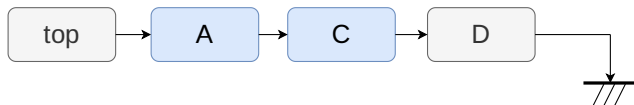


ABA Problem

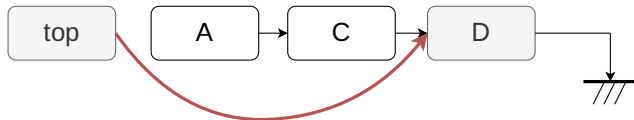
1 Thread 1 is executing pop(A)



2 Other threads execute pop(A), push(C), and push(A)



3 Thread 1's CAS succeeds



Avoiding ABA Problem

- Common workaround is to add extra “tag” to the memory address being compared
 - ▶ Tag can be a counter that tracks the number of updates to the reference
 - ▶ Can steal lower order bits of memory address or use a separate tag field if 128-bit CAS is available
- LL/SC does not suffer from the ABA problem because it checks whether a value has changed in between the interval, rather than comparing the value itself

Scalable Spin Locks

Spin Lock with TAS

```
1 class SpinLock {  
2     bool loc = false;  
3     public void lock() {  
4         while (TAS(&loc)) {  
5             // spin  
6         }  
7     }  
8     public void unlock() {  
9         loc = false;  
10    }  
11 }
```

How can we improve the performance of TAS-based spinlocks?

Test-And-Test-And-Set

- Keep reading the memory location till the location appears unlocked
- + Reduces bus traffic—why?

```
1 do {  
2   while (TATAS_GET(loc)) {}  
3 } while (TAS(loc));
```

With n threads contending for a critical section, the time per acquire-release pair is $\mathcal{O}(n)$

Spin Lock with TAS and Exponential Backoff

Adapt when to retry to reduce contention

- For example, increase the backoff with the number of unsuccessful retries (implies high contention)

```
1  class SpinLock {  
2      bool loc = false;  
3      const int MIN = ..., MUL = ..., MAX = ...;  
4      public void unlock() {  
5          loc = false;  
6      }  
7      public void lock() {  
8          int backoff = MIN;  
9          while (TAS(&loc)) {  
10             pause(backoff);  
11             backoff = min(backoff * MUL, MAX);  
12         }  
13     }  
14 }
```

Challenges with Exponential Backoff

- Adapt when to retry to reduce contention
 - ▶ For example, increase the backoff with the number of unsuccessful retries (implies high contention)

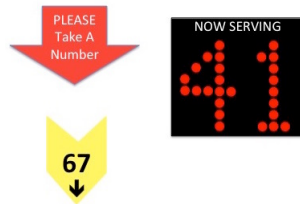
What can be some problems?

Challenges with Exponential Backoff

- Adapt when to retry to reduce contention
 - ▶ For example, increase the backoff with the number of unsuccessful retries (implies high contention)
- Critical section is potentially underutilized
- Avoid concurrent threads getting into a lockstep, backoff for a random duration, possibly doubling each time till a given maximum
- Best-performing constants depend on the host machine and the application

Ticket Lock

- TAS-based locks are unfair
- Ticket lock grants access to threads based on FCFS

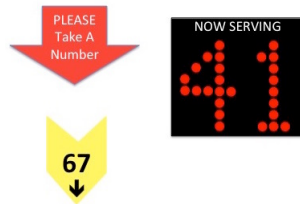


```
1 class TicketLock implements Lock {  
2     int nxt_tkt = 0;  
3     int serving = 0;  
4     public void unlock() {  
5         serving++;  
6     }
```

```
7     public void lock() {  
8         int my_tkt = FAI(&nxt_tkt);  
9         while (serving != my_tkt) {}  
10    }  
11 }  
12
```

Ticket Lock

- TAS-based locks are unfair
- Ticket lock grants access to threads based on FCFS



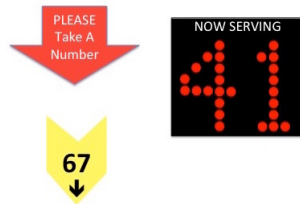
```
1 class TicketLock implements Lock {  
2     int nxt_tkt = 0;  
3     int serving = 0;  
4     public void unlock() {  
5         serving++;  
6     }
```

```
7     public void lock() {  
8         int my_tkt = FAI(&nxt_tkt);  
9         while (serving != my_tkt) {}  
10    }  
11 }
```

How is this different from Bakery's algorithm?

Ticket Lock

- TAS-based locks are unfair
- Ticket lock grants access to threads based on FCFS



```
1 class TicketLock implements Lock {  
2     int nxt_tkt = 0;  
3     int serving = 0;  
4     public void unlock() {  
5         serving++;  
6     }
```

What are some disadvantages of ticket lock?

```
7     public void lock() {  
8         int my_tkt = FAI(&nxt_tkt);  
9         while (serving != my_tkt) {}  
10    }  
11 }  
12 }
```

Queued Locks

Key Idea

- Instead of contending on a single “serving” variable, make threads wait in a queue (i.e., FCFS)
- Each thread knows its order in the queue

Implementation

- Use an array-based queue
 - ▶ Statically or dynamically allocated depending on the number of threads
- Each thread spins on its own lock (i.e., array element), and knows the successor information

Array-based Queued Lock

```
1 public class ArrayLock {  
2     AtomicInteger tail;  
3     volatile boolean[] flag;  
4     ThreadLocal<Integer> mySlot = ...;  
5     public ArrayLock(int size) {  
6         tail = new AtomicInteger(0);  
7         flag = new boolean[size];  
8         flag[0] = true;  
9     }  
10  
11 }
```

```
12 public void lock() {  
13     int slot = FAI(tail);  
14     mySlot.set(slot);  
15     while (!flag[slot]) {}  
16 }  
17 public void unlock() {  
18     int slot = mySlot.get();  
19     flag[slot] = false;  
20     flag[slot+1] = true;  
21 }  
22 }
```

- + Provides fairness
- + Invalidation traffic lower than Ticket lock

Array-based Queued Lock

```
1 public class ArrayLock {  
2     AtomicInteger tail;  
3     volatile boolean[] flag;  
4     ThreadLocal<Integer> mySlot = ...;  
5     public ArrayLock(int size) {  
6         tail = new AtomicInteger(0);  
7         flag = new boolean[size];  
8         flag[0] = true;  
9     }  
10  
11
```

```
12     public void lock() {  
13         int slot = FAI(tail);  
14         mySlot.set(slot);  
15         while (!flag[slot]) {}  
16     }  
17     public void unlock() {  
18         int slot = mySlot.get();  
19         flag[slot] = false;  
20         flag[slot+1] = true;  
21     }  
22 }
```

What could be a few disadvantages of array-based Queued locks?

MCS Queue Lock

MCS Queue lock is the state-of-art scalable FCFS lock

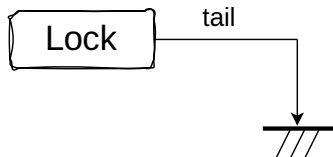
- Uses linked lists instead of arrays
- + Space required to support n threads and k locks: $\mathcal{O}(n + k)$

MCS Queue Lock

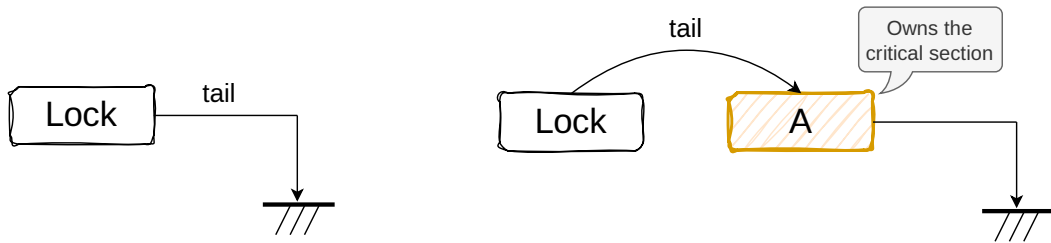
```
1  class QNode {
2      QNode next;
3      bool waiting;
4  }
5
6  public class MCSLock {
7      Node tail = null;
8      ThreadLocal<QNode> myNode = ...;
9      public void lock() {
10         QNode node = myNode.get();
11         QNode prev = swap(tail, node);
12         if (prev != null) {
13             node.waiting = true;
14             prev.next = node;
15             while (node.waiting) {}
16         }
17     }
```

```
18     public void unlock() {
19         QNode node = myNode.get();
20         QNode succ = node.next;
21         if (succ == null)
22             if (CAS(tail, node, null))
23                 return;
24         do {
25             succ = node.next;
26         } while (succ == null);
27         succ.waiting = false;
28     }
29 }
30
31
32
33
34
```

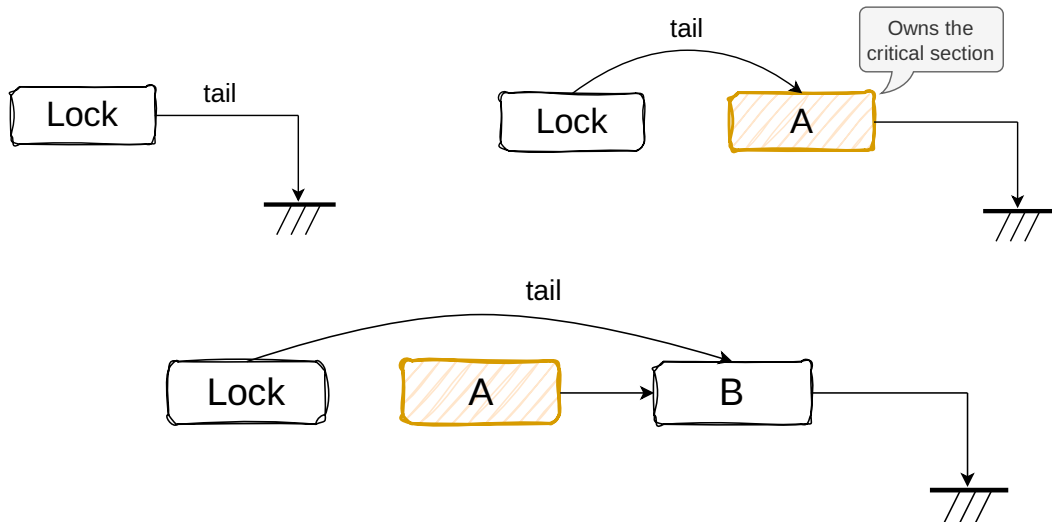
MCS Lock Operations



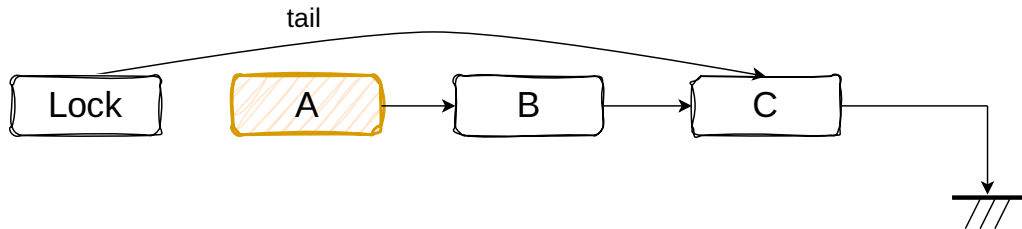
MCS Lock Operations



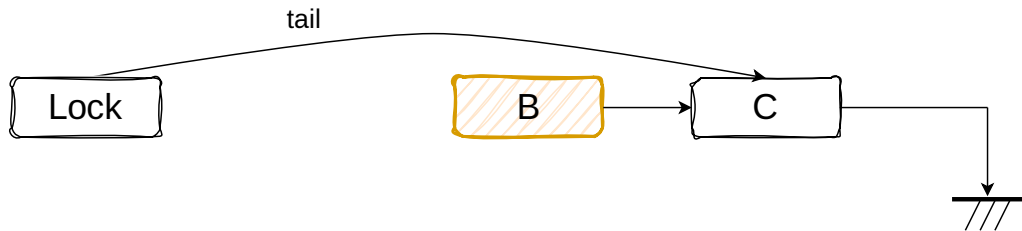
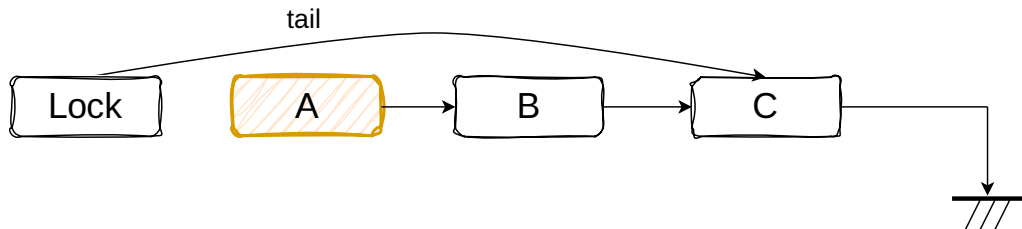
MCS Lock Operations



MCS Lock Operations



MCS Lock Operations



Properties of MCS Lock

- Threads acquire the lock in FCFS manner
- Minimizes false sharing and resource contention
- Threads joining a lock's wait queue is **wait-free**
 - ▶ Wait-freedom implies every operation has a bound on the number of steps it will take before the operation completes
 - Wait-freedom is the strongest non-blocking guarantee of progress
 - ▶ Guaranteed system-wide progress implies lock-freedom, allows individual threads to starve
 - Lock-free implies “locking up” the application in some way (e.g., deadlock and livelock), does not **only** imply absence of synchronization locks

Miscellaneous Lock Optimizations

Reentrant Locks

Reentrant locks can be re-acquired by the owner thread without causing a deadlock

- Freed after an equal number of releases

```
1 public class ParentWidget {  
2     public synchronized void doWork() {  
3         ...  
4     }  
5 }  
6  
7 public class ChildWidget extends ParentWidget {  
8     public synchronized void doWork() {  
9         ...  
10        super.doWork();  
11        ...  
12    }  
13 }
```

Lazy Initialization In Single-Threaded Context

A variable may require the initialization to be synchronized but future uses may be read-only

```
1 class Foo {  
2     private Helper helper = null;  
3  
4     public Helper getHelper() {  
5         if (helper == null)  
6             helper = new Helper();  
7         return helper;  
8     }  
9     ...  
10 }
```

Correct for single-threaded execution,
what could go wrong with multiple
threads?

Lazy Initialization In Multi-Threaded Context

```
1 class Foo {  
2     private Helper helper = null;  
3  
4     public Helper getHelper() {  
5         if (helper == null)  
6             helper = new Helper();  
7         return helper;  
8     }  
9     ...  
10 }
```

```
1 class Foo {  
2     private Helper helper = null;  
3  
4     public synchronized Helper getHelper() {  
5         if (helper == null)  
6             helper = new Helper();  
7         return helper;  
8     }  
9     ...  
10 }
```

The “Double-Checked Locking is Broken” Declaration

Lazy Initialization In Multi-Threaded Context

```
1 class Foo {  
2     private Helper helper = null;  
3  
4     public Helper getHelper() {  
5         if (helper == null)  
6             helper = new Helper();  
7         return helper;  
8     }  
9     ...  
10 }
```

```
1 class Foo {  
2     private Helper helper = null;  
3  
4     public synchronized Helper getHelper() {  
5         if (helper == null)  
6             helper = new Helper();  
7         return helper;  
8     }  
9     ...  
10 }
```

Synchronizes even after helper has been allocated.
Can we optimize the initialization pattern?

Double-Checked Locking: Possible Idea

- (i) Check if helper is initialized
 - ▶ If yes, return
 - ▶ If no, then obtain a lock
- (ii) Double check whether helper has been initialized
 - ▶ If yes, return
 - ▶ If no, initialize helper, return

```
1  class Foo {  
2      private Helper helper = null;  
3  
4      public Helper getHelper() {  
5          if (helper == null) {  
6              synchronized (this) {  
7                  if (helper == null)  
8                      helper = new Helper();  
9              }  
10         }  
11         return helper;  
12     }  
13     ...  
14 }
```

Broken Usage of Double-Checked Locking

```
1 class Foo {  
2     private Helper helper = null;  
3  
4     public Helper getHelper() {  
5         if (helper == null) {  
6             synchronized (this) {  
7                 if (helper == null)  
8                     helper = new Helper();  
9             }  
10        }  
11        return helper;  
12    }  
13    ...  
14 }
```

- The writes inside the constructor call of `Helper()` and to the field `helper` (line 8) can get reordered
- The constructor might be inlined, and the compiler could then reorder all the stores
- A partially created object may then become visible to other threads
- Even the hardware can reorder the stores

Double-Checked Locking: Broken Fix

```
1 public Helper getHelper() {  
2     if (helper == null) {  
3         Helper h;  
4         synchronized (this) {  
5             h = helper;  
6             if (h == null) {  
7                 synchronized (this) {  
8                     h = new Helper();  
9                 }  
10            }  
11            helper = h;  
12        }  
13    }  
14    return helper;  
15 }
```

- A release operation prevents operations from moving out of the critical section
- A release operation does not prevent `helper = h` (line 11) from being moved up (i.e., pulled into the critical section)

Correct Use of Double-Checked Locking

```
1  class Foo {  
2      private volatile Helper helper = null;  
3      public Helper getHelper() {  
4          if (helper == null) {  
5              synchronized (this) {  
6                  if (helper == null)  
7                      helper = new Helper();  
8              }  
9          }  
10         return helper;  
11     }  
12 }
```

Other possibilities are to use barriers in both the writer thread (the thread that initializes helper) and all reader threads

Readers-Writer Locks

Many objects are read concurrently and updated only a few times

Reader lock No thread holds the write lock

Writer lock No thread holds the reader or writer locks

```
1 public interface RWLock {  
2     public void readerLock();  
3     public void readerUnlock();  
4     public void writerLock();  
5     public void writerUnlock();  
6 }
```

Design Choices in Readers-Writer Locks

Release preference order Writer releases lock, both readers and writers are queued up

Incoming readers Writers waiting, and new readers are arriving

Downgrading Can a thread acquire a read lock without releasing the write lock?

Upgrading Can a read lock be upgraded to a write lock safely?

Readers-Writer Lock With Reader-Preference

Reader or writer preference impacts degree of concurrency

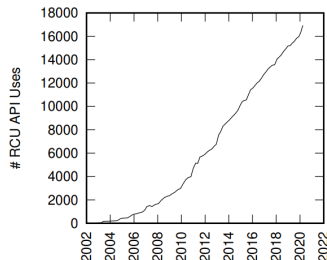
- Allows starvation of non-preferred threads

```
1 readerLock():  
2     acquire(rd)  
3     rdrs++  
4     if rdrs == 1:  
5         acquire(wr)  
6     release(rd)  
7  
8 readerUnlock():  
9     acquire(rd)  
10    rdrs--  
11    if rdrs == 0:  
12        release(wr)  
13    release(rd)
```

```
14 writerLock():  
15     acquire(wr)  
16  
17 writerUnlock():  
18     release(wr)  
19  
20  
21  
22  
23
```

RCU Locks

- Think about data structures that are mostly read, occasionally written
 - Manipulating the read counter with atomic operations is still expensive, the cost dominates performance for short critical sections
- **Idea:** Carefully update the data structure so that readers see a consistent view of data, only writers require locks
- RCU supports concurrency between a single updater and multiple readers
 - ▶ RCU is wait-free for readers and lock-free for the writer
- RCU was introduced in the Linux kernel in 2002, and is actively used throughout the kernel



What is RCU, Fundamentally?

What is RCU? – “Read, Copy, Update”

- Readers can enter a CS with just a compiler or memory barrier
- Writers cannot modify data in place, instead the write operation is split
 - (i) make a copy of the data,
 - (ii) update the local copy, and
 - (iii) publish the updated copy atomically with release semantics
- Readers can read potentially stale data
- The old data can be reclaimed only **after** all readers have completed

Using RCU Locks

Writer Thread

```
1 struct foo {  
2     int a, b, c;  
3 };  
4 struct foo *gp = NULL;  
5  
6 /* . . . */  
7  
8 p = kmalloc(sizeof(*p), GFP_KERNEL);  
9 p->a = 1;  
10 p->b = 2;  
11 p->c = 3;  
12 // cannot simply use gp = p;  
13 rcu_assign_pointer(gp, p);
```

Reader Thread

```
1 // Reader-side CS  
2  
3 // disable preemption  
4 rcu_read_lock();  
5 p = rcu_dereference(gp);  
6 // cannot simply use p = gp;  
7 if (p != NULL) {  
8     do_something(p->a, p->b, p->c);  
9 }  
10 // enable preemption  
11 rcu_read_unlock();
```

Lock Implementations in a JVM

All objects in Java are potential locks

Recursive lock lock can be acquired multiple times by the owner

Thin lock spin lock used when there is no contention, inflated to a fat lock on contention

Fat lock lock is contended or is waited upon, maintains a list of contending threads

Asymmetric Locks

Often objects are accessed by most by one thread but require synchronization for (i) occasional accesses by different threads or (ii) for potential parallelization in the future

Biased locks

- JVMs use biased locks, the acquire/release operations on the owner threads are cheaper
- Usually biased to the first owner thread
- Synchronize only when the lock is contended, need to take care of several subtle issues
- `-XX:+UseBiasedLocking` in HotSpot JVM

Monitors

Using Locks to Access a Bounded Queue

- Consider a bounded FIFO queue
- Many producer threads and one consumer thread access the queue

```
1 mutex.lock();  
2 try {  
3     queue.enq(x);  
4 } finally {  
5     mutex.unlock();  
6 }
```

What are potential challenges?

Using Locks to Access a Bounded Queue

- Consider a bounded FIFO queue
- Many producer threads and one consumer thread access the queue

```
1 mutex.lock();  
2 try {  
3     queue.enq(x);  
4 } finally {  
5     mutex.unlock();  
6 }
```

- Producers and consumers need to know about the size of the queue
- Every producer and consumer need to follow the locking convention
- The design may evolve: there can be multiple queues along with new producers and consumers

Monitors to the Rescue!

- Combination of methods, mutual exclusion locks, and condition variables
- Provides mutual exclusion for methods and the possibility to wait for a condition
 - ▶ Condition variables in monitors have an associated wait queue
 - ▶ Operations: wait, notify (or signal), and notifyAll (or broadcast)

```
1 public synchronized void enq() {  
2     que.enq(x);  
3 }
```

Condition Variables in Monitors

wait condVar, mtx

- Make the thread wait until a condition encoded by condVar is true
 - (i) Releases the monitor's mutex mtx
 - (ii) Moves the thread to condVar's wait queue
 - (iii) Puts the thread to sleep
- Steps (i)–(iii) are atomic to prevent race conditions
- When the thread wakes up, it is assumed to hold mtx

notify condVar

- Invoked by a thread to assert that condition encoded by condVar is true
- Moves one or more threads from the wait queue to the ready queue

notifyAll condVar

Moves all threads from the wait queue to the ready queue

Implementing a Monitor

Assume there is a mutex `mtx`, a condition variable `condVar` and an associated queue `condVarQ`, and an queue entry `entryQ`

monitorenter:

```
enter the method
if mtx is locked
    add this thread to entryQ
    block this thread
else
    lock mtx
```

wait:

```
add this thread to condVarQ
schedule
block this thread
```

monitorexit:

```
// choose next thread to run
schedule
unlock mtx
return from method
```

notify:

```
if condVarQ  $\neq \emptyset$ 
    remove one thread t from condVarQ
    add this thread to entryQ
    restart t
// t will occupy the monitor next
block this thread
```

Signaling Policies

There is a conflict between the signaling and signaled processes for access to the monitor

Signal and wait (SW)

- Signaling thread needs to reacquire the lock
- Signaled thread can continue execution

Signal and urgent wait (SU)

- Like SW, but signaling thread gets to go after the signaled thread
- Also called Hoare-style monitors

Signal and continue (SC)

- Signaling thread continues to hold the lock
- Java implements SC only

Signal and exit (SX)

- Signaling thread exits, signaled thread can continue execution

Bounded Buffers with Spin Locks

```
1 Queue q;  
2  
3 Mutex mtx; // protect q
```

```
3 producer:  
4 while true:  
5     data = new Data(...);  
6     acquire(mtx);  
7     while q.isFull():  
8         release(mtx);  
9         // wait  
10    acquire(mtx);  
11    q.enq(data);  
12    release(mtx);
```

```
13 consumer:  
14 while true:  
15     acquire(mtx);  
16     while q.isEmpty():  
17         release(mtx);  
18         // wait  
19     acquire(mtx);  
20     data = q.deq();  
21     release(mtx);  
22
```

Bounded Buffers with Monitors

```
1 Queue q;  
2  
3 Mutex mtx; // protect q  
4 CondVar empty, full;
```

```
5 producer:  
6 while true:  
7     data = new Data(...);  
8     acquire(mtx);  
9     while q.isFull():  
10         wait(full, mtx);  
11     q.enq(data);  
12     notify(empty);  
13     release(mtx);
```

```
14 consumer:  
15 while true:  
16     acquire(mtx);  
17     while q.isEmpty():  
18         wait(empty, mtx);  
19     data = q.deq();  
20     notify(full);  
21     release(mtx);  
22
```

References



M. Herlihy et al. The Art of Multiprocessor Programming. Chapters 1, 2, 7–8, 2nd edition, Morgan Kaufmann.



M. L. Scott and T. Brown. Shared-Memory Synchronization. Chapters 1–7, 2nd edition, Springer Cham.



Jeff Preshing. Locks Aren't Slow; Lock Contention Is.