

Das C++- Speichermodell

Rainer Grimm

Schulungen, Coaching und Technologieberatung

Multithreading mit C++11

C++'s Antwort auf die Anforderungen der Multicore-Architekturen.



Ein definiertes Speichermodell

- Atomare Operationen
- Partielle Ordnung von Operationen
- Sichtbare Effekte von Operationen

Eine standardisierte Threading-Schnittstelle

- Threads und Tasks
- Schutz und sicheres Initialisieren der Daten
- Thread-lokale Daten
- Synchronisation der Threads

Das C++-Speichermodell

Der Vertrag

Atomare Datentypen

Synchronisations- und Ordnungsbedingungen

Singleton Pattern

Sukzessive Optimierung

Das C++-Speichermodell

Der Vertrag

Atomare Datentypen

Synchronisations- und Ordnungsbedingungen

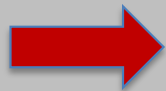
Singleton Pattern

Sukzessive Optimierung

Der Vertrag



- Entwickler respektiert die Regeln
 - atomare Operationen
 - Partielle Ordnung von Operationen
 - Speichersichtbarkeit
- System besitzt die Freiheit zu optimieren
 - Compiler
 - Prozessor
 - Speicherebenen



Hoch optimiertes Programm, das auf die Plattform zugeschnitten ist.

Der Vertrag

streng



locker

- Ein Kontrollfluß

- Tasks
- Threads
- Bedingungsvariablen

- Sequenzielle Konsistenz
- Acquire-Release-Semantik
- Relaxed-Semantik

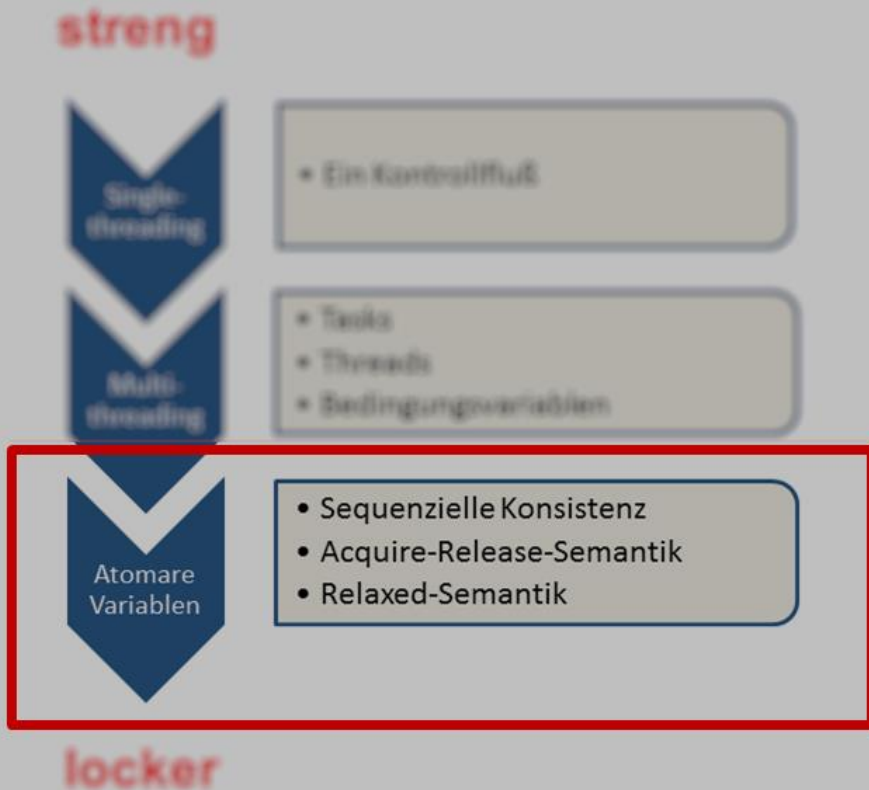
- Mehr Optimierungspotential für das System
- Anzahl der möglichen Kontrollflüsse steigt exponentiell
- Zunehmend ein ausschließliches Gebiet für Domänenexperten
- Bruch der Intuition
- Feld für Mikrooptimierung

Der Vertrag

- Sequenzielle Konsistenz
 - *Strong Memory Model*
 - Globaler Zeitgeber
 - Speichermodell für Java und C#

Bruch der Sequenziellen Konsistenz

- Acquire-Release-Semantik
 - Synchronisation zwischen Threads mit atomaren Operationen
- Relaxed-Semantik
 - *Weak Memory Model*
 - Schwache Zusicherungen



Das C++-Speichermodell

Der Vertrag

Atomare Datentypen

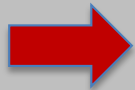
Synchronisations- und Ordnungsbedingungen

Singleton Pattern

Sukzessive Optimierung

Atomare Datentypen

Atomare Variablen sind die Grundlage für das C++-Speichermodell.



Atomare Operationen auf atomare Variablen definieren die Synchronisations- und Ordnungsbedingungen.

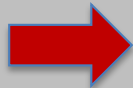
- Synchronisations- und Ordnungsbedingungen gelten für atomare Variablen und nicht atomare Variablen.
- Synchronisations- und Ordnungsbedingungen werden von höheren Abstraktionen verwendet.
 - Threads und Tasks
 - Mutexe und Locks
 - Bedingungsvariablen
 - ...

Atomar: std::atomic_flag

Das atomare Flag `std::atomic_flag`

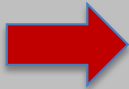
- bietet ein sehr einfaches Interface an.
 - `clear` und `test_and_set`

- ist die einzige lockfreie Datenstruktur.



Alle weiteren atomaren Datentypen für integrale Typen, Zeiger und eigene Datentypen können intern Locks verwenden.

- ist der elementarer Baustein für höhere Abstraktionen.



Spinlock

Atomar: std::atomic_flag

Spinlock

```
class Spinlock{
    std::atomic_flag  flag;
public:
    Spinlock():flag(ATOMIC_FLAG_INIT){}

    void lock(){
        while(flag.test_and_set());
    }

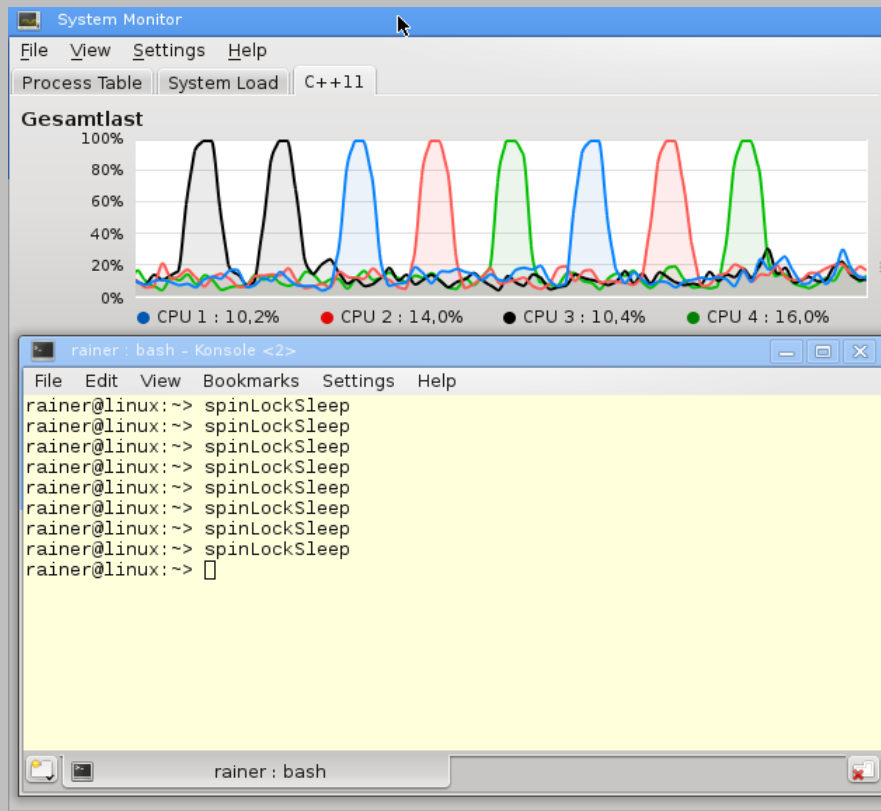
    void unlock(){
        flag.clear();
    }
};
```

Locken einer Ressource

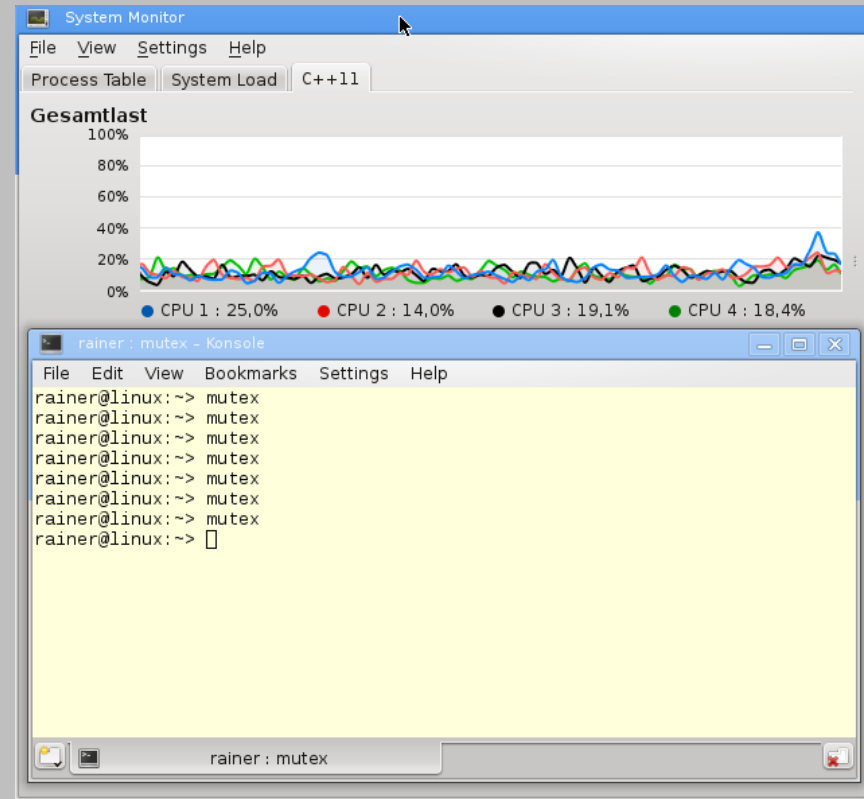
```
Spinlock spin;
// Mutex spin;
void workOnResource(){
    spin.lock();
    sleep_for(seconds(2));
    spin.unlock();
}
int main({
    thread t(workOnResource);
    thread t2(workOnResource);
    t.join();
    t2.join();
}
```

Atomar: std::atomic_flag

Spinlock

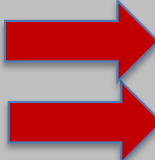


Mutex



Atomar: `std::atomic<bool>`

Der atomare Wahrheitswert `std::atomic<bool>`

- lässt sich explizit auf `true` oder `false` setzen.
- unterstützt die Funktion `compare_exchange_strong`.
 - Fundamentale Funktion für atomare Operationen.
 - Vergleicht und setzt einen Wert in einer atomaren Operation.
 - Syntax: `bool compare_exchange_strong(expected, desired)`
 - Strategie: `atom.compare_exchange_strong(exp, des)`


```
*atom == exp      *atom= des
*atom != exp      exp= *atom
```
- lässt sich als Bedingungsvariable verwenden.

Atomar: Bedingungsvariable

```
std::vector<int> mySharedWork;  
std::mutex mutex_  
std::condition_variable condVar;
```

```
bool dataReady;
```

```
void setDataReady() {  
    mySharedWork={1,0,3};  
    {  
        lock_guard<mutex> lck(mutex_  
        dataReady=true;  
    }  
    condVar.notify_one();  
}
```

```
void waitingForWork() {  
    unique_lock<mutex> lck(mutex_  
    condVar.wait(lck, []{return dataReady;});  
    mySharedWork[1]= 2;  
}
```

```
int main() {  
    thread t1(waitingForWork);  
    thread t2(setDataReady);  
  
    t1.join();  
    t2.join();  
  
    for (auto v: mySharedWork) {  
        std::cout << v << " ";  
    } // 1 2 3  
}
```

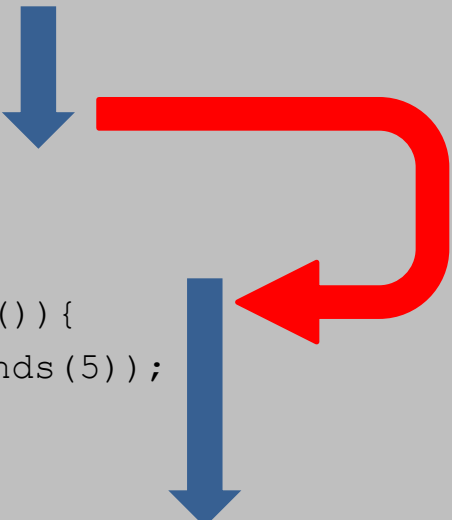
Atomar: std::atomic<bool>

```
std::vector<int> mySharedWork;  
std::atomic<bool> dataReady(false);
```

```
void setDataReady() {  
    mySharedWork={1,0,3};  
    dataReady= true;  
}
```

```
void waitingForWork() {  
    while (!dataReady.load()) {  
        sleep_for(milliseconds(5));  
    }  
    mySharedWork[1]= 2;  
}
```

```
int main() {  
    thread t1(waitingForWork);  
    thread t2(setDataReady);  
    t1.join();  
    t2.join();  
    for (auto v: mySharedWork) {  
        cout << v << " ";  
    }  
    // 1 2 3  
}
```



**sequenced-before
synchronizes-with**

Atomar: std::atomic

Alle weiteren atomaren Datentypen sind teilweise oder vollständige Spezialisierungen von

`std::atomic`.

```
std::atomic<T*>
```

```
std::atomic<Integraler Typ>
```

```
std::atomic<Eigener Typ>
```

- Für eigene Datentypen gelten Einschränkungen
 - Ihr Copy-Zuweisungsoperator und alle ihre Basisklassen müssen trivial sein.
 - Sie dürfen keine virtuellen Methoden und Basisklassen enthalten.
 - Sie müssen bitweise vergleichbar sein.

Atomare Datentypen:

Operation	atomic_flag	atomic<bool>	atomic<T*>	atomic<Integrale Typ>	atomic<Eigener Ty>
test_and_set	ja				
clear	ja				
is_lock_free		ja	ja	ja	ja
load		ja	ja	ja	ja
store		ja	ja	ja	ja
exchange		ja	ja	ja	ja
compare_exchange_weak compare_exchange_strong		ja	ja	ja	ja
fetch_add, += fetch_sub, -=			ja	ja	
fetch_or, = fetch_and, &= fetch_xor, ^=				ja	
++ --			ja	ja	

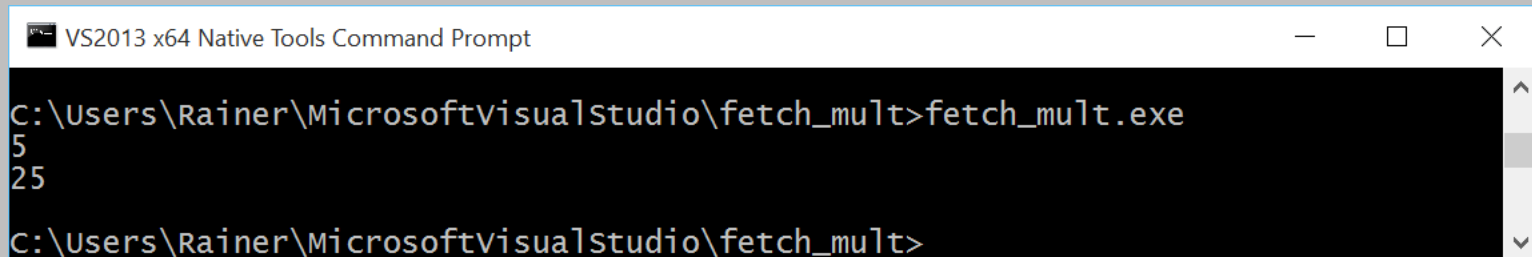


Es gibt keine Multiplikation oder Division.

Atomar: std::atomic

```
template <typename T>
T fetch_mult(std::atomic<T>& shared, T mult){
    T oldValue= shared.load();
    shared.compare_exchange_strong(oldValue, oldValue * mult);
    return oldValue;
}

int main(){
    std::atomic<int> myInt{5};
    std::cout << myInt << std::endl;
    fetch_mult(myInt,5);
    std::cout << myInt << std::endl;
}
```



```
VS2013 x64 Native Tools Command Prompt

C:\Users\Rainer\MicrosoftVisualStudio\fetch_mult>fetch_mult.exe
5
25

C:\Users\Rainer\MicrosoftVisualStudio\fetch_mult>
```

Das C++-Speichermodell

Der Vertrag

Atomare Datentypen

Synchronisations- und Ordnungsbedingungen

Singleton Pattern

Sukzessive Optimierung

Synchronisation und Ordnung

C++ kennt sechs verschiedene Speichermodelle.

```
enum memory_order {  
    memory_order_relaxed,  
    memory_order_consume,  
    memory_order_acquire,  
    memory_order_release,  
    memory_order_acq_rel,  
    memory_order_seq_cst  
};
```

- Per Default gilt die Sequenzielle Konsistenz.
 - Das Speichermodell für C# und Java.
 - `memory_order_seq_cst`
 - Implizites Argument bei atomaren Operationen

```
std::atomic<int> shared;
```

```
shared.load()  $\approx$  shared.load(std::memory_order_seq_cst);
```

Synchronisation und Ordnung

Ordnung in das Speichermodell bringt die Beantwortung zweier Fragen.

1. Für welche atomaren Operationen sind die Speichermodelle konzipiert?
2. Welche Synchronisations- und Ordnungsbedingungen definieren die Speichermodelle?

Synchronisation und Ordnung

1. Für welche atomaren Operationen sind die Speichermodelle konzipiert?

- **read-Operationen:**

`memory_order_acquire` und `memory_order_consume`

- **write-Operationen:**

`memory_order_release`

- **read-modify-write-Operationen:**

`memory_order_acq_rel` und `memory_order_seq_cst`

! `memory_order_relaxed` definiert keine Synchronisations- und Ordnungsbedingungen.

Synchronisation und Ordnung

Operation	read Operationen	write Operationen	read-modify-write Operationen
test_and_set			✓
clear		✓	
is_lock_free	✓		
load	✓		
store		✓	
exchange			✓
compare_exchange_weak compare_exchange_strong			✓
fetch_add, += fetch_sub, -=			✓
fetch_or, = fetch_and, &= fetch_xor, ^=			✓
++ --			✓

```
std::atomic<int> atom;
```

```
atom.load(std::memory_order_acq_rel)
```

```
atom.load(std::memory_order_release)
```



```
atom.load(std::memory_order_acquire)
```

```
atom.load(std::memory_order_relaxed)
```

Synchronisation und Ordnung

2. Welche Synchronisations- und Ordnungsbedingungen definieren die Speichermodelle?

- **Sequenzielle Konsistenz**

- Globale Ordnung auf allen Threads

`memory_order_seq_cst`

- **Acquire-Release-Semantik**

- Ordnung zwischen Lese- und Schreiboperationen der gleichen atomaren Variablen auf verschiedenen Threads

`memory_order_consume`, `memory_order_acquire`,
`memory_order_release` und `memory_order_acq_rel`

- **Relaxed-Semantik**

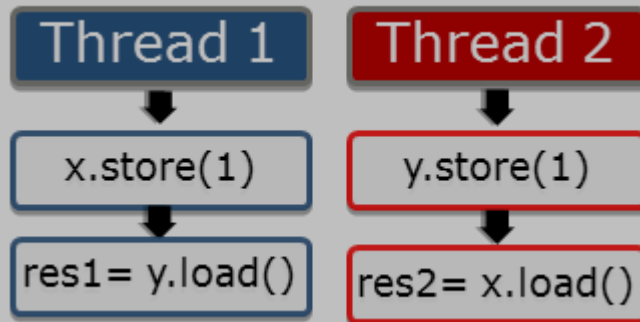
- Keine Synchronisations- oder Ordnungsbedingungen

`memory_order_relaxed`

Synchronisations und Ordnung

Sequenzielle Konsistenz (Leslie Lamport 1979)

1. Die Anweisungen eines Programms werden in der Sourcecodereihenfolge ausgeführt.
2. Es gibt eine globale Reihenfolge aller Operationen auf allen Threads.

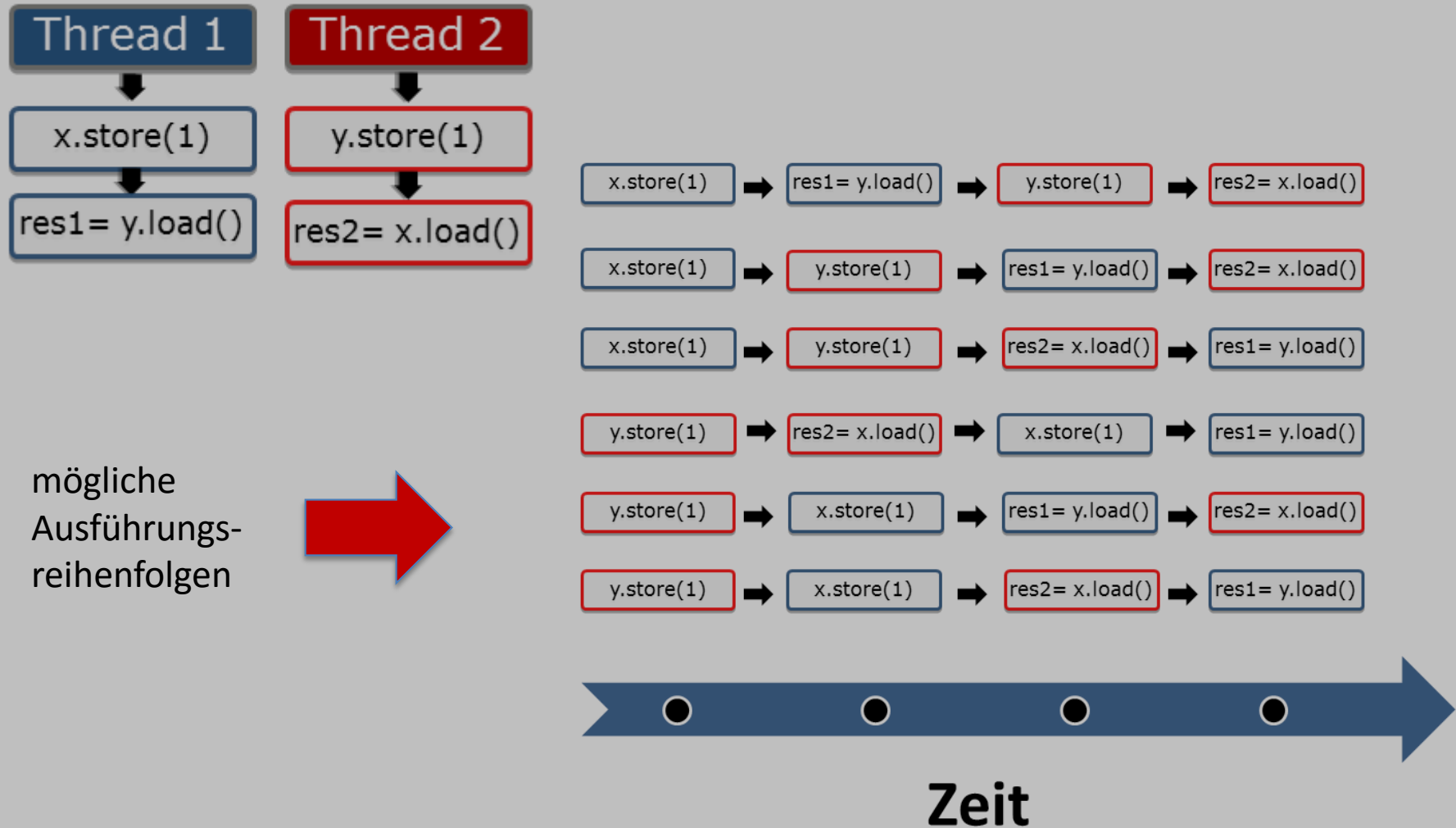


Sequenzielle Konsistenz ergibt



1. Die Befehle werden in der Reihenfolge ausgeführt, in der sie im Sourcecode stehen.
2. Jeder Thread sieht die Operationen jedes anderen Threads in der gleichen Reihenfolge (globaler Zeittakt).

Synchronisations und Ordnung



Synchronisation und Ordnung

Acquire-Release-Semantik

- Eine release-Operation auf einer atomaren Variable synchronisiert sich mit einer acquire-Operation auf der gleichen atomaren Variablen und definiert eine Ordnungsbedingung.
- Acquire-Operation:
 - Lese-Operation (`load` oder auch `test_and_set`)
- Release-Operation:
 - Schreibe-Operation (`store` oder auch `clear`)
- Ordnungsbedingungen:
 - Lese- und Schreiboperationen können nicht **vor** eine acquire-Operation verschoben werden
 - Lese- und Schreiboperationen können nicht **hinter** eine release-Operation verschoben werden

Synchronisation und Ordnung

```
void dataProducer() {  
    mySharedWork={1,0,3};  
    dataProduced.store(true, std::memory_order_release);  
}  
  
void deliveryBoy() {  
    while( !dataProduced.load(std::memory_order_acquire) );  
    dataConsumed.store(true, std::memory_order_release);  
}  
  
void dataConsumer() {  
    while( !dataConsumed.load(std::memory_order_acquire) );  
    mySharedWork[1]= 2;  
}
```



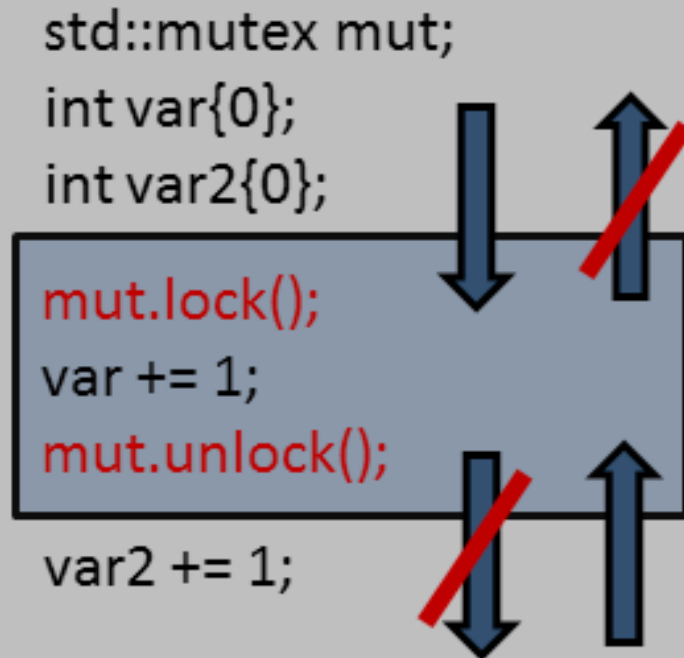
Thread 1

Thread 2

Thread 3

sequenced-before
synchronizes-with

Synchronisation und Ordnung



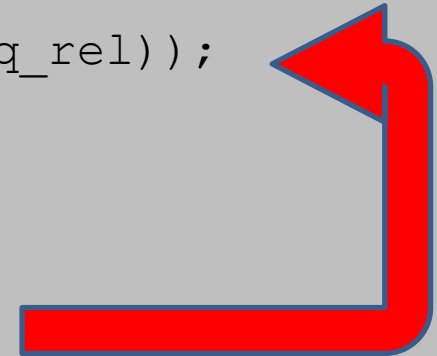
- Acquire-Operationen
 - Locken eines Mutex
 - Warten einer Bedingungsvariable
 - Starten eines Threads
- Release-Operationen
 - Unlocken eines Mutex
 - Benachrichtigung einer Bedingungsvariable
 - join-Aufruf auf einem Thread

Synchronisation und Ordnung

```
class Spinlock{
    std::atomic_flag flag;
public:
    Spinlock(): flag(ATOMIC_FLAG_INIT) {}

    void lock(){
        while(flag.test_and_set(memory_order_acq_rel));
    }

    void unlock(){
        flag.clear(std::memory_order_release);
    }
};
```



Synchronisation und Ordnung

Consume-Release-Semantik

- Consume-Release-Semantik ist die Acquire-Release-Semantik ohne Ordnungsbedingungen.
- Besitzt einen legendären Ruf
 - Sehr verständnisrersistent
 - `std::memory_order_consume` wird durch den Compiler auf `std::memory_order_acquire` abgebildet
 - Kein Compiler implementiert sie (*Ausnahme GCC*)
- Beschäftigt sich mit Datenabhängigkeiten
 - In einem Thread: *carries-a-dependency-to*
 - Zwischen Threads: *dependency-ordered-before*

Synchronisation und Ordnung

```
atomic<string*> ptr;  
int data;  
atomic<int> atoData;
```

```
void producer(){  
    string* p = new string("C++11");  
    data = 2011;  
    atoData.store(14,memory_order_relaxed);  
    ptr.store(p,memory_order_release);  
}
```

```
void consumer(){  
    string* p2;  
    while (!(p2 = ptr.load(memory_order_acquire)));  
    cout << *p2 << " " << data;  
    cout << atoData.load(memory_order_relaxed);  
}
```

```
atomic<string*> ptr,  
int data;  
atomic<int> atoData;
```

```
void producer(){  
    string* p = new string("C++11");  
    data = 2011;  
    atoData.store(14,memory_order_relaxed);  
    ptr.store(p, memory_order_release);  
}
```

```
void consumer(){  
    string* p2;  
    while (!(p2 = ptr.load(memory_order_consume)));  
    cout << *p2 << " " << data;  
    cout << atoData.load(memory_order_relaxed);  
}
```

Synchronisation und Ordnung

```
std::atomic<std::string*> ptr;  
int data;  
std::atomic<int> atoData;
```

```
void producer() {  
    std::string* p = new std::string("C++11");  
    data = 2011;  
    atoData.store(2014, std::memory_order_relaxed);  
    ptr.store(p, std::memory_order_release);  
}
```

```
void consumer() {  
    std::string* p2;  
    while (!(p2 = ptr.load(std::memory_order_consume)));  
    std::cout << "*p2: " << *p2 << std::endl;  
    std::cout << "data: " << data << std::endl;  
    std::cout << "atoData: " << atoData.load(std::memory_order_relaxed) << std::endl;  
}
```

dependency-ordered-before
carries-a-dependency-to



Synchronisation und Ordnung

Relaxed Semantik

- Es gelten keine Synchronisations- und Ordnungsbedingungen. Nur die Atomizität der atomaren Operationen ist gewährleistet.
- Regeln
 - Atomare Operationen mit strengeren Speicherordnungen werden verwendet, um atomare Operationen mit Relaxed-Semantik und nicht atomare Operationen zu ordnen.
 - Operationen in einem Thread werden in Sourcecodeorder ausgeführt (*sequenced-before*).
 - Typische Anwendungsfälle  atomare Zähler (`shared_ptr`)

 Threads können die Operationen in einem anderen Thread in anderer Reihenfolge wahrnehmen.

Synchronisation und Ordnung

```
std::atomic<int> cnt = {0};  
void f(){  
    for (int n = 0; n < 1000; ++n) {  
        cnt.fetch_add(1, std::memory_order_relaxed);  
    }  
}  
int main(){  
    std::vector<std::thread> v;  
    for (int n = 0; n < 10; ++n){  
        v.emplace_back(f);  
    }  
    for (auto& t : v) {  
        t.join();  
    }  
    std::cout << "Final counter value is " << cnt << '\n';  
}
```

Das C++-Speichermodell

Der Vertrag

Atomare Datentypen

Synchronisations- und Ordnungsbedingungen

Singleton Pattern

Sukzessive Optimierung

Singleton

```
mutex myMutex;

class MySingleton{
public:
    static MySingleton& getInstance(){
        lock_guard<mutex> myLock(myMutex);
        if( !instance ) instance= new MySingleton();
        return *instance;
    }
private:
    MySingleton();
    ~MySingleton();
    MySingleton(const MySingleton&)= delete;
    MySingleton& operator=(const MySingleton&)= delete;
    static MySingleton* instance;
};

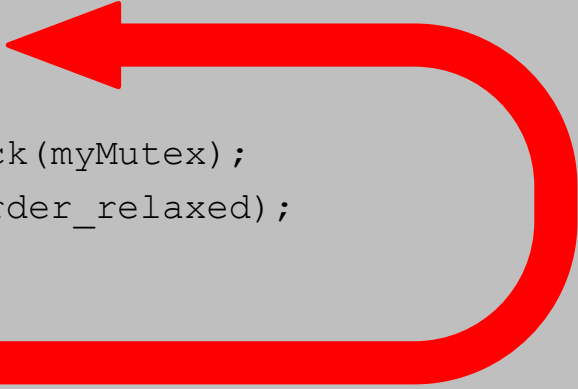
MySingleton::MySingleton()= default;
MySingleton::~~MySingleton()= default;
MySingleton* MySingleton::instance= nullptr;
...
MySingleton::getInstance();
```



Performanzproblem

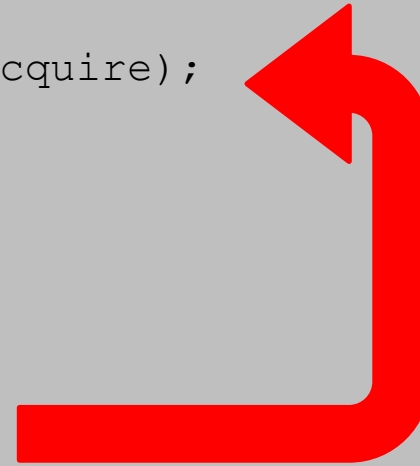
Sequenzielle Konsistenz

```
class MySingleton{
public:
    static MySingleton* getInstance(){
        MySingleton* sin= instance.load();
        if ( !sin ){
            std::lock_guard<std::mutex> myLock(myMutex);
            sin= instance.load(std::memory_order_relaxed);
            if( !sin ){
                sin= new MySingleton();
                instance.store(sin);
            }
        }
        return sin;
    }
private:
    static std::atomic<MySingleton*> instance;
    static std::mutex myMutex;
    . . .
}
```



Acquire-Release-Semantik

```
class MySingleton{
public:
    static MySingleton* getInstance(){
        MySingleton* sin= instance.load(std::memory_order_acquire);
        if ( !sin ){
            std::lock_guard<std::mutex> myLock(myMutex);
            sin= instance.load(std::memory_order_relaxed);
            if( !sin ){
                sin= new MySingleton();
                instance.store(sin,std::memory_order_release);
            }
        }
        return sin;
    }
    . . .
}
```



Meyers Singleton

```
class MySingleton{
public:
    static MySingleton& getInstance(){
        static MySingleton instance;
        return instance;
    }
private:
    MySingleton()= default;
    ~MySingleton()= default;
    MySingleton(const MySingleton&)= delete;
    MySingleton& operator=(const MySingleton&)= delete;
};
```



Setzt Microsoft Visual Studio 2015 voraus.

Singleton: Der Performanztest

```
constexpr auto tenMill= 10'000'000;
class MySingleton{ ... }

std::chrono::duration<double> getTime(){
    auto begin= std::chrono::system_clock::now();
    for ( size_t i= 0; i <= tenMill; ++i) MySingleton::getInstance();
    return std::chrono::system_clock::now() - begin;
};

int main(){
    auto fut1= std::async(std::launch::async,getTime);
    auto fut2= std::async(std::launch::async,getTime);
    auto fut3= std::async(std::launch::async,getTime);
    auto fut4= std::async(std::launch::async,getTime);
    auto total= fut1.get() + fut2.get() + fut3.get() + fut4.get();
    std::cout << total.count() << std::endl;
}
```

Singleton: Der Performanztest

- Wie oft?
 - Je vier Thread greifen 10'000'000 Mal auf ein Singleton zu
- Welche Hardware?
 - Mein PC (Linux) besitzt 4 Kerne, mein Laptop (Windows) 2 Kerne
- Was wird gemessen?
 - Die Zeiten der vier Threads werden addiert
- Welcher Compiler?
 - Aktueller GCC
 - Visual Studio 15
 - Auf beiden Plattformen lief der Test mit maximaler und ohne Optimierung
- Was ist gut?
 - Single-Threaded dient als Referenzwert

Singleton: Der Performanztest

Compiler	Optimierung	Single Threaded	std::lock_guard (Mutex)	Sequenzielle Konsistenz	Acquire-Release Semantik	Meyers Singleton
GCC	nein	0.09	18.15	0.56	0.50	0.10
GCC	ja	0.03	12.47	0.09	0.07	0.04
cl.exe	nein	0.09	23.40	1.33	1.37	0.16
cl.exe	ja	0.02	15.48	0.07	0.07	0.03



Die Details inklusiv der Verwendung von `std::call_once` sind im Artikel [Threadsichers Initialisieren eins Singletons](#).

Singleton: Der Performanztest



- Meine Erkenntnisse
 - Das Singleton Pattern weckt viele Emotionen.
 - Der Compiler optimiert die Aufrufe von `MySingleton::getInstance()` weg.
 - Meyers Singleton ist die schnellste und einfachste Implementierung.

By Watchduck (a.k.a. Tilman Piesk) - Own work, CC BY 3.0, <https://commons.wikimedia.org/w/index.php?curid=10876384>

Das C++-Speichermodell

Der Vertrag

Atomare Datentypen

Synchronisations- und Ordnungsbedingungen

Singleton Pattern

Sukzessive Optimierung

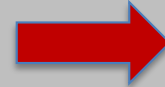
Probleme?

```
int x= 0;  
int y= 0;
```

```
void writing() {  
    x= 2000;  
    y= 11;  
}
```

```
void reading() {  
    cout << "y: " << y << " ";  
    cout << "x: " << x << endl;  
}
```

```
int main() {  
    thread thread1(writing);  
    thread thread2(reading);  
    thread1.join();  
    thread2.join();  
};
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

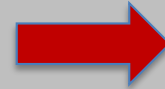
Probleme?

```
int x= 0;
int y= 0;

void writing() {
    x= 2000;
    y= 11;
}

void reading() {
    cout << "y: " << y << " ";
    cout << "x: " << x << endl;
}

int main() {
    thread thread1(writing);
    thread thread2(reading);
    thread1.join();
    thread2.join();
};
```



y	x	z
0	0	
11	0	
0	2000	
11	2000	

Undefined behaviour

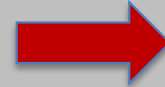
Mutex

```
int x= 0;  
int y= 0;  
mutex mut;
```

```
void writing(){  
    lock_guard<mutex> guard(mut);  
    x= 2000;  
    y= 11;  
}
```

```
void reading(){  
    lock_guard<mutex> guard(mut)  
    cout << "y: " << y << " ";  
    cout << "x: " << x << endl;  
}  
...
```

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

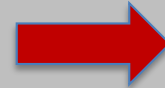
Mutex

```
int x= 0;  
int y= 0;  
mutex mut;
```

```
void writing(){  
    lock_guard<mutex> guard(mut);  
    x= 2000;  
    y= 11;  
}
```

```
void reading(){  
    lock_guard<mutex> guard(mut)  
    cout << "y: " << y << " ";  
    cout << "x: " << x << endl;  
}  
...
```

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	✓
11	0	
0	2000	
11	2000	✓

Probleme?

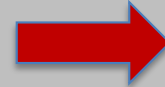
```
volatile x= 0;  
volatile y= 0;
```

```
void writing() {  
    x= 2000;  
    y= 11;  
}
```

```
void reading() {  
    cout << y << " ";  
    cout << x << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

Probleme?

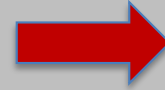
```
volatile x= 0;
volatile y= 0;

void writing() {
    x= 2000;
    y= 11;
}

void reading() {
    cout << y << " ";
    cout << x << endl;
}

...

thread thread1(writing);
thread thread2(reading);
```



y	x	z
0	0	
11	0	
0	2000	
11	2000	

Undefined behaviour

Java versus C++

Java volatile == C++ atomic

- `std::atomic`
 - Schutz der Daten vor gemeinsamen Zugriff mehrerer Threads
- `volatile`
 - Zugriff auf speziellen Speicher, auf dem Lesen und Schreiben nicht optimiert werden darf

Atomare Variablen

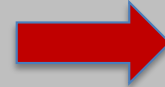
```
atomic<int> x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x.store(2000);  
    y.store(11);  
}
```

```
void reading() {  
    cout << y.load() << " ";  
    cout << x.load() << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

Atomare Variablen

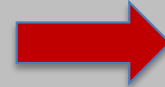
```
atomic<int> x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x.store(2000);  
    y.store(11);  
}
```

```
void reading() {  
    cout << y.load() << " ";  
    cout << x.load() << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	✓
11	0	
0	2000	✓
11	2000	✓

Acquire-Release Semantik

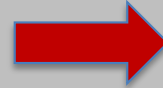
```
atomic<int> x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x.store(2000, memory_order_relaxed);  
    y.store(11, memory_order_release);  
}
```

```
void reading() {  
    cout << y.load(memory_order_acquire) << " ";  
    cout << x.load(memory_order_relaxed) << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

Acquire-Release Semantik

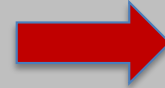
```
atomic<int> x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x.store(2000, memory_order_relaxed);  
    y.store(11, memory_order_release);  
}
```

```
void reading() {  
    cout << y.load(memory_order_acquire) << " ";  
    cout << x.load(memory_order_relaxed) << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	✓
11	0	
0	2000	✓
11	2000	✓

Nicht atomare Variablen

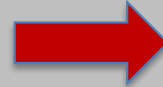
```
int x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x= 2000;  
    y.store(11, memory_order_release);  
}
```

```
void reading() {  
    cout << y.load(memory_order_acquire) << " ";  
    cout << x << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

Nicht atomare Variablen

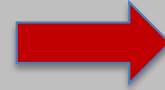
```
int x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x= 2000;  
    y.store(11, memory_order_release);  
}
```

```
void reading() {  
    cout << y.load(memory_order_acquire) << " ",  
    cout << x << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

Undefined behaviour

Relaxed Semantik

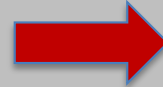
```
atomic<int> x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x.store(2000, memory_order_relaxed);  
    y.store(11, memory_order_relaxed);  
}
```

```
void reading() {  
    cout << y.load(memory_order_relaxed) << " ";  
    cout << x.load(memory_order_relaxed) << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	
11	0	
0	2000	
11	2000	

Relaxed Semantik

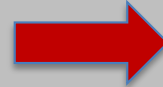
```
atomic<int> x= 0;  
atomic<int> y= 0;
```

```
void writing() {  
    x.store(2000, memory_order_relaxed);  
    y.store(11, memory_order_relaxed);  
}
```

```
void reading() {  
    cout << y.load(memory_order_relaxed) << " ";  
    cout << x.load(memory_order_relaxed) << endl;  
}
```

...

```
thread thread1(writing);  
thread thread2(reading);
```



y	x	Yes
0	0	✓
11	0	✓
0	2000	✓
11	2000	✓

Die Lösung: CppMem

CppMem: Interactive C/C++ memory model

Model
☒ standard ☐ preferred ☐ release_acquire ☐ tot ☐ relaxed_only

Program
 examples/MP_message_passing ▼ MP+na_rel+acq_na.c ▼

Execution
☒ C

```
// MP+na_rel+acq_na
// Message Passing, of data held in non-atomic x,
// with release/acquire synchronisation on y.
// Question: is the read of x required to see the new data value 1
// rather than the initial state value 0?
int main() {
  int x=0; atomic_int y=0;
  {{{ { x=2000;
        y.store(1, memory_order_release); }
    ||| { r1=y.load(memory_order_acquire);
        r2=x; } }}}
  return 0;
}
```

run reset help 8 executions; 2 consistent, only 1 race free

Computed executions

Display Relations

- ☒ sb ☐ asw ☐ dd ☐ cd
☒ rf ☒ mo ☒ sc ☒ lo
☐ hb ☐ vse ☐ ithb ☒ sw ☐ rs ☐ hrs ☒ dob ☐ cad
☒ unsequenced_races ☒ data_races

Display Layout

- ☐ dot ☐ neato_par ☒ neato_par_init ☐ neato_downwards
☐ tex
 edit display options

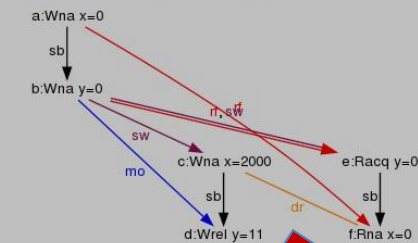
Execution candidate no. 1 of 8

previous consistent previous candidate next candidate next consistent 1 goto

Model Predicates

consistent_race_free_execution = false

- ☒ consistent_execution = true
- ☒ assumptions = true
- ☒ well_formed_threads = true
- ☒ well_formed_rf = true
- ☒ locks_only_consistent_locks = true
- ☒ locks_only_consistent_lo = true
- ☒ consistent_mo = true
- ☒ sc_accesses_consistent_sc = true
- ☒ sc_fenced_sc_fences_headed = true
- ☒ consistent_hb = true
- ☒ consistent_rf = true
- ☒ det_read = true
- ☒ consistent_non_atomic_rf = true
- ☒ consistent_atomic_rf = true
- ☒ coherent_memory_use = true
- ☒ rmw_atomicity = true
- ☒ sc_accesses_sc_reads_restricted = true
- unsequenced_races are absent
- data_races are present
- indeterminate_reads are absent
- locks_only_bad_mutexes are absent

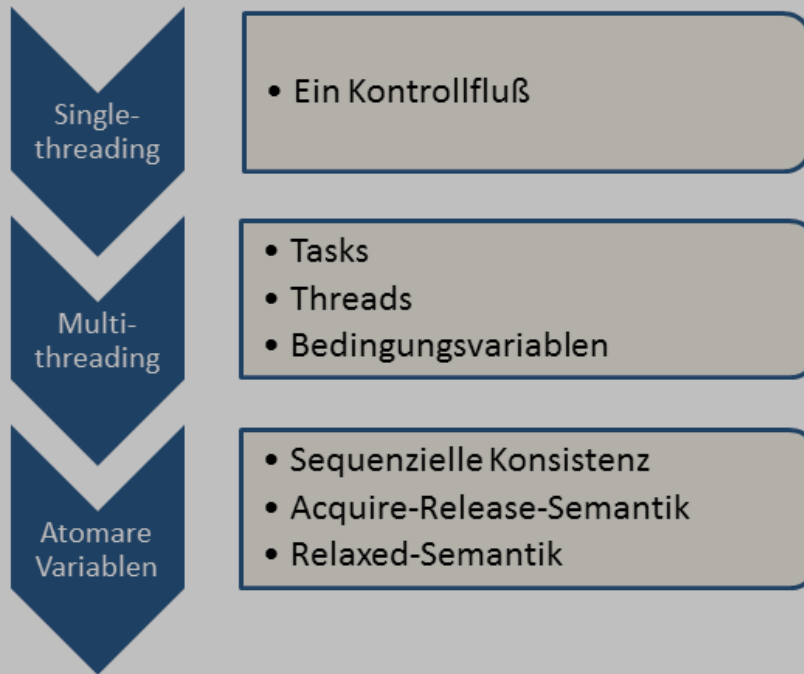


Files: out.exc, out.dot, out.dsp, c

[CppMem](#)

Das C++-Speichermodell

streng



locker

- Mehr Optimierungspotential für das System
- Anzahl der möglichen Kontrollflüsse steigt exponentiell
- Zunehmend ein ausschließliches Gebiet für Domänenexperten
- Bruch der natürlichen Intuition
- Feld für Mikrooptimierung

Weitere Informationen

- **Modernes C++:** Schulungen, Coaching und Technologieberatung durch Rainer Grimm
 - www.ModernesCpp.de
- Blogs zu modernem C++
 - www.grimm-jaud.de (Deutsch)
 - www.ModernesCpp.com (Englisch)
- Kontakt
 - @rainer_grimm
 - schulungen@grimm-jaud.de





```
#include <iostream>

int main(){

    std::cout << "myVec: ";
    std::vector<int> myVec(20);
    std::iota(myVec.begin(), myVec.end(), 0);
    std::cout << "myVec: ";
    for ( auto i: myVec) std::cout << i << " ";
    std::cout << "\n";

    std::function< bool(int)> myBindPred = bind( std::logical_not<>,
    myVec.erase(std::remove_if(myVec.begin(), myVec.end(), myBindPred),
    myVec.end());

    std::cout << "myVec: ";
    for ( auto i: myVec) std::cout << i << " ";
    std::cout << "\n\n";

    std::vector<int> myVec2(20);
    std::iota(myVec2.begin(), myVec2.end(), 0);

    std::cout << "myVec2: ";
    for ( auto i: myVec2) std::cout << i << " ";
    std::cout << "\n\n";
```

Rainer Grimm

Schulungen, Coaching und Technologieberatung