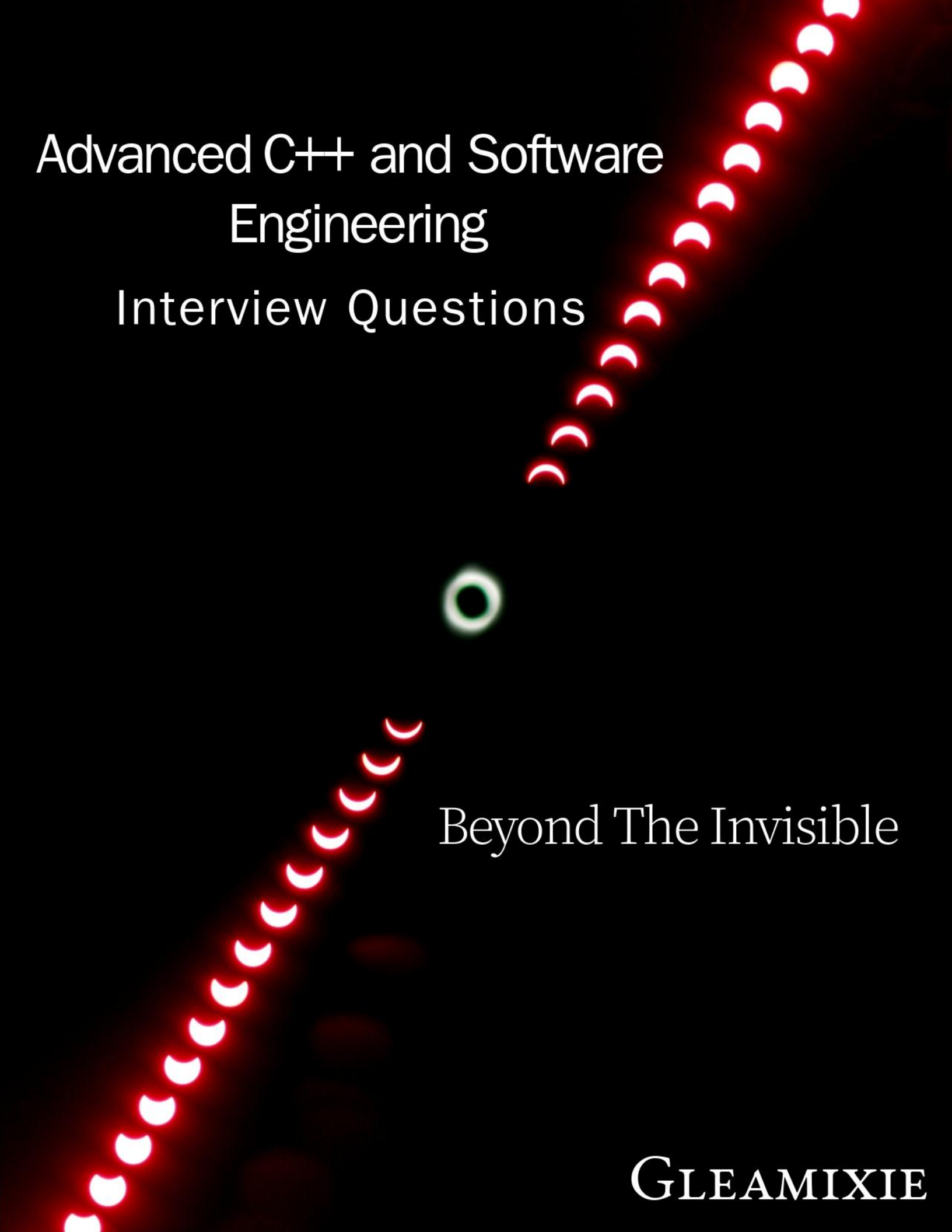


A decorative pattern of red crescent moons, some solid and some glowing, arranged in two curved lines that sweep across the cover from the top right and bottom left towards the center.

# Advanced C++ and Software Engineering Interview Questions

Beyond The Invisible

GLEAMIXIE

# Table of Contents

---

## **Preface**

## **Introduction**

## **General**

- What will print this code ?
- What is a volatile keyword in C++ ?
- What is the auto keyword in C++ ?
- What will happen if the function returns a reference and assigns it to auto ?
- What will be the output of this code ?
- What is the difference between a pointer and a reference in C++ ?
- What will be the output of this code ?
- What is the difference Between struct and class in C++ ?
- What are Pass by Value, Pass by Pointer, and Pass by Reference differences ?
- How to not allow creating objects in Stack ?
- What is a mutable keyword in C++ ?
- How to not allow creating objects in Heap ?
- What is an Explicit keyword in C++ ?
- What is the difference between Enum and Enum Class ?
- How are enum values assigned by default ?
- What is the difference between Shallow Copy and Deep Copy ?
- What is the Extern keyword in C++ ?
- What is the static keyword ?
- How can you ensure that a static variable is shared across multiple translation units ?
- Can you explain the usage of static member variables in a class, and how does initialization work for them ?
- What is the goto operator in C++ ?
- What are lambdas in C++ ?
- What is the capture list in lambdas ?
- Can you have nested lambdas in C++ ?
- What is functor in C++ ?

## **Type Casting**

- What are the different types of type casts in C++ ?

How does a C-style cast differ from other C++ casts ?  
How does static\_cast differ from C-style casting ?  
Is downcasting, upcasting safe with static\_cast ?  
Can we cast two unrelated class objects using static\_cast ?  
When should you use dynamic\_cast ?  
What happens if dynamic\_cast fails ?  
Is there a requirement for using dynamic\_cast ?  
Can dynamic\_cast be used for non-polymorphic types ?  
What is the performance difference between static\_cast and dynamic\_cast ?  
What will be the output of this code ?  
Are there any compile time optimizations related to dynamic\_cast ?  
Where have you used const\_cast ?  
What is the purpose of reinterpret\_cast in C++ ?  
Can you use reinterpret\_cast to cast between pointers of unrelated classes ?  
Is it possible to cast Struct to Integer ?  
Can you use reinterpret\_cast to cast between two unrelated classes  
and access members of the destination class ?

### **Classes In C++ And OOP**

What is the first argument of class constructors ?  
What are virtual tables ?  
What is the size of an empty class in C++ ?  
What is Polymorphism, and how is it achieved in C++ ?  
What is a dynamic dispatch ?  
What is a virtual function ?  
What are the advantages of OOP ?  
How many types of inheritance are there in C++ ?  
What is a diamond problem in C++ ?  
What will be the output of this code ?  
What is virtual destructor used for ?  
Why are all destructors not virtual by default ?  
What are the usages of the friend function and best practices ?  
What type of inheritance are there in C++ ?  
What will happen if the virtual function is called from the constructor?  
What are interface classes in C++ ?  
What is an abstract class in C++ ?  
What are pure virtual functions in C++ ?

### **Move Semantics**

What is a Move Constructor ?

What is the rvalue reference ?

What is perfect forwarding ?

How does move semantics differ from copy semantics ?

What happens if you try to use an object after it has been moved from ?

Can a move constructor throw an exception ?

Can you use move semantics with primitive data types like int ?

How std::move work under the hood ?

What is the X value ?

How does move semantics interact with smart pointers, such a std::unique\_ptr and std::shared\_ptr ?

Discuss the use of move semantics in the context of custom serialization and deserialization processes.

What are smart pointers in C++ ?

How does std::unique\_ptr ensure exclusive ownership ?

What is the purpose of std::enable\_shared\_from\_this ?

Can you create a std::shared\_ptr from a std::unique\_ptr ?

What is the difference between std::make\_shared and using new to create an object ?

Explain the role of the control block in std::shared\_ptr.

How does it keep track of references ?

What is the purpose of std::owner\_less and how is it used with smart pointers ?

How does std::default\_delete work as a default deleter for std::unique\_ptr ?

What is a circular reference, and why can it be a potential issue in C++ ?

What is data structure padding in C++ ?

Why is data structure padding needed ?

What is structure packing in C++ ?

What is structure alignment in C++ ?

What is the default alignment of a struct ?

What will be the size of the struct ?

What is the size of struct Dummy ?

What will be the output of this code ?

## **Concurrency, Multithreading and Processes**

What is a thread ?

What is a process ?

What is the difference between Thread and Process ?

What is the difference between a process and a program ?

What is multithreading ?

What is a race condition in multithreading, and how can you prevent it ?

What synchronization primitives are there in C++ ?

Can you explain the concept of thread pool ?

What is a condition variable, and how is it used in C++ multithreading ?

What is the difference between Mutex and Semaphore ?

What is the difference between user space threads and kernel space threads ?

What is the difference between Mutexes and Atomics in C++ ?

What are the advantages of Atomic over Mutexes ?

What is the difference between lock-free and wait-free ?

What is the ABA problem in concurrent programming, and how can it be mitigated ?

What is Context Switching ?

What are thread local storage (TLS) variables and when would you use them ?

Explain the Happens-Before relationship in the context of memory consistency.

What are starvation and deadlock in concurrent programming, and how can they be prevented or mitigated ?

Explain the concept of memory barriers and their importance in concurrent programming.

How many memory orders are available in C++ atomic operations ?

What are memory orders in C++ atomic operations important ?

Describe the characteristics of `memory_order_relaxed` in C++ atomic operations.

Explain the purpose of `memory_order_acquire` in C++ atomic operations. When should you use it ?

How does `memory_order_consume` differ from other memory orders in C++ atomic operations ?

How do memory orders impact the performance of multithreaded programs in C++ ?

What is default memory ordering in C++ Atomics ?

What is the purpose of `memory_order_seq_cst` in C++ atomic operations ?

Explain the concept of memory coherence in the context of multithreaded programming and how memory barriers relate to it.

How to achieve proper synchronization using acquire and release memory orders together ?

What is the use of `std::atomic_thread_fence` ?

What is a spinlock ?

How would you implement spinlock in C++ ?

Why is Spinlock implemented using `atomic_flag`, not `atomic bool` ?

What is thread affinity ?

What are magic static variables ?

What is process priority, and how does it affect process scheduling ?

What is a zombie process ?

How does an operating system manage processes ?

How many process schedulers do you know in Linux ?

What is IPC ?

How many types of IPC are there ?

What is the fastest IPC ?

What is RCU in concurrent programming ?

### **Template Metaprogramming**

What is the difference between function template specialization and function overloading ?

How can you implement type traits using C++ templates, and why are they essential for template metaprogramming ?

Discuss some practical applications of template metaprogramming in real-world C++ codebases.

What is the difference between `typename` and `class` in template declarations ?

How would you implement fibonacci using templates ?

Write a metafunction to calculate the factorial of a given number at compile time.

What is SFINAE ?

### **STL (Standard Template Library)**

What are the advantages of `std::vector` ?

What is the time complexity of common operations in `std::vector` ?

What is `std::map` and how is it implemented under the hood ?

What does amortized complexity mean in the context of `std::vector` operations ?

How is the underlying data structure of `std::priority_queue` implemented ?

What is `std::unordered_map` ?

What is the average time complexity for insertion, retrieval, and deletion operations in `std::unordered_map` ?

What is the role of the hash function in `std::unordered_map` ?

What is the purpose of the load factor in `std::unordered_map` ?

Is it possible to use complex types, such as user defined classes, as keys in `std::unordered_map` ?

What is the behavior of `std::unordered_map` when the hash function

or does the key equality function change after elements have been inserted ?  
Why is it not possible to keep references in `std::vector` ?  
Is there a workaround to keep references in `std::vector` ?  
Is it possible and good practice to extend from STL containers ?  
What is `std::allocator` ?  
When should you use a custom allocator for `std::vector` ?  
What is the difference between `set` and `multiset` ?  
How many types of Iterators are there in STL ?  
What is a contiguous iterator in STL ?  
What is the difference between `random_access_iterator` and `contiguous_iterator` ?  
What is the difference between a mutable iterator and an immutable iterator ?  
What is the difference between iterator and pointer ?  
Are there performance differences when using iterator or reverse iterator ?  
What is `std::sort` default algorithm using under the hood ?  
What are some categories of algorithms provided by STL algorithm library ?  
How does `std::reverse` differ from `std::reverse_copy` ?  
What is collision resolution in hash tables ?  
What is separate chaining in collision resolution ?  
When is open addressing preferred over separate chaining ?

### **Exceptions**

How many levels of safety are there in C++ ?  
What is the `std::nothrow` object and how is it used in exception handling ?  
What is the role of `std::unexpected` and `std::terminate` in exception handling ?  
Can you catch an exception by value and by reference at the same time ?  
Can you explain what happens if an exception is thrown within a destructor ?  
What is the role of `std::exception_ptr` in exception handling ?  
What is the difference between a rethrown exception and a propagated exception ?  
What will be the output of this code ?  
What will print this code ?

### **Design Patterns and Architectural Concepts**

What is the rule of zero in C++ ?  
What is the rule of three in C++ ?  
What is the rule of five in C++ ?  
What is Copy and Swap Idiom in C++ ?  
What is CRTP ?  
What is a virtual friend function ?  
What is policy based design ?

What is RAII (Resource Acquisition Is Initialization), and how is it useful in C++ ?

### **Performance Optimization**

What is the difference between compile time and runtime optimization ?

What is Small String Optimization (SSO) in C++ and how does it optimize string handling ?

What is copy-on-write (COW) optimization in C++ strings, and how does it work ?

What is Empty Base Class Optimization (EBCO) ?

Explain Copy Elision in C++ and how it optimizes code.

How can you optimize memory usage in C++ for containers like `std::vector` ?

What is the const-correctness principle, and how can it aid code optimization ?

What is loop unrolling, and how can it optimize code ?

What is Loop Fusion, and how can it optimize code ?

Explain the concept of branch prediction and how it can affect code optimization.

What is SIMD ?

What is loop tiling (loop blocking) and how does it enhance cache efficiency ?

Explain the concept of data prefetching in optimization. How does it improve memory access patterns ?

What is `__restrict` in C++ ?

What is the strict aliasing rule in C++ ?

What is narrowing type punning, and when is it typically used ?

What is the use of `std::memcpy` in C++ ?

### **Modern C++ Standards C++17, C++20**

What is `std::optional` and what is its usage ?

What is a fold expression in C++17 ?

What is structured binding in C++17 ?

What is `std::variant` in C++17 and how is it used ?

What is the C++17 `std::string_view` and how does it differ from `std::string` ?

What is Class Template Argument Deduction introduced in C++17 ?

What is `std::clamp` ?

What is alternative operator representation in C++ ?

What is the execution policy in C++17 ?

### **Computer Architecture and Advanced Memory Management**

What is memory fragmentation ?

What is virtual address space ?

What is the difference between physical and virtual addresses ?

What is ASLR ?

What is cache coherency ?

How many cache coherency protocols are there, list some.

What is the MOESI protocol, and how does it differ from MESI ?

What is the purpose of the Modified state in MESI and MOESI protocols ?

How does the Exclusive state differ from the Modified state in MESI and MOESI ?

Why is cache coherency important in multiprocessor systems ?

When does a cache line transition from Shared to Invalid in the MESI protocol ?

What is a cache line or cache block ?

Explain the concept of cache hit and cache miss.

What is cache associativity, and how does it affect cache performance ?

What is a write-through cache and a write-back cache ?

How does cache prefetching improve cache performance ?

What is cache eviction, and how is it managed in a cache system ?

What is cache line size, and how does it affect cache performance ?

What are Direct-mapping, Associative Mapping & Set-Associative Mappings?

What is False Sharing ?

What is cache contention, and how does it relate to false sharing ?

Explain cache line alignment and its importance in optimizing cache usage.

Explain cache coloring and how it can be used to mitigate cache coherency issues.

What is cache thrashing ?

What is pipelining in computer architecture ?

Explain the concept of out-of-order execution.

What is the purpose of a hardware interrupt in a CPU, and how does it differ from a software interrupt ?

Explain the advantages and disadvantages of symmetric multiprocessing (SMP) and asymmetric multiprocessing (AMP) in multi-core processors.

What is the role of Instruction-Level Parallelism (ILP) in processor design and how can ILP be leveraged for mission critical applications ?

What is NUMA architecture ?

Can you explain the concept of memory affinity in NUMA systems ?

## **Operating System**

What are page tables ?

Explain how page tables work.

What is TLB ?

What is the difference between a one-level and a two-level Page Table, and what are the advantages of using a two-level Page Table ?

Discuss the concept of TLB miss and the steps taken by the CPU when a TLB miss occurs.

What are Huge Pages in Linux, and when might they be advantageous ?

Can you explain the difference between the TLB and the Page Cache in Linux ?

Explain TLB shutdown and its significance in multi-processor or multi-core systems.

Discuss the benefits and challenges of using Transparent Huge Pages (THP) in Linux memory management.

What is TLB preloading, and how does it help improve memory access performance in certain scenarios ?

How does Linux implement demand paging, and what role do Page Tables play in this memory management strategy ?

### **Compiler and Linker Internals**

Explain the compilation process in C++ and the roles of the preprocessor, compiler, assembler, and linker.

What is the purpose of symbol resolution in the linking phase, and how does the linker accomplish it ?

Explain the differences between static linking and dynamic linking.

What are the advantages and disadvantages of each ?

What is name mangling in C++, and why is it used by the compiler ?

What are the common techniques used by the compiler to optimize C++ code during the compilation process ?

Explain the concept of relocation in the context of linking, and how the linker resolves relocations.

What is profile-guided optimization (PGO) in G++ ?

What is the One Translation Unit (OTU) Rule, and how does it relate to C++ code compilation and linking process ?

# Preface

---

This book is a culmination of my 10 years of experience in the field of software engineering.

While I have chosen to remain anonymous, I assure you that the questions and answers compiled in this book are designed to provide a comprehensive understanding of C++ and software engineering, especially in the context of job interviews.

The idea for this book was born out of my own experiences in preparing for interviews and the realization that there was a need for a focused, comprehensive resource on C++ and software engineering interview questions.

This book is intended for anyone preparing for a job interview where C++ is a requirement, whether you're a new graduate or an experienced professional seeking to refresh your knowledge.

While being an interview book, it is most likely you will find a new interesting topic to research/learn further about it and enhance your knowledge.

I hope that this book serves as a valuable tool in your interview preparation and contributes to your success.

**Happy Reading and Good Luck to Your Job Interview !**

You can contact me through the email address if you have any questions or suggestions regarding the book. Also please share with me the results of your job interview, I would be happy to hear from you! ❤️

Email: [gleamixie@gmail.com](mailto:gleamixie@gmail.com)

© 2025 by [Gleamixie]. All rights reserved.

No part of this publication may be reproduced, distributed, or transmitted in any form or by any means, including photocopying, recording, or other electronic or mechanical methods, without the prior written permission of the publisher, except in the case of brief quotations embodied in critical reviews and certain other noncommercial uses permitted by copyright law.

# Introduction

---

Welcome to the exciting journey of preparing C++ and Software Engineering interview.

C++ is a fundamental computer programming language that has been a cornerstone in the software industry for many years. It stands out for its performance and control and has been widely used in different software applications, including game development and high-performance application development.

This book is designed around an extensive range of C++ and software engineering topics with questions as well as distinct detailed answers. Our goal is not merely to give answers to various questions. We strive to provide you with the knowledge and understanding of various topics. Each unique question represents an excellent opportunity to enhance your understanding of the C++ and software engineering.

Each chapter dives into a specific topic, tackling questions and answers that span from fundamental to advanced concepts.

This format lets you

- 1) Read Sequentially, starting from any chapter and progress through increasingly complex topics.
- 2) Target Specific Areas, skipping directly to chapters on topics where you need more practice, offering a customized learning journey.

Remember, the key to success in any interview is preparation. This book is a tool to assist that preparation, and I hope it brings you one step closer to acing your C++ interview.

Let's get started!

# General

---

## Question 1 | General

What will print this code ?

```
#include <iostream>

class Dummy {
public:
    void test(int i) {
        std::cout << "int: " << i << std::endl;
    }
    void test(const int i) {
        std::cout << "const int: " << i << std::endl;
    }
};

int main() {
    Dummy d;
    int i = 10;
    d.test(i);
    return 0;
}
```

### Answer

Compile error.

Compiler resolves that there is no difference between functions.

To overload, either the `int` parameter must be reference or pointer.

## Question 2 | General

### What is a volatile keyword in C++ ?

#### Answer

The `volatile` keyword is used as a type qualifier for a variable, which means that the value of the variable may change at any time, even if that change has nothing to do with the current program's code .

It is used to prevent the compiler from optimizing away reads and writes to the variable, ensuring that they always occur as expected.

#### Example

```
#include <iostream>
#include <thread>

volatile bool stopFlag = false;

void modifyVolatileVariable() {
    std::cout << "Thread: Modifying volatile variable...\n";
    std::this_thread::sleep_for(std::chrono::seconds(2));
    stopFlag = true;
    std::cout << "Thread: Modification complete.\n";
}

int main() {
    std::thread t(modifyVolatileVariable);
    std::cout << "Main: Waiting volatile variable to change..."
              << std::endl;
    // Wait until the volatile variable is changed
    while (!stopFlag) {
        // Do some work in the main thread...
    }
    std::cout << "Main: Volatile variable has been changed.\n";
    t.join();

    return 0;
}
```

In this example, the `stopFlag` variable is marked as `volatile` because it is being modified by a separate thread. Without the `volatile` keyword, the compiler might optimize the loop in the main function (e.g. when optimization `O3` flag is enabled) by assuming that `stopFlag` doesn't change within the loop, leading to an infinite loop. The use of `volatile` ensures that the compiler fetches the latest value of `stopFlag` from memory every time it is accessed.

### Question 3 | General

What is the `auto` keyword in C++ ?

#### *Answer*

The `auto` keyword in C++ is used to declare a variable whose type is automatically deduced from the initializer.

For example, the following code declares `temp` variable `a` of type `int` because the initializer is an `int`.

```
auto temp = 10;
```

It works at compile time.

### Question 4 | General

What will happen if the function returns a reference and assigns it to `auto` ?

#### *Answer*

Choose `auto` `x` when you want to work with copies.

Choose `auto` `&x` when you want to work with original items and may modify them.

Choose `auto` `const` `&x` when you want to work with original items and will not modify them.

## Question 5 | General

What will be the output of this code ?

```
#include <iostream>

int main() {
    int x = 5;
    std::cout << ++x + x++ << std::endl;
    return 0;
}
```

### Answer

#### *Undefined Behavior*

Modifying a variable multiple times within a single expression, without a sequence point in between, leads to undefined behavior.

In the expression `++x + x++`, we are modifying `x` both with the pre-increment (`++x`) and the post-increment (`x++`) operators within the same expression. According to the C++ language rules, there's no guaranteed order of evaluation for these increments, and therefore, the behavior becomes undefined.

## Question 6 | General

**What is the difference between a pointer and a reference in C++ ?**

### ***Answer***

**1. Nullability**

Pointers can be assigned the value `nullptr` to indicate they are not pointing to anything. References must always refer to an existing variable and cannot be `null`.

**2. Reassignment**

Pointers can be reassigned to point to different memory addresses. References once initialized, cannot be re-assigned to refer to a different variable.

**3. Arithmetic Operations**

Pointers can be subjected to arithmetic operations (addition, subtraction) or move through memory addresses. References do not have arithmetic properties.

**4. Accessing Value**

To access the value pointed to by a pointer, you use the dereference operator (\*). For references, you directly use the reference variable's name.

## Question 7 | General

What will be the output of this code ?

```
#include <iostream>

int main() {
    int arr [] = {10, 20, 30};
    int* ptr = arr;
    std::cout << *(ptr + 1) << " " << *(ptr + 5) << std::endl;

    return 0;
}
```

### Answer

The output of this code is "20" followed by **unpredictable behavior**.

In the line `*(ptr + 1)`, we're correctly accessing the second element of the array, which is 20. However, in `*(ptr + 5)`, we're accessing memory beyond the bounds of the array. This is undefined behavior and can lead to unpredictable results.

The code might crash, output garbage values, or seem to work fine, but it's not reliable or safe.

## Question 8 | General

**What is the difference Between struct and class in C++ ?**

### ***Answer***

Both **struct** and **class** are used to define user defined data types that can hold variables and functions.

However, there is a key difference in their default access control.

**struct** - In a struct, all members are public by default. This means that all members of the struct are accessible from outside the struct.

**class** - In a class, all members are private by default. This means that members of the class are not directly accessible from outside the class.

**struct** - Inherits publicly by default. If you derive a struct from a base struct, the members of the base struct remain public in the derived struct.

**class** - Inherits privately by default. If you derive a class from a base class, the members of the base class remain private in the derived class.

Historically, struct has been used to represent simple data structures that group variables together. Members are typically treated as data only.

Class is used to define more complex objects that combine data and functions. It is suitable for implementing object oriented programming concepts.

## Question 9 | General

**What are Pass by Value, Pass by Pointer, and Pass by Reference differences ?**

### Answer

Passing by value involves sending a copy of the actual argument's value to the function. The function works with this copy, so any modifications made to the parameter within the function do not affect the original argument. This method is suitable when it is required to keep the original value unchanged in the caller's context.

```
void modifyValue (int x) {  
    x = 20; // This change doesn't affect the caller's value  
}
```

Passing by pointer involves passing a memory address to the function.

```
void modifyValue (int* ptr) {  
    *ptr = 20; // This change affects the caller's value  
}
```

This allows the function to directly modify the value at that memory location, impacting the original value in the caller's context.

Passing by reference allows you to pass the actual argument by reference, meaning any modifications made to the parameter within the function directly affect the original value in the caller's context. This provides a convenient and efficient way to work with variables meanwhile maintaining the original values.

```
void modifyValue (int& ref) {  
    ref = 20; // This change affects the caller's value  
}
```

**Pass by value** creates a copy of the argument.

**Pass by pointer** allows modification of the original value using a pointer.

**Pass by reference** directly works with the original value and is often preferred for efficiency and convenience.

## Question 10 | General

How to not allow creating objects in Stack ?

### ***Answer***

Make class destructor **Protected** or **Private**.

By making the destructor protected or private, it is possible to prevent objects from being destroyed. This discourages creating objects on the stack because they cannot be properly cleaned up.

## Question 11 | General

What is a mutable keyword in C++ ?

### ***Answer***

The **mutable** keyword is used to modify the behavior of class members within a const member function. When a class member is marked as **mutable**, it can be modified even if the containing object is const.

## Question 12 | General

How to not allow creating objects in Heap ?

### *Answer*

By deleting the operator `new`.

### *Example*

```
class Dummy {  
    public:  
        Dummy() {}  
        void* operator new(size_t) = delete;  
        ~Dummy() {}  
};
```

## Question 13 | General

### What is an Explicit keyword in C++ ?

#### Answer

The `explicit` keyword is used to indicate that a particular constructor or conversion operator should not participate in automatic type conversions or implicit conversions. It prevents the compiler from automatically invoking that constructor or conversion operator during certain operations, helping to prevent unintended and potentially dangerous implicit type conversions.

#### Example

```
#include <iostream>
#include <cstring>

class MyString {
public:
    // Constructor marked as explicit
    explicit MyString(const char* str) {
        data = new char[std::strlen (str) + 1];
        strcpy(data, str);
    }

    void print() {
        std::cout << data << std::endl;
    }

private:
    char* data;
};

int main() {

    MyString myStr1 = "Hello World"; // Error
    MyString myStr2("Hello C++");    // OK
    myStr2.print();

    return 0;
}
```

## Question 14 | General

**What is the difference between Enum and Enum Class ?**

### ***Answer***

**Enum classes** (also called scoped enumerations) makes enumerations both strongly typed and strongly scoped.

Class enum doesn't allow implicit conversion to int, and also doesn't compare enumerators from different enumerations.

**Plain enums** - where enumerator names are in the same scope as the enum and their values implicitly convert to integers and other types

**Enum classes** - where enumerator names are local to the enum and their values do not implicitly convert to other types and access is done via scope resolution operator.

## Question 15 | General

How are enum values assigned by default ?

### Answer

Enum values are assigned **integer** values by default, starting from **0** for the first member and incrementing by **1** for subsequent members.

### Example

```
#include <iostream>

enum Color {
    Red,        // Assigned the value 0
    Green,       // Assigned the value 1
    Blue        // Assigned the value 2
};

int main() {
    Color selectedColor = Green;
    if (selectedColor == 1) {
        std::cout << "The selected color is Green."
                    << std::endl;
    } else {
        std::cout << "The selected color is not Green."
                    << std::endl;
    }
    return 0;
}
```

## Question 16 | General

### What is the difference between Shallow Copy and Deep Copy ?

#### **Answer**

The default copy constructor performs a shallow copy, where member variables of the new object are copied directly from the original object.

If the class contains pointers or dynamically allocated memory, a shallow copy might lead to issues where multiple objects share the same memory, resulting in unintended behavior.

#### *Example*

```
class ShallowCopy {
private:
    char* data;

public:
    ShallowCopy(const char* input) {
        data = new char[strlen(input) + 1];
        strcpy(data, input);
    }

    // Shallow copy constructor
    ShallowCopy(const ShallowCopy& other) : data(other.data) {}

    const char* getData() const {
        return data;
    }

    ~ShallowCopy() {
        delete[] data;
    }
};

class DeepCopy {
private:
    char* data;
```

```
public:
    DeepCopy(const char* input) {
        data = new char[strlen(input) + 1];
        strcpy(data, input);
    }

    // Deep copy constructor
    DeepCopy(const DeepCopy& other) {
        data = new char[strlen(other.data) + 1];
        strcpy(data, other.data);
    }

    const char* getData() const {
        return data;
    }

    ~DeepCopy() {
        delete[] data;
    }
};
```

## Question 17 | General

**What is the Extern keyword in C++ ?**

### ***Answer***

The `extern` keyword is used to declare a variable, function, or object that is defined in another translation unit (in another source file).

It informs the compiler that the declaration is not defining the entity but is just providing information about its existence and type. This allows you to use the same variable or function in multiple files without causing linker errors due to duplicate definitions.

#### **1. Declaration without Definition**

When you declare a variable or function using `extern`, you are providing a forward declaration without actually defining the entity.

The definition should be present in another translation unit.

#### **2. Global Scope**

`extern` is often used at the global scope to indicate that the variable or function is defined elsewhere in the program.

#### **3. Multiple Files**

It is particularly useful when we need to share variables or functions across multiple source files, enabling to use them in a modular and organized manner.

#### **4. Avoid Duplicate Definitions**

Without `extern`, if the variable is defined in one source file and included its definition in another source file, it would result in duplicate definition errors during the linking phase.

## Example

In the *main.cpp* file.

```
#include <iostream>

// Forward declaration (shared.cpp)
extern int sharedValue;

int main () {
    std::cout << "Shared Value: " << sharedValue << std::endl;
    return 0;
}
```

In the *shared.cpp* file.

```
// Definition of sharedValue
int sharedValue = 42;
```

## Question 18 | General

What is the static keyword ?

### **Answer**

The `static` keyword is used with variables, functions, and classes to modify their behavior and scope. Depending on where it is used, static can have different meanings.

#### **1. Static Variables**

When used with a variable inside a function or block, the static keyword changes the variable's storage duration and scope. A static variable retains its value between function calls and persists throughout the program's execution.

#### **2. Static Functions**

When used with a function, the static keyword restricts the function's linkage to the translation unit where it is defined. It essentially makes the function have internal linkage, meaning it can't be accessed from other translation units.

#### **3. Static Class Members**

When used with class members, the static keyword indicates that the member is shared among all instances of the class rather than being specific to each instance.

#### **4. Static Local Variables**

When used with a variable inside a function or block, the static keyword creates a variable with static storage duration that retains its value between function calls.

## Question 19 | General

**How can you ensure that a static variable is shared across multiple translation units ?**

### ***Answer***

If we want a single instance of a static variable to be shared across multiple translation units, we should define it in a header file and include that header in all relevant source files.

This ensures consistent memory allocation for the variable.

## Question 20 | General

**Can you explain the usage of static member variables in a class, and how does initialization work for them ?**

### ***Answer***

Static member variables in a class are shared among all instances of the class.

They are declared using the static keyword within the class definition. When a static member variable is initialized within the class, it must be a constant expression. If not initialized within the class, the variable needs to be defined and initialized outside the class definition to allocate memory for it.

```
class MyClass {  
public:  
    // Initialized within the class  
    static const int constantStaticVar = 42;  
  
    // Not initialized within the class  
    static int nonConstantStaticVar;  
  
    // Compilation Error  
    static int nonConstantStaticVar = 22;  
};
```

```
// Defined and initialized outside the class  
int MyClass::nonConstantStaticVar = 43;
```

## Question 21 | General

What is the goto operator in C++ ?

### Answer

The `goto` statement is used to transfer the control of a program to a labeled statement within the same function. It provides an unconditional jump from one part of the code to another.

The use of the goto statement is generally discouraged and considered bad practice, as it can make the code harder to understand and maintain.

```
#include <iostream>  
  
int main() {  
    int value = 22;  
  
    if(value < 42) {  
        // Jump to the label if the condition is true  
        goto lessThanTen;  
    }  
  
    std::cout << "Value is greater than or equal to 10."  
              << std::endl;  
    return 0;  
  
lessThanTen:  
    std::cout << "Value is less than 10." << std::endl;  
    return 0;  
}
```

## Question 22 | General

What are lambdas in C++ ?

### ***Answer***

A lambda is an anonymous function that you can define inline within your code.

Lambdas provide a concise way to define small, self-contained functions that can be used directly at the point where they are defined. Lambdas are often used for short, one-off functions that are passed as arguments to algorithms, functions, or other constructs.

**Syntax** `[](){}`

## Question 23 | General

What is the capture list in lambdas ?

### Answer

Capture list is `[]` in lambda expression `[](){};`

### Example

```
int x = 43;
int y = 22;

// No captured variables. x and y are not seen inside the lambda function.
auto f0 = [](){};

// Capture all outside variables by reference.
auto f1 = [&](){};

// Capture all outside variables by copy, and y with reference.
auto f2 = [=, &y](){};

// Capture all outside variables by reference, y with copy.
auto f3 = [&, y](){};

// Capture all outside variables by copy.
auto f4 = [=](){};
```

## Question 24 | General

Can you have nested lambdas in C++ ?

### Answer

Yes, it is possible to define nested lambdas, where an inner lambda can capture variables from an outer lambda's scope.

## Question 25 | General

### What is functor in C++ ?

#### **Answer**

A functor refers to an object or class that overloads the function call `operator()` to allow it to be used like a function. Functors are instances of classes that act as function-like types, providing a way to encapsulate behavior and state. They are often used for various purposes, such as custom sorting, as arguments to algorithms, or in scenarios where it is required to provide a callable object with specific behavior.

#### *Example*

```
#include <iostream>

class Adder {
public:
    Adder(int value) : v(value) {}

    int operator() (int x) const {
        return x + v;
    }

private:
    int v;
};

int main() {
    Adder adder(5);
    int n = 6;
    int res = adder(n);
    std::cout << "Result: " << res << std::endl;
    return 0;
}
```

# Type Casting

---

## Question 26 | Type Casting

What are the different types of type casts in C++ ?

**Answer**

1. C-style Cast - `(type)` value
2. Static Cast - `static_cast<type>` (expression)
3. Dynamic Cast - `dynamic_cast<type>` (expression)
4. Const Cast - `const_cast<type>` (expression)
5. Reinterpret Cast - `reinterpret_cast<type>` (expression)

## Question 27 | Type Casting

How does a C-style cast differ from other C++ casts ?

**Answer**

The C-style cast attempts a sequence of different casts to achieve the desired conversion, which can sometimes result in unintended conversions.

Other C++ casts (e.g. `static_cast`, `dynamic_cast`) provide more precise control over the conversion and offer better type safety.

## Question 28 | Type Casting

How does `static_cast` differ from C-style casting ?

### Answer

Unlike C-style casting, `static_cast` provides more type safety. It performs a more restricted set of conversions and disallows potentially unsafe conversions that C-style casting might allow.

## Question 29 | Type Casting

Is downcasting, upcasting safe with `static_cast` ?

### Answer

Upcasting is safe. It involves converting a pointer or reference from a derived class type to its base class type. This is safe because a derived class always contains all the members of its base class.

Downcasting can be safe or unsafe. It involves converting a pointer or reference from a base class type to its derived class type. It is safe if we are sure about the runtime type of the object being pointed to.

However, if we are not certain about the runtime type, downcasting can be unsafe and might lead to undefined behavior.

```
class B {};  
  
class D : public B {};  
  
void f(B* pb, D* pd) {  
    // Not safe, D can have fields/methods that are not in B.  
    D* pd2 = static_cast<D*>(pb);  
  
    // Safe conversion, D always contains all of B.  
    B* pb2 = static_cast<B*>(pd);  
}
```

## Question 30 | Type Casting

Can we cast two unrelated class objects using `static_cast` ?

**Answer**

Yes.

*Example*

```
#include <iostream>

class A {
public:
    int value;
    A(int v) : value(v) {}
};

class B {
public:
    double value;
    B(double v) : value(v) {}

    // Constructor to convert from A to B
    B(const A& a) : value(static_cast<double>(a.value)) {}
};

int main () {
    A a(10);
    // Implicitly using the constructor
    B b = static_cast<B>(a);
    // Output: 10.0
    std::cout << b.value << std::endl;
    return 0;
}
```

## Question 31 | Type Casting

When should you use `dynamic_cast` ?

### *Answer*

You should use `dynamic_cast` when you want to perform a type conversion involving polymorphic classes and you need to ensure that the conversion is valid at runtime.

## Question 32 | Type Casting

What happens if `dynamic_cast` fails ?

### *Answer*

If `dynamic_cast` fails to cast a pointer to the desired type, it returns a `NULL` for pointers or throws a `std::bad_cast` exception for references.

## Question 33 | Type Casting

Is there a requirement for using `dynamic_cast` ?

### *Answer*

To use `dynamic_cast`, the base class in the hierarchy must have at least one virtual function. This enables the creation of the vtable necessary for runtime type identification.

## Question 34 | Type Casting

Can `dynamic_cast` be used for non-polymorphic types ?

### **Answer**

No, `dynamic_cast` can only be used for types involved in a polymorphic hierarchy with at least one virtual function.

## Question 35 | Type Casting

What is the performance difference between `static_cast` and `dynamic_cast` ?

### **Answer**

`static_cast` is faster than `dynamic_cast` because it doesn't involve runtime type checks. `dynamic_cast` requires runtime type information (RTTI) and performs a runtime check for validity, which incurs a performance overhead.

## Question 36 | Type Casting

What will be the output of this code ?

```
#include <iostream>

class BaseA {
public:
    virtual void print() {
        std::cout << "BaseA" << std::endl;
    }
};

class BaseB {}

int main ()
{
    BaseA* baseAPtr = new BaseA;
    BaseB* baseBPtr = dynamic_cast<BaseB*> (baseAPtr);

    if (baseBPtr != 0) {
        std::cout << "Succeed to cast" << std::endl;
    } else {
        std::cout << "Failed to cast" << std::endl;
    }

    delete baseAPtr;
    return 0;
}
```

*Answer*

It will print **Failed to cast**.

## Question 37 | Type Casting

Are there any compile time optimizations related to `dynamic_cast` ?

### **Answer**

Yes, some compilers can optimize away certain `dynamic_cast` operations when it is known that the cast will always succeed due to the code logic.

However, these optimizations might not cover all cases.

## Question 38 | Type Casting

Where have you used `const_cast` ?

### **Answer**

When working with third party libraries that have incorrectly declared `const` on their functions or methods, you might need to use `const_cast` to work around these issues. This is particularly the case when the library's interface cannot be modified.

## Question 39 | Type Casting

What is the purpose of `reinterpret_cast` in C++ ?

### **Answer**

`reinterpret_cast` is used to convert one pointer type to another, even if the types are unrelated. It allows for low level type conversions and should be used with extreme caution due to its potential for introducing undefined behavior.

## Question 40 | Type Casting

Can you use `reinterpret_cast` to cast between pointers of unrelated classes ?

### Answer

Yes, `reinterpret_cast` can be used to cast between pointers of unrelated classes. Despite that, this can lead to undefined behavior unless the cast is explicitly defined by your program's memory layout and platform characteristics.

## Question 41 | Type Casting

Is it possible to cast Struct to Integer ?

### Answer

Yes.

### Example

```
#include <iostream>

struct Dummy {
    int x;
};

int main () {
    Dummy d;
    d.x = 10;
    int x = reinterpret_cast<int&> (d);
    std::cout << x << std::endl;

    return 0;
}
```

## Question 42 | Type Casting

**Can you use `reinterpret_cast` to cast between two unrelated classes and access members of the destination class ?**

### ***Answer***

No, `reinterpret_cast` should not be used to cast between unrelated classes and access members of the destination class. This is highly unsafe and can lead to undefined behavior. The behavior is not well-defined since the memory layouts and member offsets of unrelated classes can vary widely.

Attempting to access members of the destination class using such a cast can result in accessing memory locations that do not correspond to the members of the destination class, causing unexpected behavior or crashes.

# Classes In C++ And OOP

---

## Question 43 | Classes In C++ And OOP

What is the first argument of class constructors ?

### Answer

The first argument of class constructors in C++ is the implicit *this* pointer, which refers to the current instance of the class. It allows the constructor to access and initialize the data members and methods of that specific instance.

It is important to note that the *this* pointer is an implicit pointer provided by the C++ language, and you can't see it with the naked eye when you write code.

It's a behind-the-scenes mechanism that the compiler uses to manage object instances and their members. It is not explicitly passed or manipulated the *this* pointer from the code. It is automatically available within class methods to reference the current object instance.

## Question 44 | Classes In C++ And OOP

What are virtual tables ?

### Answer

A vtable is an array of function pointers.

Each class that contains at least one virtual function has an associated vtable. The vtable stores pointers to the most derived versions of the virtual functions defined in the class and its base classes. When a class is instantiated, an additional hidden pointer is added to the object's memory layout. This pointer points to the vtable of the object's class.

When a virtual function is called through a pointer or reference, the compiler uses the vpointer to access the appropriate vtable and determine which function to execute. This allows the correct version of the function to be invoked based on the object's actual runtime type.

## Question 45 | Classes In C++ And OOP

What is the size of an empty class in C++ ?

### **Answer**

The size of an empty class is **one** byte.

This is because C++ mandates that every object of a class type must have a distinct memory address. Therefore, even an empty class needs some memory space to ensure that each object has a unique address.

## Question 46 | Classes In C++ And OOP

What is Polymorphism, and how is it achieved in C++ ?

### **Answer**

Polymorphism is the ability of an object to take on different forms.

The polymorphism is achieved through two mechanisms.

1. **Compile time polymorphism**

This is achieved through function overloading and operator overloading.

2. **Runtime polymorphism**

This is achieved through virtual functions and dynamic dispatch.

**Compile time** polymorphism occurs when a function or operator is defined with multiple signatures, each of which takes a different set of arguments. The compiler will choose the correct signature to call based on the arguments that are passed to the function or operator.

**Runtime** polymorphism occurs when a function call is made to a virtual function. A virtual function is a function that is declared in a base class and overridden in a derived class. When a virtual function is called, the compiler will not choose the correct version of the function to call until the program is running. The decision of which version of the function to call is made at runtime, based on the type of the object that is being referred to.

## **Question 47 | Classes In C++ And OOP**

**What is a dynamic dispatch ?**

### ***Answer***

Dynamic dispatch is the process of choosing which version of a virtual function to call at runtime. Dynamic dispatch is achieved by the runtime system, which keeps track of the types of objects that are being created and used.

## **Question 48 | Classes In C++ And OOP**

**What is a virtual function ?**

### ***Answer***

A virtual function is a function that is declared in a base class and overridden in a derived class.

When a virtual function is called, the compiler will not choose the correct version of the function to call until the program is running. The decision of which version of the function to call is made at runtime, based on the type of the object that is being referred to.

## Question 49 | Classes In C++ And OOP

What are the advantages of OOP ?

### ***Answer***

#### **1. Abstraction**

OOP allows you to abstract the data and methods of an object, so that you only need to worry about the interface of the object. This makes your code easier to understand and maintain.

#### **2. Data hiding**

OOP allows you to hide the data and methods of an object, so that they cannot be accessed by other objects. This makes your code more secure.

#### **3. Code reuse**

OOP allows you to reuse code by creating classes and objects. This can save you time and effort when writing code.

#### **4. Modularization**

OOP allows you to break down your code into smaller, more manageable modules. This makes your code easier to understand and maintain.

## Question 50 | Classes In C++ And OOP

How many types of inheritance are there in C++ ?

### **Answer**

C++ supports **five** types of inheritance.

**1. Single inheritance**

A class can inherit from only one base class.

**2. Multilevel inheritance**

A class can inherit from another class that already inherits from another class.

**3. Hierarchical inheritance**

A class can inherit from multiple base classes, but each base class is only inherited from once.

**4. Hybrid inheritance**

A combination of hierarchical and multiple inheritance.

**5. Virtual inheritance**

A class can inherit from another class, but the derived class can only have one copy of the base class's data members and methods.

1. Single inheritance is the most common type of inheritance in C++. It is a straightforward way to reuse code and create new classes.

2. Multilevel inheritance is a more complex type of inheritance. It allows to create classes that inherit from other classes that inherit from other classes.

3. Hierarchical inheritance is a type of inheritance where a class can inherit from multiple base classes, but each base class is only inherited from once.

4. Hybrid inheritance is a combination of hierarchical and multiple inheritance. This can be used to create more complex class hierarchies.

5. Virtual inheritance is a type of inheritance where a class can inherit from another class, but the derived class can only have one copy of the base class's data members and methods.

## Question 51 | Classes In C++ And OOP

What is a diamond problem in C++ ?

### Answer

The diamond problem is a problem that can occur in C++ when a class inherits from two base classes that have a common base class. The problem arises because the derived class inherits two copies of the data members and methods from the common base class. This can lead to ambiguity and errors.

The diamond problem is illustrated by the following class hierarchy.

```
class Base1 {  
public:  
    int x;  
};  
  
class Base2 {  
public:  
    int y;  
};  
  
class Derived : public Base1, public Base2 {  
};
```

The Derived class inherits the **x** data member from Base1 and the **y** data member from Base2. This means that the Derived class has two copies of the **x** and **y** data members.

If the Derived class has a function that accesses the **x** or **y** data member, the compiler will not know which copy of the data member to use. This can lead to ambiguity and errors.

There are a few ways to avoid the diamond problem. One way is to use virtual inheritance. Virtual inheritance ensures that the derived class only inherits one copy of the data members and methods from the common base class.

Another way to avoid the diamond problem is to use multiple inheritance carefully. When using multiple inheritance, it is important to make sure that the base classes do not have any overlapping data members or methods.

## Question 52 | Classes In C++ And OOP

What will be the output of this code ?

```
class Base {
public:
    virtual void f() {
        cout << "Base::f" << endl;
    }
};

class Derived : public Base {
public:
    void f() override {
        cout << "Derived::f" << endl;
    }
};

int main() {
    Base* b = new Derived();
    b->f();
    return 0;
}
```

### Answer

It will print **Derived::f**

## Question 53 | Classes In C++ And OOP

What is virtual destructor used for ?

### Answer

A **virtual destructor** is a destructor that is declared as a virtual function.

A destructor is a function that is called when an object is destroyed. The virtual destructor is used to ensure that the correct destructor is called when an object is destroyed, even if the object is a pointer to a base class.

### Example

```
class Base {
public:
    virtual ~Base() {}
};

class Derived : public Base {
public:
    virtual ~Derived() {}
};

int main () {
    Base* b = new Derived();
    delete b;
    return 0;
}
```

The Base class has a virtual destructor. The Derived class also has a virtual destructor. When the b variable is deleted, the runtime system will call the ~Derived() destructor then ~Base() destructor, even though the b variable is a pointer to a Derived object. If there is no virtual keyword in ~Base(), the call to ~Derived() destructor will not happen.

The virtual destructor is important because it ensures that the correct destructor is called, even if the object is a pointer to a base class.

This is important because it prevents memory leaks and other errors.

## Question 54 | Classes In C++ And OOP

**Why are all destructors not virtual by default ?**

### ***Answer***

Virtual destructors can lead to performance overhead, as the runtime system has to search the inheritance hierarchy to find the correct destructor to call.

Non virtual destructors provide predictable behavior, especially in cases where inheritance and polymorphism are not involved. The destructor of the actual class type is always called, eliminating ambiguity.

## Question 55 | Classes In C++ And OOP

**What are the usages of the friend function and best practices ?**

### ***Answer***

To implement operators that are overloaded for a class.

To implement helper functions that are used by the members of a class.

To allow a function to access the private members of a class for debugging or testing purposes.

#### **1. Avoid Overexposing Data**

Friend functions allow controlled access to private members without exposing them to the public interface. This helps maintain encapsulation and data hiding.

#### **2. Operator Overloading**

Friend functions are often used to overload operators for classes where one or both operands are of a user-defined type. For example, overloading << for custom output or + for custom addition. By using a friend function for operator overloading, you can provide access to the necessary private members without exposing them through public member functions.

Sometimes it is not possible to declare operator overload as a member function, like in IO operator << and operator >>. The first parameter of these functions has to be ostream or istream,

which are library classes and it is not possible to extend them, declaring such functions as friend gives them access to private variables of the class.

## **Question 56 | Classes In C++ And OOP**

**What type of inheritance are there in C++ ?**

### ***Answer***

There are 4 types of inheritance.

#### **1. Public inheritance**

In the public inheritance, the public members of the base class are inherited as public members of the derived class, and the protected members of the base class are inherited as protected members of the derived class. Public inheritance is the most common type of inheritance in C++.

#### **2. Protected inheritance**

In the protected inheritance, the public and protected members of the base class are inherited as protected members of the derived class. Protected inheritance is used when we want to prevent the derived class from accessing the public members of the base class.

#### **3. Private inheritance**

In the private inheritance, the public and protected members of the base class are inherited as private members of the derived class. Private inheritance is used when we want to prevent the derived class from accessing any of the members of the base class.

#### **4. Virtual inheritance**

In the virtual inheritance, the base class is inherited multiple times, but each derived class has its own copy of the base class's members. Virtual inheritance is used to avoid the diamond problem, which is a situation where two derived classes inherit from the same base class, and the base class has a virtual function.

## Question 57 | Classes In C++ And OOP

**What will happen if the virtual function is called from the constructor?**

### ***Answer***

Calling a virtual function in a constructor can lead to unexpected behavior and potential issues. When a virtual function is called in a constructor, it is important to understand how the C++ object's construction process works and how virtual dispatch is affected.

### **Construction Order**

When an object of a derived class is created, the base class constructor is called before the derived class constructor. This means that during the base class constructor's execution, the object is not yet fully constructed as a derived object.

### **Virtual Dispatch**

Virtual dispatch relies on the type information of the fully constructed object to determine which function to call. During the base class constructor's execution, the object's type is still that of the base class, even if it's being constructed as a derived object.

Given these points, calling a virtual function in a constructor can lead to the following issues.

#### **1. Incorrect Dispatch**

When a virtual function is called in a base class constructor, it is dispatched based on the type of the base class, not the derived class.

This means that the base class's version of the virtual function will be invoked, even if the derived class has overridden the function.

#### **2. Undefined Behavior**

If the virtual function being called accesses members or behavior that are specific to the derived class and those members are not yet constructed, it can lead to undefined behavior or unexpected results.

## Example

```
class Base {
public:
    Base() {
        callVirtualFunction();
    }

    virtual void callVirtualFunction() {
        std::cout << "Base::callVirtualFunction()"
                  << std::endl;
    }
};

class Derived : public Base {
public:
    Derived() {}

    void callVirtualFunction() override {
        std::cout << "Derived::callVirtualFunction()"
                  << std::endl;
    }
};

int main() {
    Derived d;
    return 0;
}
```

## Question 58 | Classes In C++ And OOP

What are interface classes in C++ ?

### **Answer**

This concept is achieved using abstract classes and pure virtual functions. An abstract class is a class that cannot be instantiated and serves as a blueprint for other classes. It contains at least one pure virtual function, which has no implementation in the base class and must be overridden by derived classes. And it should not contain any method that has implementation.

## Question 59 | Classes In C++ And OOP

What is an abstract class in C++ ?

### **Answer**

An abstract class in C++ is a class that cannot be instantiated on its own and is meant to serve as a base or blueprint for other classes. It provides a common interface and may include both regular member functions with implementations and pure virtual functions with no implementations.

Key points about abstract classes.

#### **1. Cannot Be Instantiated**

It is not possible to create objects of an abstract class directly. It is meant to be inherited by other classes, which will provide implementations for its pure virtual functions.

#### **2. Pure Virtual Functions**

An abstract class typically contains one or more pure virtual functions.

A pure virtual function is declared using the syntax below.

```
virtual returnType functionName (parameters) = 0;
```

These functions have no implementation in the abstract class and must be overridden by derived classes.

### 3. Concrete Member Functions

An abstract class can also have regular member functions with implementations. These functions provide common behavior that can be inherited by derived classes.

### 4. Inheritance and Polymorphism

Derived classes that inherit from an abstract class must implement the pure virtual functions. This enforces a contract and ensures that every derived class adheres to the same interface. Instances of derived classes can be treated polymorphically using pointers or references of the abstract class type.

## Question 60 | Classes In C++ And OOP

What are pure virtual functions in C++ ?

### ***Answer***

A pure virtual function is a virtual function declared in a base class with the syntax below.

```
virtual returnType functionName (parameters) = 0;
```

It is a function that has no implementation in the base class and is meant to be overridden by derived classes.

# Move Semantics

---

## Question 61 | Move Semantics

### What is a Move Constructor ?

#### Answer

A move constructor in C++ is a special member function that allows the efficient transfer of ownership of resources from one object to another. It is used to implement move semantics, which is a technique to improve performance when objects need to be transferred or returned from functions.

Key points about move constructors.

#### 1. Purpose

Move constructors are used to efficiently transfer the resources owned by one object to another, avoiding unnecessary copying and memory allocation.

#### 2. rvalue References

A move constructor takes an rvalue reference (&&) as its parameter, indicating that it accepts temporary objects or objects about to go out of scope.

#### 3. Stealing Resources

Instead of copying the resources, a move constructor **steals** the resources from the source object by transferring ownership. The source object is typically left in a valid but unspecified state.

#### 4. Implementation

A move constructor is usually implemented by swapping the internal resource pointers or data members between the source and destination objects.

#### 5. `std::move`

To invoke the move constructor explicitly, you can use the `std::move` function, which casts an object to an rvalue reference. This indicates that you are willing to move from the object, rather than copy.

## Question 62 | Move Semantics

What is the rvalue reference ?

### *Answer*

An rvalue reference is a reference used to bind to temporary values i.e., rvalues. When used correctly, this reference type allows you to determine the distinction between lvalues on the left side and rvalues on the right side of the assignment .

rvalue references are essential to move semantics, perfect forwarding, among other advanced capabilities. The type specifier identifying an rvalue reference is two ampersands `&&`.

You'll also require rvalue references for other purposes in your codebase, such as implementing move constructors, move assignment operators, generating move-only types, and functions that need to forward their arguments without essential copying.

## Question 63 | Move Semantics

### What is perfect forwarding ?

#### **Answer**

**Perfect forwarding** is a technique in C++ that allows you to forward function arguments from one function to another while preserving their original value categories.

It might be useful when you have a wrapper function that needs to pass its arguments to another function without losing information about whether the arguments are temporary (rvalues) or have names (lvalues).

Perfect forwarding is often used in generic programming scenarios and when implementing forwarding functions like constructors or factory functions.

Key points about perfect forwarding.

#### **1. Maintaining Value Categories**

Perfect forwarding ensures that the arguments are passed on to the target function with the same value categories as they were received. lvalues remain lvalues, and rvalues remain rvalues.

#### **2. Templates and rvalue References**

Perfect forwarding is achieved using templates and rvalue references (&&). Templated forwarding functions are used to deduce the argument types and use rvalue references to forward them.

#### **3. `std::forward`**

The `std::forward` function is used to forward the arguments while maintaining their value categories. It takes care of casting the arguments back to their original value categories.

## Example

```
#include <iostream>
#include <utility>

class Data {
public:
    Data(int value) : data(value) {
        std::cout << "Constructor called with value: "
                    << value << std::endl;
    }

    // Copy constructor
    Data(const Data& other) : data(other.data) {
        std::cout << "Copy Constructor called" << std::endl;
    }

    // Move constructor
    Data(Data&& other) noexcept : data(other.data) {
        std::cout << "Move Constructor called" << std::endl;
    }

private:
    int data;
};

// Forwarding function template
template <typename T>
void forwardData(T&& data) {
    processData(std::forward<T>(data));
}

// Function to process the forwarded data
void processData(Data&& data) {
    std::cout << "Processing rvalue data" << std::endl;
}

void processData(Data& data) {
    std::cout << "Processing lvalue data" << std::endl;
}
```

```
int main() {  
    Data d(42);  
    forwardData(d); // lvalue forwarding  
    forwardData(Data(22)); // rvalue forwarding  
    return 0;  
}
```

## Question 64 | Move Semantics

How does move semantics differ from copy semantics ?

### **Answer**

**Copy semantics** involve creating a duplicate of an object's resources, while move semantics involve transferring ownership of the resources from one object to another.

**Move semantics** are more efficient and can be used when the original object is no longer needed.

## Question 65 | Move Semantics

What happens if you try to use an object after it has been moved from ?

### **Answer**

If you use an object after it has been moved from, the behavior is **undefined**. It is important to ensure that an object's state is not accessed after it has been moved.

## Question 66 | Move Semantics

Can a move constructor throw an exception ?

### Answer

Yes, a move constructor can throw an exception. It is generally recommended to ensure that move constructors do not throw exceptions to maintain the **nothrow** guarantee associated with move operations.

## Question 67 | Move Semantics

Can you use move semantics with primitive data types like int ?

### Answer

While move semantics can technically be applied to primitive data types, it is not as impactful as with larger objects. Move semantics are most beneficial when dealing with objects that manage resources or have complex copy constructors.

## Question 68 | Move Semantics

How `std::move` work under the hood ?

### Answer

The first thing to note is that `std::move()` doesn't actually move anything. It changes an expression from being an **lvalue** (such as a named variable) to being an **xvalue**. An **xvalue** tells the compiler – you can plunder me, move anything I'm holding and use it elsewhere since I'm going to be destroyed.

## Example

```
template <typename T>
typename std::remove_reference<T>::type&&
move (T&& arg) noexcept {
    return static_cast<typename std::remove_reference<T>::type&&>(arg);
}
```

## Question 69 | Move Semantics

What is the X value ?

**Answer**

`Xvalue` represents an `expiring` value, the result of evaluating certain expressions like function calls that return an `rvalue` reference or expressions involving `std::move`.

## Question 70 | Move Semantics

How does move semantics interact with smart pointers, such a `std::unique_ptr` and `std::shared_ptr` ?

**Answer**

The `std::unique_ptr` can be efficiently moved between objects, transferring ownership of the dynamically allocated resource. With `std::shared_ptr`, the move operations involve updating the reference counts appropriately to maintain shared ownership.

## Question 71 | Move Semantics

Discuss the use of move semantics in the context of custom serialization and deserialization processes.

**Answer**

Move semantics can be leveraged in custom serialization and deserialization processes to efficiently transfer serialized data between objects. When loading or deserializing data, move semantics allow the efficient construction of objects, reducing unnecessary copies. Similarly, when serializing data, move semantics enable the transfer of data without duplicating it.

# Memory Management

---

## Question 72 | Memory Management

What are smart pointers in C++ ?

### *Answer*

Smart pointers automatically manage the lifetime of dynamically allocated objects. They provide a safer and more convenient alternative to raw pointers by automatically releasing memory when it is no longer needed.

## Question 73 | Memory Management

How does `std::unique_ptr` ensure exclusive ownership ?

### *Answer*

`std::unique_ptr` enforces exclusive ownership by disallowing copying. It can only be moved, which transfers ownership to the new `std::unique_ptr`. When the unique pointer goes out of scope or is explicitly deleted, the managed object is automatically deleted.

## Question 74 | Memory Management

What is the purpose of `std::enable_shared_from_this` ?

### Answer

`std::enable_shared_from_this` is a base class that allows an object managed by `std::shared_ptr` to safely create additional `std::shared_ptr` instances that share ownership of the same object. It provides a way to obtain `std::shared_ptr` instances without risking premature object deletion.

## Question 75 | Memory Management

Can you create a `std::shared_ptr` from a `std::unique_ptr` ?

### Answer

Yes, you can move ownership from a `std::unique_ptr` to a `std::shared_ptr`.

This is useful when you want to share ownership after creating an object with `std::unique_ptr`.

```
#include <memory>

int main() {

    std::unique_ptr<int> uniquePtr(new int(10));
    std::shared_ptr<int> sharedPtr = std::move(uniquePtr);

    return 0;
}
```

## Question 76 | Memory Management

What is the difference between `std::make_shared` and using `new` to create an object ?

### Answer

The key differences between `std::make_shared` and `new` are.

`std::make_shared` allocates both the object and its control block in a single memory allocation, potentially reducing memory overhead.

Using `new` to create an object and then creating a `std::shared_ptr` involves two separate memory allocations: **one for the object** and another **for the control**.

```
std::shared_ptr<MyClass> ptr1 = std::make_shared<MyClass>();  
std::shared_ptr<MyClass> ptr2 (new MyClass);
```

## Question 77 | Memory Management

Explain the role of the control block in `std::shared_ptr`.  
How does it keep track of references ?

### Answer

The control block associated with `std::shared_ptr` stores the reference count and a pointer to the managed object. When a new `std::shared_ptr` is created, the control block is shared among all instances pointing to the same object.

The reference count is incremented when a new `std::shared_ptr` is created and decremented when it goes out of scope. When the reference count reaches zero, the control block and the managed object are deleted.

## Question 78 | Memory Management

What is the purpose of `std::owner_less` and how is it used with smart pointers ?

### Answer

`std::owner_less` is a comparison function that can be used to compare the ownership relationships of smart pointers. It is particularly useful in scenarios where you want to sort or order smart pointers in a container based on their ownership relationships.

## Question 79 | Memory Management

How does `std::default_delete` work as a default deleter for `std::unique_ptr` ?

### Answer

`std::default_delete` is a template class used as the default deleter for `std::unique_ptr`. It's specialized for specific types and uses the `delete` operator to deallocate memory when the `std::unique_ptr` goes out of scope.

For example, `std::unique_ptr<int>` uses `std::default_delete<int>` as its default deleter.

## Question 80 | Memory Management

What is a circular reference, and why can it be a potential issue in C++ ?

### Answer

A circular reference, also known as a cyclic dependency, occurs when two or more objects have references to each other, forming a cycle. In C++, this can become an issue when memory is not properly managed, leading to memory leaks and unexpected behavior. Circular references are problematic because they can prevent the automatic deallocation of memory by smart pointers. When objects hold strong references to each other, their reference counts won't reach zero even when they're no longer accessible, preventing the memory from being released.

## Question 81 | Memory Management

What is data structure padding in C++ ?

### **Answer**

[Data structure padding](#) also known as data structure alignment is the concept of adding extra bytes between structure members to ensure proper alignment and improve memory access efficiency.

Padding aligns members according to their data types and memory requirements.

## Question 82 | Memory Management

Why is data structure padding needed ?

### **Answer**

[Data structure padding](#) is needed to align data in a memory-efficient way that improves access speed. CPUs perform better with aligned data.

Without padding, memory reads and writes might require multiple memory accesses, leading to performance degradation.

## Question 83 | Memory Management

What is structure packing in C++ ?

### **Answer**

[Structure packing](#) refers to the process of arranging the members of a struct in memory to minimize padding. This is often done to reduce memory consumption.

However, it can result in slower memory access due to unaligned memory access penalties on certain architectures.

## Question 84 | Memory Management

What is structure alignment in C++ ?

### Answer

**Structure alignment** refers to aligning the members of a struct according to their data types' alignment requirements. It ensures that members are accessed efficiently by CPUs. The alignment is determined by the largest data type in the struct.

## Question 85 | Memory Management

What is the default alignment of a struct ?

### Answer

The default alignment of a struct is determined by the alignment requirement of its largest member. As an example, if a struct contains an `int` and a `double`, the alignment will be determined by the `double`.

## Question 86 | Memory Management

What will be the size of the struct ?

```
struct Dummy
{
    private:
        int i;    // 4 bytes
        double d; // 8 bytes
        bool b;   // 1 byte
        double x; // 8 bytes
};
```

### Answer

32 Bytes.

## Question 87 | Memory Management

What is the size of struct Dummy ?

```
struct Dummy
{
    private:
        int x;
        double d;
        bool b;
        double do;
    public:
        virtual void f() {}
};
```

*Answer*

40 bytes. (v\_ptr\* added as so added 8 bytes).

## Question 88 | Memory Management

What will be the output of this code ?

```
#include <iostream>
#include <memory>

class Base {
public:
    Base() {
        std::cout << "Base class constructor" << std::endl;
    }
    ~Base() {
        std::cout << "Base class destructor" << std::endl;
    }
};

class Derived : public Base {
public:
    Derived() {
        std::cout << "Derived class constructor" << std::endl;
    }
    ~Derived() {
        std::cout << "Derived class destructor" << std::endl;
    }
};

int main() {
    std::shared_ptr <Base> b2d = std::make_shared<Derived>();
    b2d.reset();
}
```

### Answer

Base class constructor  
Derived class constructor  
Derived class destructor  
Base class destructor

# Concurrency, Multithreading and Processes

---

## Question 89 | Concurrency, Multithreading and Processes

What is a thread ?

### ***Answer***

A [thread](#) is the smallest unit of a CPU's execution that can be scheduled by the operating system's scheduler. It represents a single sequence of instructions and has its own program counter, stack, and set of registers.

Threads within a process share the same memory space and resources, allowing them to communicate and cooperate with each other more efficiently than separate processes.

## Question 90 | Concurrency, Multithreading and Processes

What is a process ?

### ***Answer***

A [process](#) in computing refers to an independent and self contained program that runs in its own isolated memory space. It is the fundamental unit of execution in an operating system.

Each process has its own set of resources, including memory, file handles, and system information, and runs as an independent entity.

## Question 91 | Concurrency, Multithreading and Processes

What is the difference between Thread and Process ?

### **Answer**

#### **1. Independence**

Processes are independent of each other. Each process runs in its own separate memory space and is isolated from other processes. They cannot directly access each other's memory.

#### **2. Shared Memory**

Threads within the same process share the same memory space. They can directly access and modify the process's memory, which simplifies communication between threads.

#### **3. Resource Intensive**

Processes are more resource intensive because they have their memory space, file handles, and other resources. Creating a new process involves duplicating the entire process structure.

#### **4. Resource Efficiency**

Threads are more resource efficient compared to processes since they share resources like memory and file handles with other threads in the same process.

#### **5. Isolation**

Processes provide strong isolation, making them more robust. If one process crashes, it typically does not affect other processes. Threads within the same process are less isolated. If one thread crashes due to a bug, it can potentially affect other threads within the same process.

#### **6. Creation Time**

Creating a new thread within an existing process is faster compared to creating a new process. This is because threads share the same memory space as the parent process. Creating a new process involves more overhead and takes longer compared to creating a new thread.

## Question 92 | Concurrency, Multithreading and Processes

What is the difference between a process and a program ?

### *Answer*

A **program** is a set of instructions stored on disk, while a process is an instance of a program in execution. A program becomes a process when it's loaded into memory and executed.

## Question 93 | Concurrency, Multithreading and Processes

What is multithreading ?

### *Answer*

**Multithreading** is a concurrent execution technique where multiple threads run independently within a single process.

Each thread represents a separate flow of execution and can perform tasks concurrently, allowing for efficient utilization of multi core processors.

## Question 94 | Concurrency, Multithreading and Processes

What is a race condition in multithreading, and how can you prevent it ?

### *Answer*

A **race condition** occurs when multiple threads access shared data concurrently, and at least one of them modifies the data. This can lead to unpredictable behavior and bugs.

To prevent race conditions, you can use synchronization mechanisms like mutexes or other atomic operations to ensure that only one thread accesses or modifies the shared data at a time.

## Question 95 | Concurrency, Multithreading and Processes

What synchronization primitives are there in C++ ?

### Answer

There are different synchronization primitives available to help manage the concurrent execution of threads and protect shared resources.

Here are some common synchronization primitives.

1. **Mutex** (`std::mutex`)

Mutex is a fundamental synchronization primitive. It allows only one thread at a time to access a protected resource. Threads attempting to lock a mutex will block until the mutex is unlocked by the owning thread.

2. **Recursive Mutex** (`std::recursive_mutex`)

A recursive mutex is similar to a regular mutex but allows the same thread to lock it multiple times without causing a deadlock. Each lock must be matched with an equivalent number of unlocks.

3. **Timed Mutex** (`std::timed_mutex`)

Timed mutex is an extension of the basic mutex that allows a thread to attempt to lock it with a specified timeout. If the mutex cannot be acquired within the specified time, the thread can take alternative actions.

4. **Lock Guard** (`std::lock_guard`)

`std::lock_guard` is a scoped locking mechanism that automatically locks and unlocks a mutex within a specific scope. It ensures that the mutex is released when the `std::lock_guard` goes out of scope.

5. **Unique Lock** (`std::unique_lock`)

`std::unique_lock` is similar to `std::lock_guard` but offers more flexibility. It allows you to manually lock and unlock the associated mutex and provides features like deferred locking and condition variables.

6. **Condition Variable** (`std::condition_variable`)

Condition variables are used to block one or more threads until a certain condition is met. They are often used in combination with mutexes to build synchronization patterns like producer-consumer scenarios.

#### 7. Atomic Operations (`std::atomic`)

Atomic operations ensure that **read-modify-write** operations on shared variables are executed atomically without interference from other threads. They are provided through the `<atomic>` header and include atomic types like `std::atomic_int`.

#### 8. Semaphore (`std::semaphore` in C++20)

Semaphore is a synchronization primitive introduced in C++20 that allows you to control access to a shared resource. It can be used to limit the number of threads that can access the resource simultaneously.

#### 9. Futures and Promises (`std::future`, `std::promise`)

Futures and promises provide a way to communicate between threads by allowing one thread to retrieve a value computed by another thread asynchronously. They are used in scenarios where one thread produces a result, and another thread consumes it.

## Question 96 | Concurrency, Multithreading and Processes

Can you explain the concept of thread pool ?

### *Answer*

A **thread pool** is a collection of worker threads that are created and managed in advance. These threads are kept alive throughout the lifetime of an application and can be reused to perform tasks concurrently.

Thread pools are used to manage the overhead of thread creation and destruction, making them more efficient for tasks that require multithreading.

## Question 97 | Concurrency, Multithreading and Processes

What is a condition variable, and how is it used in C++ multithreading ?

### **Answer**

A condition variable (`std::condition_variable`) is a synchronization primitive used to block a thread until a specific condition is met. It is used in conjunction with a mutex to protect shared data. Threads can wait on a condition variable until another thread signals that the condition is satisfied, allowing them to proceed.

## Question 98 | Concurrency, Multithreading and Processes

What is the difference between Mutex and Semaphore ?

### **Answer**

#### **Mutex**

A mutex is primarily used to protect shared resources from concurrent access by multiple threads. It allows only one thread to access the protected resource at a time, ensuring exclusive access.

#### **Semaphore**

A semaphore is a more versatile synchronization primitive used to control access to a shared resource or to manage the number of threads allowed to perform certain operations. Semaphores can allow multiple threads to access a resource simultaneously, depending on their count.

## Question 99 | Concurrency, Multithreading and Processes

What is the difference between user space threads and kernel space threads ?

### **Answer**

User space threads and kernel threads are two different approaches to managing threads in a multithreaded program.

Here are the main differences between them.

User space threads are managed entirely by the application or user-level thread library without kernel intervention. The operating system kernel is unaware of the existence of user-level threads. Thread creation, scheduling, and synchronization are all handled in user space.

Kernel threads, on the other hand, are managed and scheduled by the operating system kernel. Each kernel thread corresponds directly to a thread of execution in the application. The kernel is responsible for creating, scheduling, and managing kernel threads.

## Question 100 | Concurrency, Multithreading and Processes

What is the difference between Mutexes and Atomics in C++ ?

### **Answer**

Mutexes are primarily used for providing exclusive access to shared resources, ensuring that only one thread can access the protected resource at a time. They are commonly used to prevent data races and implement critical sections.

Atomics are used for low level, lock-free synchronization operations on individual variables. They are designed to work without the need for locks or mutexes, making them suitable for scenarios where performance is critical.

## Question 101 | Concurrency, Multithreading and Processes

What are the advantages of Atomic over Mutexes ?

### Answer

**Atomic** operations are typically faster than mutex-based synchronization because they avoid the overhead associated with locking and unlocking mutexes. This can lead to better overall program performance, especially in scenarios with frequent read and write operations on shared variables.

**Mutexes** protect **entire critical sections**, which can lead to contention among threads if they are trying to access different parts of the same resource concurrently. Atomics allow for **finer-grained** synchronization at the level of individual variables, reducing contention and enabling more parallelism.

**Mutexes** can lead to deadlocks when threads are waiting for resources that are held by other threads. Atomic operations do not introduce the possibility of deadlocks because they operate on individual variables without blocking.

## Question 102 | Concurrency, Multithreading and Processes

What is the difference between lock-free and wait-free ?

### Answer

A data structure or algorithm is considered **lock-free** if it guarantees that at least one thread will make progress (complete its operation) within a finite number of steps, regardless of the number of threads and their contention for resources.

A data structure or algorithm is considered **wait-free** if it guarantees that every thread will complete its operation within a finite number of steps, regardless of contention and the number of threads.

## Question 103 | Concurrency, Multithreading and Processes

What is the ABA problem in concurrent programming, and how can it be mitigated ?

### ***Answer***

The **ABA** problem occurs when a value at a memory location is changed from **A** to **B** and then back to **A**, but a thread may incorrectly assume that the value hasn't changed.

To mitigate the ABA problem, you can use techniques like **double-check** locking, atomic versions of **compare-and-swap** (CAS) with tagged values, or memory reclamation mechanisms such as **hazard** pointers.

## Question 104 | Concurrency, Multithreading and Processes

### What is Context Switching ?

#### Answer

**Context switching** is a fundamental concept in multithreading and operating system design. It refers to the process of saving and restoring the state of a CPU or thread so that it can switch from one task (or thread) to another.

Context switching allows a system to efficiently multitask by giving the illusion that multiple threads or processes are executing simultaneously, even though only one thread can execute on a single CPU core at any given moment.

Here's how context switching typically works.

#### 1. Saving State

When the operating system decides to switch from one thread to another (for example, due to time-sharing or a thread voluntarily yielding the CPU), it saves the current thread's execution context. This includes the values of CPU registers, program counter, stack pointer, and other relevant information. This context is saved in a data structure called a **"context"** or **"thread control block"**.

#### 2. Loading State

The operating system then loads the saved context of the next thread that should run. This includes setting the CPU registers and program counter to the values stored in the context of the new thread.

#### 3. Switching Execution

With the state loaded, the CPU now begins executing instructions from the new thread's code. This gives the appearance of concurrent execution.

#### 4. Repeat

The process continues, allowing multiple threads to execute in a time-sliced or concurrent manner.

## Question 105 | Concurrency, Multithreading and Processes

What are thread local storage (TLS) variables and when would you use them ?

### ***Answer***

Thread local storage variables are variables that have a separate instance for each thread.

They are used when you need to store thread specific data, such as thread specific configuration settings or cached data. C++ provides `thread_local` keyword for declaring thread local variables.

## Question 106 | Concurrency, Multithreading and Processes

Explain the Happens-Before relationship in the context of memory consistency.

### ***Answer***

The **Happens-Before** relationship is a concept in memory consistency. It defines the order in which memory operations appear to occur to ensure that threads observe a consistent view of shared data. For example, if thread A releases a lock and thread B acquires the same lock, the memory operations in thread A that happened before releasing the lock are guaranteed to be visible to thread B after it acquires the lock.

**Happens-Before** is essential for avoiding data races.

## Question 107 | Concurrency, Multithreading and Processes

**What are starvation and deadlock in concurrent programming, and how can they be prevented or mitigated ?**

### ***Answer***

**Starvation** occurs when a thread is unable to make progress due to other threads constantly accessing shared resources.

**Deadlock** occurs when two or more threads are unable to proceed because they are each waiting for a resource held by another. To prevent or mitigate starvation and deadlock, techniques like fair scheduling, resource ordering, and deadlock detection algorithms can be employed.

## Question 108 | Concurrency, Multithreading and Processes

**Explain the concept of memory barriers and their importance in concurrent programming.**

### ***Answer***

Memory barriers (**fences**) are synchronization primitives used to control the order of memory operations in multithreaded programs. They ensure that certain memory reads and writes become visible to other threads in the expected order. Memory barriers are used to achieve correct memory synchronization, preventing data races, and enforcing memory consistency across threads.

## Question 109 | Concurrency, Multithreading and Processes

How many memory orders are available in C++ atomic operations ?

### **Answer**

There are **six** memory orders available in C++ atomic operations.

1. `std::memory_order_seq_cst`
2. `std::memory_order_relaxed`
3. `std::memory_order_consume`
4. `std::memory_order_acquire`
5. `std::memory_order_release`
6. `std::memory_order_acq_rel`

## Question 110 | Concurrency, Multithreading and Processes

What are memory orders in C++ atomic operations important ?

### **Answer**

Memory orders in C++ atomic operations specify the constraints on memory accesses and the visibility of operations in a multithreaded environment. They are important for controlling the synchronization, consistency, and performance of concurrent programs.

## Question 111 | Concurrency, Multithreading and Processes

Describe the characteristics of `memory_order_relaxed` in C++ atomic operations.

### Answer

`std::memory_order_relaxed` provides the fewest constraints on memory operations. It allows memory accesses to be reordered freely, making it suitable for situations where strict synchronization is not required. It offers the best performance but provides minimal synchronization guarantees.

## Question 112 | Concurrency, Multithreading and Processes

Explain the purpose of `memory_order_acquire` in C++ atomic operations. When should you use it ?

### Answer

`std::memory_order_acquire` ensures that all memory operations performed before the atomic operation (including reads) are visible to the current thread. It's used when you want to establish a `happens-before` relationship between previous operations and the current atomic operation, ensuring proper synchronization.

## Question 113 | Concurrency, Multithreading and Processes

How does `memory_order_consume` differ from other memory orders in C++ atomic operations ?

### Answer

`std::memory_order_consume` is a memory order with a special purpose.

It is used for situations where a read operation depends on a preceding read, but not on writes. It ensures that the preceding reads are visible but does not provide strong synchronization with respect to writes.

## Question 114 | Concurrency, Multithreading and Processes

How do memory orders impact the performance of multithreaded programs in C++ ?

### **Answer**

Memory orders can have a significant impact on performance.

Stricter memory orders (e.g., `std::memory_order_seq_cst`, `std::memory_order_acquire` and `std::memory_order_release`) provide stronger synchronization but may introduce more overhead. Using a relaxed memory order (e.g., `std::memory_order_relaxed`) can improve performance but may require careful design to avoid data races.

## Question 115 | Concurrency, Multithreading and Processes

What is default memory ordering in C++ Atomics ?

### **Answer**

Default memory order is `std::memory_order_seq_cst`.

## Question 116 | Concurrency, Multithreading and Processes

What is the purpose of `memory_order_seq_cst` in C++ atomic operations ?

### **Answer**

`std::memory_order_seq_cst` is a memory order that provides the strongest synchronization guarantees. It ensures that all memory operations appear to be executed in a globally agreed-upon order, following the program's logical order. It prevents reordering of both reads and writes and enforces a total order on memory operations.

## Question 117 | Concurrency, Multithreading and Processes

**Explain the concept of memory coherence in the context of multithreaded programming and how memory barriers relate to it.**

### ***Answer***

Memory coherence refers to the property that all threads in a multithreaded program observe a consistent view of memory.

Memory barriers are essential for maintaining memory coherence. They ensure that memory operations are ordered and synchronized, preventing inconsistencies that could arise due to out-of-order execution.

## Question 118 | Concurrency, Multithreading and Processes

**How to achieve proper synchronization using acquire and release memory orders together ?**

### ***Answer***

Proper synchronization in multithreaded programming can be achieved by using acquire and release memory orders together in a coordinated way.

This technique is known as **acquire-release synchronization**.

Release Operation (**`std::memory_order_release`**)

The releasing thread performs a write operation with **`std::memory_order_release`**. This ensures that all previous writes to memory by the releasing thread are visible to other threads before the release operation itself. It acts as a one-way barrier, ensuring that writes before it are visible to other threads.

Acquire Operation (**`std::memory_order_acquire`**)

The acquiring thread performs a read operation with **`std::memory_order_acquire`**. This ensures that all subsequent reads by the acquiring thread occur after the acquire operation,

thus preventing reads from happening before the acquire operation. It acts as a one-way barrier, ensuring that reads after it are visible only after the acquire operation.

### Example

```
#include <atomic>
#include <thread>
#include <cassert>

std::atomic<bool> x(false);
std::atomic<bool> y(false);
std::atomic<int> z(0);

void write_thread() {
    x.store(true, std::memory_order_release);
    y.store(true, std::memory_order_release);
}

void read_thread() {
    while (!y.load(std::memory_order_acquire)) {}

    if (x.load(std::memory_order_acquire)) {
        ++z;
    }
}

int main() {

    std::thread a(write_thread);
    std::thread b(read_thread);

    a.join();
    b.join();

    // z must be incremented
    assert(z.load() != 0);

    return 0;
}
```

## Question 119 | Concurrency, Multithreading and Processes

What is the use of `std::atomic_thread_fence`?

### Answer

`std::atomic_thread_fence` is a C++ Standard Library function that allows you to insert a memory barrier at a specific point in your code based on the provided memory barrier.

Memory barriers are used in multithreaded programming to enforce memory ordering constraints between threads. These barriers ensure that memory operations before the barrier are visible and ordered correctly with respect to memory operations after the barrier.

### Example

```
#include <atomic>
#include <thread>
#include <cassert>

bool x = false;
std::atomic<bool> y = false;
std::atomic<int> z = 0;

void write_thread() {
    // Write to the shared variables
    x = true;
    std::atomic_thread_fence(std::memory_order_release);
    y.store(true, std::memory_order_relaxed);
}

void read_thread() {
    // Spin until y becomes true
    while (! y.load(std::memory_order_relaxed)) {}

    std::atomic_thread_fence(std::memory_order_acquire);

    // If y is true, then x must be true as well
    if (x) {
        ++z;
    }
}

int main() {
```

```
std::thread a(write_thread);  
std::thread b(read_thread);  
  
a.join();  
b.join();  
  
assert(z.load() != 0);  
  
return 0;  
}
```

## Question 120 | Concurrency, Multithreading and Processes

What is a spinlock ?

### **Answer**

A **spinlock** is a synchronization mechanism used in multithreaded programming to protect shared resources from simultaneous access by multiple threads.

Unlike traditional locks which put a thread to sleep if it cannot acquire the lock immediately, a spinlock repeatedly checks the lock's availability in a tight loop, **spinning** until it can acquire the lock. In other words, it keeps the thread actively waiting instead of yielding the CPU to the operating system.

## Question 121 | Concurrency, Multithreading and Processes

How would you implement spinlock in C++ ?

**Answer**

```
#include <atomic>

class SpinLock {
public:
    SpinLock() {
        // Initialize the flag as unlocked
        flag_.clear();
    }

    void lock() {
        while(flag_.test_and_set(std::memory_order_acquire)) {
            // Spin until the lock is acquired...
        }
    }

    void unlock() {
        // Release the lock
        flag_.clear(std::memory_order_release);
    }

private:
    std::atomic_flag flag_;
};
```

## Question 122 | Concurrency, Multithreading and Processes

Why is Spinlock implemented using `atomic_flag`, not `atomic bool` ?

**Answer**

`std::atomic_flag`

Guarantees to occupy only one byte of memory, making it very memory efficient.

`std::atomic<bool>`

May occupy more space depending on the platform and compiler, potentially consuming more memory.

## Question 123 | Concurrency, Multithreading and Processes

What is thread affinity ?

**Answer**

**Thread affinity** refers to the association of a thread in a multithreaded program with a specific processing resource, typically a CPU core. It determines on which core or set of cores a thread should execute.

## Question 124 | Concurrency, Multithreading and Processes

What are magic static variables ?

**Answer**

**Magic static** variables refer to local static variables inside functions. In C++11 and later standards, the initialization of local static variables is guaranteed to be thread-safe.

## Question 125 | Concurrency, Multithreading and Processes

What is process priority, and how does it affect process scheduling ?

**Answer**

**Process priority** is a value assigned to a process, indicating its importance. Higher priority processes are given preference by the scheduler and get more CPU time.

## Question 126 | Concurrency, Multithreading and Processes

What is a zombie process ?

**Answer**

A **zombie process** is a terminated process that has completed execution but still has an entry in the process table. It exists until its parent process acknowledges its termination, after which it's removed from the table.

## Question 127 | Concurrency, Multithreading and Processes

How does an operating system manage processes ?

**Answer**

The operating system manages processes through **process control blocks** (PCBs), which contain information about each process, including its state, program counter, registers, and more.

The operating system's scheduler determines which process gets CPU time.

## Question 128 | Concurrency, Multithreading and Processes

How many process schedulers do you know in Linux ?

### Answer

#### 1. Completely Fair Scheduler (CFS)

The CFS is the default scheduler in modern Linux kernels (2.6.23 onwards). It uses a concept called **fair scheduling** to ensure that all processes get a fair share of CPU time. Processes are assigned a **weight** based on their priority, and the scheduler strives to allocate CPU time in proportion to these weights. CFS is designed to be more scalable and efficient than its predecessor, the O(1) scheduler.

#### 2. Round Robin Scheduler

The **Round Robin (RR) scheduler** assigns an equal time quantum (or time slice) to each process in a circular order. When a process's time quantum expires, it's moved to the back of the queue, allowing the next process to run. RR is straightforward and ensures that no process hogs the CPU for an extended period. However, it may not be as efficient in managing processes with varying workloads.

#### 3. FIFO Scheduling

**FIFO scheduling** is one of the simplest scheduling algorithms. It follows a strict **first-come first-served** order, where the process that arrives first gets the CPU first. Once a process begins execution, it continues until it completes or voluntarily yields the CPU. There is no concept of priority or fairness in FIFO scheduling. While this approach is straightforward, it can lead to poor utilization of CPU resources because long running processes can block others from executing, resulting in potential starvation.

#### 4. Real Time Schedulers

Linux supports real time scheduling policies for applications that require strict timing guarantees.

The two primary real-time schedulers are **SCHED\_FIFO** (First-In-First-Out) and **SCHED\_RR** (Round Robin).

These policies prioritize real-time tasks over normal tasks and allow them to meet specific deadlines.

## 5. Priority-Based Scheduling

Linux processes can be assigned different priorities ranging from **-20** (highest) to **19** (lowest) using the **nice** command.

The scheduler gives preferential treatment to processes with higher priorities.

## 6. Cgroup-Based Scheduling

**Control Groups** (cgroups) allow administrators to group processes together and allocate resources to these groups. The cgroup subsystem can influence process scheduling by setting resource limits, such as CPU shares and memory usage, for specific groups.

## 7. Deadline Scheduler (SCHED\_DEADLINE)

This is a real-time scheduling policy designed for applications with hard deadlines. Processes under the **SCHED\_DEADLINE** policy specify their execution time and deadline, and the scheduler guarantees that they meet their deadlines.

## Question 129 | Concurrency, Multithreading and Processes

What is IPC ?

### ***Answer***

IPC stands for Inter-Process Communication. It is a set of techniques and mechanisms that allow processes to communicate and share data with each other in a computing environment. IPC is crucial for enabling collaboration and coordination between processes running concurrently on the same machine or across a network.

## Question 130 | Concurrency, Multithreading and Processes

How many types of IPC are there ?

### ***Answer***

Here's a list of the common IPC mechanisms.

#### **1. Unnamed Pipes (Pipe)**

Unnamed pipes provide a simple form of IPC for communication between related processes. They are unidirectional and are typically used for communication between a parent process and its child.

Example of using unnamed pipes in C++ for IPC between a parent and child process. *This example assumes a POSIX environment.*

```
#include <iostream>
#include <unistd.h>
#include <sys/wait.h>

int main() {
    int pipefd[2];
    pid_t cpid;
    char buf;

    if (pipe(pipefd) == -1) {
        perror("pipe");
        exit(EXIT_FAILURE);
    }

    cpid = fork();
    if (cpid == -1) {
        perror("fork");
        exit(EXIT_FAILURE);
    }

    if (cpid == 0) { /* Child reads from pipe */
        close(pipefd[1]); /* Close unused write end */

        while (read(pipefd[0], &buf, 1) > 0) {
            write(STDOUT_FILENO, &buf, 1);
        }

        write(STDOUT_FILENO, "\n", 1);
        close(pipefd[0]);
        _exit(EXIT_SUCCESS);
    } else { /* Parent writes to pipe */
        close(pipefd[0]); /* Close unused read end */
        write(pipefd[1], "Hello Child", 12);
        close(pipefd[1]); /* Reader will see EOF */
        wait(NULL); /* Wait for child */
        exit(EXIT_SUCCESS);
    }
}
```

In this example, the parent process creates a pipe and then forks a child process. The parent writes the string “Hello Child” to the pipe, and the child reads this string from the pipe and writes it to stdout. The pipe is **unidirectional**, so data written by the parent is read by the child.

## 2. Named Pipes (FIFOs)

Named pipes, also known as FIFOs, are similar to unnamed pipes but have a name in the file system. They allow bidirectional communication between processes and can be used for IPC between unrelated processes.

This example includes two programs, one that writes to the pipe first and then reads from it, and another that reads from the pipe first and then writes to it.

Program 1 | Writes First | program1.cpp

```
#include <fcntl.h>
#include <sys/stat.h>
#include <sys/types.h>
#include <unistd.h>
#include <string.h>
#include <stdio.h>

int main() {
    int fd;
    const char * m_fifo = "/tmp/myfifo";
    mkfifo(m_fifo, 0666);
    char w_data[80], r_data[80];
    while (1) {
        fd = open(m_fifo, O_WRONLY);
        fgets(w_data, 80, stdin);
        write(fd, w_data, strlen(w_data) + 1);
        close(fd);
        fd = open(m_fifo, O_RDONLY);
        read(fd, r_data, sizeof(r_data));
        printf("Process-2: %s\n", r_data);
        close(fd);
    }
    return 0;
}
```

Program 2 | Reads First | program2.cpp

```

#include <fcntl.h>
#include <sys/stat.h>
#include <sys/types.h>
#include <unistd.h>
#include <string.h>
#include <stdio.h>

int main() {
    int fd;
    const char * m_fifo = "/tmp/myfifo";
    mkfifo(m_fifo, 0666);
    char r_data[80], w_data[80];
    while(1) {
        fd = open(m_fifo, O_RDONLY);
        read(fd, r_data, 80);
        printf("Process-1: %s\n", r_data);
        close(fd);
        fd = open(m_fifo, O_WRONLY);
        fgets(w_data, 80, stdin);
        write(fd, w_data, strlen(w_data) + 1);
        close(fd);
    }
    return 0;
}

```

### 3. Shared Memory

Shared memory allows multiple processes to share a region of memory. It provides fast and efficient communication because processes can read and write directly to shared memory locations.

In this example the process one will attach to shared memory using specific key, write to it and then detach from it, later the second process will attach to shared memory, read the data, detach from shared memory and close the shared memory using `shmctl` function with the `IPC_RMID` command, this will mark the segment to be destroyed. The segment will actually be destroyed only after all processes have detached it (i.e., no process is attached to it).

Program 1 | Writes to the shared memory segment

```

#include <sys/ipc.h>
#include <sys/shm.h>
#include <stdio.h>
#include <string.h>

```

```

#include <pthread.h>

#define SHM_KEY 1234

struct shared_use_st {
    pthread_mutex_t mutex;
    char data[1024];
};

int main() {
    int shmid = shmget(SHM_KEY, sizeof(struct shared_use_st),
                      0666 | IPC_CREAT);
    if (shmid == -1) {
        perror("Shared memory");
        return 1;
    }

    struct shared_use_st* shared_stuff =
        (struct shared_use_st*) shmat(shmid, NULL, 0);

    if (shared_stuff == (void*) -1) {
        perror("Shared memory attach");
        return 1;
    }

    pthread_mutexattr_t attr;
    pthread_mutexattr_init(&attr);
    pthread_mutexattr_setpshared(&attr, PTHREAD_PROCESS_SHARED);
    pthread_mutex_init(&shared_stuff->mutex, &attr);

    pthread_mutex_lock(&shared_stuff->mutex);
    strncpy(shared_stuff->data, "Hello from process 1",
            sizeof(shared_stuff->data));
    pthread_mutex_unlock(&shared_stuff->mutex);

    // Detach from the shared memory segment
    if (shmdt(shared_stuff) == -1) {
        perror("shmdt");
        return 1;
    }

    return 0;
}

```

Program 2 | Reads from shared memory segment

```

#include <sys/ipc.h>
#include <sys/shm.h>
#include <stdio.h>
#include <string.h>
#include <pthread.h>

#define SHM_KEY 1234

struct shared_use_st {
    pthread_mutex_t mutex;
    char data[1024];
};

int main() {
    int shmid = shmget(SHM_KEY, sizeof(struct shared_use_st),
                      0666 | IPC_CREAT);
    if (shmid == -1) {
        perror("Shared memory");
        return 1;
    }

    struct shared_use_st* shared_stuff =
        (struct shared_use_st*) shmat(shmid, NULL, 0);

    if (shared_stuff == (void*) -1) {
        perror("Shared memory attach");
        return 1;
    }

    pthread_mutex_lock(&shared_stuff->mutex);
    printf("Received message: %s\n", shared_stuff->data);
    pthread_mutex_unlock(&shared_stuff->mutex);

    // Detach from the shared memory segment
    if (shmdt(shared_stuff) == -1) {
        perror("shmdt");
        return 1;
    }

    // Delete the shared memory segment
    if (shmctl(shmid, IPC_RMID, NULL) == -1) {
        perror("shmctl");
        return 1;
    }

    return 0;
}

```

#### 4. Message Queues

Message queues enable processes to send and receive messages asynchronously. They are often used for communication between unrelated processes and can support more complex messaging patterns.

##### Example

*sender.cpp*

```
#include <mqueue.h>
#include <iostream>
#include <cstring>
#include <stdlib.h>

int main() {
    mqd_t mq;
    struct mq_attr attr;
    char* buffer;

    // Open the message queue
    mq = mq_open("/test_queue",
                O_CREAT | O_WRONLY, 0644, NULL);

    if (mq == (mqd_t) -1) {
        perror("mq_open");
        exit(1);
    }

    // Query the message queue attributes
    if (mq_getattr(mq, &attr) == -1) {
        perror("mq_getattr");
        exit(1);
    }

    // Allocate a buffer of the appropriate size
    buffer = new char[attr.mq_msgsize];

    std::cout << "Send a message: ";
    std::cin.getline(buffer, attr.mq_msgsize);

    // Send the message
```

```

    if (mq_send(mq, buffer, strlen(buffer) + 1, 0) == -1) {
        perror("mq_send");
        exit(1);
    }

    // Free the buffer
    delete[] buffer;

    // Close the message queue
    if (mq_close(mq) == -1) {
        perror("mq_close");
        exit(1);
    }

    return 0;
}

```

*receiver.cpp*

```

#include <mqueue.h>
#include <iostream>
#include <cstring>
#include <cstdlib>

int main() {
    mqd_t mq;
    struct mq_attr attr;

    // Open the message queue
    mq = mq_open("/test_queue",
                O_CREAT | O_RDONLY, 0644, NULL);
    if (mq == (mqd_t) -1) {
        perror("mq_open");
        exit(1);
    }

    // Query the message queue attributes
    if (mq_getattr(mq, &attr) == -1) {
        perror("mq_getattr");
        exit(1);
    }

    // Allocate a buffer of the appropriate size
    char* buffer = new char[attr.mq_msgsize];

    // Receive the message

```

```

    if (mq_receive(mq, buffer, attr.mq_msgsize, NULL) == -1) {
        perror("mq_receive");
        exit(1);
    }

    std::cout << "Received message: " << buffer << std::endl;

    // Free the buffer
    delete[] buffer;

    // Close the message queue
    if (mq_close(mq) == -1) {
        perror("mq_close");
        exit(1);
    }

    // Unlink the message queue
    if (mq_unlink("/test_queue") == -1) {
        perror("mq_unlink");
        exit(1);
    }

    return 0;
}

```

## 5. Semaphore

Semaphores are synchronization primitives used to coordinate access to shared resources or protect critical sections of code. They can be used for signaling between processes.

### Example

*shared.hpp*

```

#pragma once
#include <semaphore.h>

const char* const SEM_MUTEX_NAME = "/sem_mutex";
const char* const SEM_BUFFER_COUNT_NAME = "/sem_buffer_count";
const char* const SEM_BUFFER_SPACE_NAME = "/sem_buffer_space";
const char* const SHARED_MEM_NAME = "/shared_mem";
const int BUFFER_SIZE = 10;

```

*producer.cpp*

```

#include <iostream>
#include <fcntl.h>
#include <sys/mman.h>
#include <semaphore.h>
#include <unistd.h>
#include <cstdlib>

#include "shared.hpp"

int main() {
    sem_t* mutex = sem_open(SEM_MUTEX_NAME, O_CREAT, 0666, 1);
    sem_t* buffer_count = sem_open(SEM_BUFFER_COUNT_NAME,
                                   O_CREAT, 0666, 0);
    sem_t* buffer_space = sem_open(SEM_BUFFER_SPACE_NAME,
                                   O_CREAT, 0666, BUFFER_SIZE);

    int shm_fd = shm_open(SHARED_MEM_NAME,
                          O_CREAT | O_RDWR, 0666);

    ftruncate(shm_fd, BUFFER_SIZE * sizeof(int));

    int* buffer = (int*) mmap(0, BUFFER_SIZE * sizeof(int),
                              PROT_WRITE, MAP_SHARED,
                              shm_fd, 0);

    srand(time(NULL));

    for (int i = 0; i < 10; ++i) {
        sem_wait(buffer_space);
        sem_wait(mutex);
        int num = rand() % 100;
        buffer[i % BUFFER_SIZE] = num;
        std::cout << "Produced: " << num << std::endl;

        sem_post(mutex);
        sem_post(buffer_count);
        sleep(1); // Sleep for a second to simulate work
    }

    munmap(buffer, BUFFER_SIZE * sizeof(int));
    close(shm_fd);
    sem_close(mutex);
    sem_close(buffer_count);
    sem_close(buffer_space);
}

```

```

    return 0;
}

```

### consumer.cpp

```

#include <iostream>
#include <fcntl.h>
#include <sys/mman.h>
#include <semaphore.h>
#include <unistd.h>

#include "shared_data.h"

int main() {
    sem_t* mutex = sem_open(SEM_MUTEX_NAME, O_CREAT, 0666, 1);
    sem_t* buffer_count = sem_open(SEM_BUFFER_COUNT_NAME,
                                    O_CREAT, 0666, 0);
    sem_t* buffer_space = sem_open(SEM_BUFFER_SPACE_NAME,
                                    O_CREAT, 0666, BUFFER_SIZE);

    int shm_fd = shm_open(SHARED_MEM_NAME,
                          O_CREAT | O_RDWR, 0666);
    int* buffer = (int*) mmap(0, BUFFER_SIZE * sizeof(int),
                              PROT_READ, MAP_SHARED, shm_fd, 0);

    for (int i = 0; i < 10; ++i) {
        sem_wait(buffer_count);
        sem_wait(mutex);

        int num = buffer[i % BUFFER_SIZE];
        std::cout << "Consumed: " << num << std::endl;

        sem_post(mutex);
        sem_post(buffer_space);
        sleep(1); // Sleep for a second to simulate work
    }

    munmap(buffer, BUFFER_SIZE * sizeof(int));
    close(shm_fd);
    sem_close(mutex);
    sem_close(buffer_count);
    sem_close(buffer_space);

    return 0;
}

```

You can compile these programs with `g++ -pthread producer.cpp -o producer` and `g++ -pthread consumer.cpp -o consumer`, then run them in separate terminal windows. The producer will generate random numbers and put them into the buffer, and the consumer will read these numbers from the buffer. They use semaphores to ensure that they don't access the buffer at the same time.

## 6. Shared Files

Processes can communicate by reading and writing to shared files. This approach is often used when data persistence is required.

## 7. Sockets

Sockets enable communication between processes over a network, making them suitable for distributed computing and networked applications. They can be used for both inter-process and inter-host communication.

## 8. Signals

Processes can send and receive signals to notify each other about specific events or conditions. Signals also can be used for inter-process communication as well as for handling exceptions and interrupts.

# Question 131 | Concurrency, Multithreading and Processes

What is the fastest IPC ?

**Answer**

Shared Memory.

**Shared memory** allows processes to read and write to a shared region of memory directly. Since there's minimal copying of data involved, it can offer excellent performance, especially for large data transfers.

The shared memory requires careful synchronization to avoid data corruption.

## Question 132 | Concurrency, Multithreading and Processes

What is RCU in concurrent programming ?

### ***Answer***

**RCU** (Read-Copy-Update) is a synchronization primitive used in concurrent programming to manage data shared among multiple threads or processes efficiently.

RCU is primarily designed for read-heavy workloads, where reads far outnumber writes. It's often used in situations where you need to minimize the overhead of read-side synchronization.

# Template Metaprogramming

---

## Question 133 | Template Metaprogramming

**What is the difference between function template specialization and function overloading ?**

### ***Answer***

Function template specialization is used to provide a customized implementation of a function template for specific data types.

Function overloading, on the other hand, involves defining multiple functions with the same name but different parameter lists.

## Question 134 | Template Metaprogramming

**How can you implement type traits using C++ templates, and why are they essential for template metaprogramming ?**

### ***Answer***

Type traits are implemented using templates to provide compile time information about types. They are essential for template metaprogramming because they allow you to perform type-based conditional compilation and provide generic solutions that work with various types.

## Question 135 | Template Metaprogramming

**Discuss some practical applications of template metaprogramming in real-world C++ codebases.**

### **Answer**

Template metaprogramming is used in various real-world scenarios, such as creating efficient container libraries, implementing serialization frameworks, generating code for code generators, and optimizing numerical libraries. For example, Boost MPL (Meta-Programming Library) uses template metaprogramming to provide compile time algorithms for working with types and sequences.

## Question 136 | Template Metaprogramming

**What is the difference between typename and class in template declarations ?**

### **Answer**

Both typename and class can be used interchangeably to declare template parameters. The typename is more versatile and can also be used for nested dependent types.

There are some cases where it is required to use `class` instead of `template`.

```
template < template < typename, typename >  
class Dummy, typename Type >
```

When specifying a template template, the class keyword must be used as above, it is not interchangeable with typename in this case (since C++17 both keywords are allowed in this case).

You also must use class when explicitly instantiating a template.

```
template class Dummy<int>;
```

## Question 137 | Template Metaprogramming

How would you implement fibonacci using templates ?

**Answer**

```
#include <iostream>

template <int N>
struct Fibonacci {
    static const int value = Fibonacci<N - 1>::value +
                             Fibonacci<N - 2>::value;
};

template <>
struct Fibonacci <0> {
    static const int value = 0;
};

template <>
struct Fibonacci <1> {
    static const int value = 1;
};

int main() {
    constexpr int result = Fibonacci<22>::value;
    std::cout << "Fibonacci (22) = " << result << std::endl;
    return 0;
}
```

## Question 138 | Template Metaprogramming

Write a metafunction to calculate the factorial of a given number at compile time.

### Answer

```
#include <iostream>

template <int N>
struct Factorial {
    static const int value = N * Factorial<N - 1>::value;
};

template <>
struct Factorial<0> {
    static const int value = 1;
};

int main() {
    constexpr int result = Factorial<22>::value;
    std::cout << "Factorial of (22) is: "
               << result << std::endl;
    return 0;
}
```

## Question 139 | Template Metaprogramming

### What is SFINAE ?

#### Answer

SFINAE stands for **Substitution Failure Is Not An Error**.

It is a principle in C++ template metaprogramming that allows the compiler to gracefully handle situations where a substitution of template arguments fails during the compilation process. Instead of resulting in a compilation error, SFINAE causes the compiler to consider other function/template overloads that may be a better match.

SFINAE is particularly useful in cases where you want to enable or disable a template function or specialization based on certain conditions without causing a compilation error.

#### Example

```
#include <iostream>
#include <type_traits>

template <typename T>
typename std::enable_if<std::is_integral<T>::value,
                        void>::type printValue(T value) {
    std::cout << "Integral value: " << value << std::endl;
}

template <typename T>
typename std::enable_if<std::is_floating_point<T>::value,
                        void>::type printValue(T value) {
    std::cout << "Floating-point value: "
              << value << std::endl;
}

int main() {
    printValue (42);
    printValue (3.14);
    printValue ("Hello");
    return 0;
}
```

# STL (Standard Template Library)

---

## Question 140 | Standard Template Library

What are the advantages of `std::vector` ?

### *Answer*

1. **Fast Random Access**

Elements can be accessed in constant time using an index.

2. **Dynamic Resizing**

It can grow or shrink dynamically as elements are added or removed.

3. **Contiguous Memory**

Elements are stored in adjacent memory locations, improving cache locality.

4. **Iteration**

It supports range based for loops and STL algorithms.

5. **Standard Library Compatibility**

It is a part of the C++ Standard Library, ensuring portability.

6. **Automatic Memory Management**

Memory allocation and deallocation are managed by the container.

## Question 141 | Standard Template Library

What is the time complexity of common operations in `std::vector` ?

### Answer

Accessing an element is  $O(1)$

Insertion at the end (**amortized**) is  $O(1)$

Insertion at any other position is  $O(n)$  (**linear time due to shifting**)

Deletion at the end is  $O(1)$

Deletion at any other position is  $O(n)$  (**linear time due to shifting**)

Searching for an element is  $O(n)$  (**linear time due to sequential search**)

## Question 142 | Standard Template Library

What is `std::map` and how is it implemented under the hood ?

### Answer

`std::map` is an associative container in the C++ Standard Library that stores elements as key-value pairs, where each key is unique. It is implemented as a **balanced binary search tree** (usually a red-black tree) that maintains its elements in sorted order based on the keys. This allows for efficient lookups, insertions, and deletions while maintaining logarithmic time complexity.

The key features of `std::map` include.

#### 1. Sorted Order

Elements are automatically sorted based on the keys, providing fast access in a sorted manner.

#### 2. Unique Keys

Each key can only appear once in the map.

#### 3. Balanced Tree

The tree structure is balanced to ensure relatively equal heights of subtrees, maintaining efficient operations.

#### 4. Logarithmic Complexity

Insertions, deletions, and lookups have a time complexity of  $O(\log n)$ , making it suitable for large datasets.

## Question 143 | Standard Template Library

**What does amortized complexity mean in the context of `std::vector` operations?**

### ***Answer***

In the context of `std::vector`, **amortized complexity** refers to the average cost per operation over a sequence of operations, taking into account both the costly and inexpensive operations. While individual operations on a `std::vector` may have varying time complexities (such as insertion or reallocation), amortized analysis provides insight into the average performance over a series of operations.

Let's consider `std::vector`'s dynamic array resizing behavior during insertions. When the capacity of a `std::vector` is exhausted, a reallocation is triggered to allocate a larger memory block and copy the existing elements. While this reallocation is an expensive operation, it occurs infrequently.

As a result, the amortized complexity of inserting an element remains constant over multiple insertions, even though some individual insertions might involve reallocations.

## Question 144 | Standard Template Library

How is the underlying data structure of `std::priority_queue` implemented ?

### Answer

The underlying data structure of `std::priority_queue` is implemented as a binary heap, specifically a binary max heap by default. A binary heap is a complete binary tree where each parent node's value is greater than or equal to the values of its children. This structure ensures that the highest priority element is at the top, allowing efficient access and removal.

## Question 145 | Standard Template Library

What is `std::unordered_map` ?

### Answer

`std::unordered_map` is a container provided by the C++ Standard Library that implements an associative array, also known as a hash map. It stores key-value pairs, where each key is unique and is associated with a value.

Unlike `std::map`, which uses a binary search tree, `std::unordered_map` uses a hash table for fast access and retrieval.

## Question 146 | Standard Template Library

What is the average time complexity for insertion, retrieval, and deletion operations in `std::unordered_map` ?

### Answer

On average, the time complexity for insertion, retrieval, and deletion operations in `std::unordered_map` is  $O(1)$ , assuming a good hash function and uniform distribution of keys.

In the worst case, these operations can take  $O(n)$  time due to collisions.

## Question 147 | Standard Template Library

What is the role of the hash function in `std::unordered_map` ?

### *Answer*

The hash function in `std::unordered_map` is responsible for mapping keys to indices in the hash table. A good hash function should distribute the keys uniformly across the table to minimize collisions and ensure efficient access.

The C++ Standard Library provides default hash functions for common types, but custom hash functions can also be used.

## Question 148 | Standard Template Library

What is the purpose of the load factor in `std::unordered_map` ?

### *Answer*

The load factor in `std::unordered_map` is the ratio of the number of stored elements to the number of buckets in the hash table. A higher load factor increases the likelihood of collisions, while a lower load factor reduces memory usage but may lead to inefficient hash table utilization.

Rehashing occurs when the load factor exceeds a certain threshold to maintain good performance.

## Question 149 | Standard Template Library

Is it possible to use complex types, such as user defined classes, as keys in `std::unordered_map` ?

### *Answer*

Yes, you can use complex types, including user defined classes, as keys in `std::unordered_map`. However, you need to ensure that the custom hash function and equality function are well defined and provide consistent behavior for those types.

## Question 150 | Standard Template Library

What is the behavior of `std::unordered_map` when the hash function or does the key equality function change after elements have been inserted ?

### Answer

Changing the hash function or key equality function after elements have being inserted into an `std::unordered_map` leads to undefined behavior.

This is because the hash values and bucket assignments would become inconsistent, and elements might not be retrievable or deletable properly.

## Question 151 | Standard Template Library

Why is it not possible to keep references in `std::vector` ?

### Answer

`std::vector<int&>` or any other vector of references is not allowed in C++ because references cannot be reseated to refer to a different object after initialization.

This goes against the nature of how vectors manage their elements and how references work.

Vectors need to be able to reallocate memory and move elements around when they grow or shrink. This process is incompatible with references, which are designed to be bound to a specific object and cannot be rebound.

## Question 152 | Standard Template Library

Is there a workaround to keep references in `std::vector` ?

### Answer

Yes, you can use `std::reference_wrapper` as a workaround to store references in `std::vector`. `std::reference_wrapper` is a class template that provides a reference-like object that can be stored in containers and supports reassignment.

## Example

```
#include <iostream>
#include <vector>
#include <functional>

int main() {
    int a = 5;
    int b = 10;
    int c = 15;

    std::vector<std::reference_wrapper<int>> refVec;
    refVec.push_back(std::ref(a));
    refVec.push_back(std::ref(b));
    refVec.push_back(std::ref(c));

    for (auto ref : refVec) {
        ref += 1;
    }

    for (const auto ref : refVec) {
        std::cout << ref << " " << std::endl;
    }

    return 0;
}
```

## Question 153 | Standard Template Library

**Is it possible and good practice to extend from STL containers ?**

### **Answer**

While it is technically possible to inherit from STL containers, it's generally not recommended due to the following reasons.

#### **1. Violation of Liskov Substitution Principle (LSP)**

Inheriting from STL containers can break the Liskov Substitution Principle, which is a fundamental principle of object oriented programming.

This principle states that objects of a superclass should be replaceable with objects of its subclasses without affecting the correctness of the program. Inheriting from STL containers could lead to unexpected behavior when code expects specific container behavior.

#### **2. Undefined Behavior**

Many STL containers have non virtual destructors, which means that deleting a derived object through a pointer to the base class can lead to undefined behavior.

#### **3. Implementation Details**

STL containers have internal data structures and implementation details that may change between different versions of compilers and libraries. Inheriting from them might expose you to those changes and make your code less portable.

#### **4. Alternative Approaches**

If you need to extend the functionality of an STL container, consider using composition (holding a member variable of the container type) or creating free functions that work with the container.

#### **5. Wrapper Classes**

If you need to extend the behavior of an STL container, consider creating a wrapper class that uses the container internally and exposes the functionality you need. This allows to encapsulate the container and its behavior without inheriting from it.

## Question 154 | Standard Template Library

What is `std::allocator` ?

### Answer

`std::allocator` is a class template in the C++ Standard Library that provides a basic interface for memory allocation and deallocation. It is often used as the default allocator for standard containers like `std::vector`, `std::list`, and `std::map` when you don't specify a custom allocator. The purpose of `std::allocator` is to allocate and deallocate memory for objects in a way that is compatible with the C++ Standard Library's requirements.

Here's a basic overview of how `std::allocator` works.

#### 1. Memory Allocation

The `std::allocator` class provides a member function called `allocate`, which is used to allocate memory for a specified number of objects. When you create a container, like `std::vector`, it uses `std::allocator` to allocate memory for its elements. The `allocate` function is called with the number of elements to allocate space for.

#### 2. Memory Deallocation

The `std::allocator` class also provides a member function called `deallocate`, which is used to deallocate memory previously allocated by the `allocate` function. This is essential for managing memory efficiently and preventing memory leaks.

#### 3. Construction and Destruction

`std::allocator` itself doesn't handle the construction and destruction of objects. These responsibilities usually lie with the container that uses the allocator.

For example, when you add an element to a `std::vector`, the vector will construct the object in the allocated memory using placement `new`.

#### 4. Rebinding

Allocators can be rebound to allocate memory for objects of different types. This is important for standard containers that need to allocate memory for objects of various types. The allocator template can be specialized or rebound to allocate memory for a specific type.

## Question 155 | Standard Template Library

When should you use a custom allocator for `std::vector` ?

### *Answer*

You should use a custom allocator for `std::vector` if you need to improve the performance, memory usage, or portability of your program. You should also use a custom allocator if you need to control the way memory is allocated and deallocated.

## Question 156 | Standard Template Library

What is the difference between set and multiset ?

### *Answer*

The main difference between set and multiset is that set does not allow duplicate elements, while multiset allows duplicate elements.

## Question 157 | Standard Template Library

How many types of Iterators are there in STL ?

### *Answer*

1. Input iterator - read elements
2. Output iterator - write elements
3. Forward iterator - read and move forward
4. Bidirectional iterator - read, move forward, and move backward
5. Random access iterator - read, write, move forward, move backward, and jump to any element

## Question 158 | Standard Template Library

What is a contiguous iterator in STL ?

### Answer

The `contiguous_iterator` concept was introduced in C++17 and replaces the `random_access_iterator` concept.

The `contiguous_iterator` concept is a refinement of the `random_access_iterator` concept, and it adds the following requirements.

*The iterator can be dereferenced to obtain a reference to element it points to.*

*The iterator can be incremented and decremented by one.*

*The iterator can be compared to other iterators for equality and inequality.*

*The iterator can be used to compute the distance between two iterators.*

*The iterator can be used to construct a pointer to the element it points to.*

## Question 159 | Standard Template Library

What is the difference between `random_access_iterator` and `contiguous_iterator` ?

### Answer

The main difference between `random_access_iterator` and `contiguous_iterator` is that `contiguous_iterator` adds the requirement that the iterator can be used to construct a pointer to the element it points to.

This is because not all random access iterators are pointers.

For example, the iterator returned by the `std::vector::begin()` method is a contiguous iterator, but it is not a pointer.

## Question 160 | Standard Template Library

**What is the difference between a mutable iterator and an immutable iterator ?**

### ***Answer***

A mutable iterator can be used to modify the elements that it points to.

An immutable iterators cannot be used to modify the elements that it points to.

## Question 161 | Standard Template Library

**What is the difference between iterator and pointer ?**

### ***Answer***

#### **Safety and range checking**

Iterator comes with built in range checking mechanisms, preventing access to invalid memory locations during container traversal. This enhances the safety and robustness of code.

Pointer does not inherently provide range checking, making it susceptible to issues like buffer overflows and accessing invalid memory locations. Memory safety is the responsibility of the programmer.

#### **Type of Access**

Iterator offers logical access to elements within a container, abstracting away the specifics of memory addresses. It provides a higher level of readability and ease of use.

Pointer provides direct access to memory locations, allowing for more granular control over memory. However, this direct access comes at the cost of increased complexity and potential for errors.

## Question 162 | Standard Template Library

**Are there performance differences when using iterator or reverse iterator ?**

### ***Answer***

Yes, there can be performance differences when using an iterator or a reverse iterator.

A reverse iterator iterates over a container in reverse order. This means that it has to traverse the container from the end to the beginning. This can be slower than iterating over the container in forward order, because the reverse the iterator has to do more work.

## Question 163 | Standard Template Library

**What is `std::sort` default algorithm using under the hood ?**

### ***Answer***

The default algorithm used by `std::sort` is the **IntroSort** algorithm, which is a hybrid sorting algorithm that combines **quicksort** and **heapsort**.

**IntroSort** starts by using quicksort, which is a divide-and-conquer algorithm that is generally very efficient for sorting small and medium-sized arrays.

However, quicksort can be unstable, which means that it can change the relative order of equal elements in the array. If the array is large, IntroSort switches to heapsort, which is a stable sorting algorithm that is generally slower than quicksort but is guaranteed to preserve the relative order of equal elements in the array.

## Question 164 | Standard Template Library

What are some categories of algorithms provided by STL algorithm library ?

### **Answer**

The STL algorithm library can be categorized into several groups.

1. **Non-modifying algorithms**

These algorithms don't modify the elements but provide information or test conditions, such as `std::find`, `std::count`, and `std::all_of`.

2. **Modifying algorithms**

These algorithms modify the elements in a range, like `std::copy`, `std::transform`, and `std::fill`.

3. **Sorting and related algorithms**

Algorithms for sorting and rearranging elements, including `std::sort`, `std::stable_sort`, and `std::rotate`.

4. **Binary search algorithms**

Algorithms that perform binary searches on sorted ranges, such as `std::binary_search`, `std::lower_bound`, and `std::upper_bound`.

5. **Set algorithms**

Algorithms that work with sorted ranges and perform set operations like intersection, union, and difference, such as `std::set_intersection`, `std::set_union`, and `std::set_difference`.

## Question 165 | Standard Template Library

How does `std::reverse` differ from `std::reverse_copy` ?

### **Answer**

`std::reverse` modifies the elements in-place within the given range, reversing their order. `std::reverse_copy`, on the other hand, creates a reversed copy of the original range, leaving the original range unchanged.

## Question 166 | Standard Template Library

What is collision resolution in hash tables ?

### **Answer**

**Collision resolution** is the process of handling situations where two or more keys hash to the same location in a hash table. It's necessary to prevent data loss or overwriting when multiple keys collide.

## Question 167 | Standard Template Library

What is separate chaining in collision resolution ?

### **Answer**

**Separate chaining** is a collision resolution technique in hash tables where each slot of the hash table contains a linked list to store multiple key-value pairs that hash to the same location.

Collisions are resolved by appending elements to the list.

## Question 168 | Standard Template Library

When is open addressing preferred over separate chaining ?

### *Answer*

[Open addressing](#) is preferred when memory usage needs to be minimized, and it's known that the hash table won't be densely populated. It's also preferred when control over memory allocation is necessary.

# Exceptions

---

## Question 169 | Exceptions

How many levels of safety are there in C++ ?

### *Answer*

1. **No guarantee**

This is the lowest level of exception safety. A function with no guarantee can do anything if an exception is thrown, including leaking resources, leaving objects in an inconsistent state, or crashing the program.

2. **Basic guarantee**

A function with the basic guarantee guarantees that no resources will be leaked if an exception is thrown. However, the function may leave objects in an inconsistent state.

3. **Strong guarantee**

A function with the strong guarantee guarantees that the program will be left in a consistent state if an exception is thrown. This means that all objects will be in the same state as they were before the function was called, and no resources will be leaked.

4. **Nothrow guarantee**

A function with the nothrow guarantee guarantees that it will never throw an exception. This is the highest level of exception safety.

## Question 170 | Exceptions

**What is the `std::nothrow` object and how is it used in exception handling ?**

### ***Answer***

The `std::nothrow` object is used to indicate that memory allocation functions (`new`, `new[]`) should not throw exceptions if memory allocation fails.

If memory cannot be allocated, these functions return a null pointer instead of throwing an exception.

This can be useful in scenarios where you want to handle memory allocation failures manually.

## Question 171 | Exceptions

**What is the role of `std::unexpected` and `std::terminate` in exception handling ?**

### ***Answer***

`std::unexpected` is a function that's called when an exception is thrown from a function that's not declared with an exception specification.

`std::terminate` is a function that's called when the exception handling mechanism cannot find a suitable catch block to handle an exception, leading to program termination.

You can customize the behavior of these functions by using `set_unexpected` and `set_terminate` functions.

## Question 172 | Exceptions

**Can you catch an exception by value and by reference at the same time ?**

### ***Answer***

Yes, it's possible to catch an exception both by value and by reference. But it is worth considering that if you catch an exception by value first and then by reference, the catch block with the reference will never be executed.

This is because the catch block with the value parameter already handles the exception, and the control flow does not reach subsequent catch blocks.

## Question 173 | Exceptions

**Can you explain what happens if an exception is thrown within a destructor ?**

### ***Answer***

Throwing exceptions from destructors can lead to unexpected behavior and potential resource leaks. Here are some considerations.

#### **1. Stack Unwinding**

When an exception is thrown within a destructor, the C++ runtime initiates a process known as **stack unwinding**. Stack unwinding means that the destructor of the object currently being destroyed is immediately terminated, and the control flow jumps to the nearest matching catch block, if one exists.

#### **2. Other Destructors**

If the destructor of an object is interrupted due to an exception, it doesn't mean that other destructors will be skipped. All previously constructed objects within the same scope will have their destructors called as part of the stack unwinding process.

#### **3. Memory Leaks**

If the destructor that threw an exception had already performed resource deallocation (e.g.

releasing memory, closing files, or releasing other resources), and the exception interrupted the destructor, it could potentially lead to resource leaks.

#### 4. Exception Handling

If an exception is thrown within a destructor and there's a suitable catch block in the call stack (e.g., in the destructor's caller or higher level functions), the exception can be caught and handled.

If no suitable catch block is found during stack unwinding, the program may terminate.

#### 5. Terminate Function

If an uncaught exception occurs during stack unwinding (i.e., there is no appropriate catch block to handle the exception), the `std::terminate` function is called, which by default causes the program to terminate.

#### 6. Nested Exceptions

If multiple objects within the same scope throw exceptions in their destructors, the order in which the exceptions are handled during stack unwinding is the reverse order of their construction. This is because objects are destructed in the reverse order of their construction.

## Question 174 | Exceptions

What is the role of `std::exception_ptr` in exception handling?

### *Answer*

`std::exception_ptr` is a type that holds a captured exception object. It allows you to capture an exception and rethrow it later in a different context. This is useful for propagating exceptions across threads or for logging and reporting purposes.

## Question 175 | Exceptions

**What is the difference between a rethrown exception and a propagated exception ?**

### **Answer**

A rethrown exception is an exception that is thrown from a catch block and then thrown again from another catch block. A propagated exception is an exception that is thrown from a try block and then propagated to the calling function.

## Question 176 | Exceptions

**What will be the output of this code ?**

```
#include <iostream>

int main() {
    try {
        throw 22;
    }
    catch (char* e) {
        std::cout << "Caught the exception " << e << std::endl;
    }
    catch (...) {
        std::cout << "Default Exception" << std::endl;
    }
    return 0;
}
```

### **Answer**

Default Exception

An integer value `22` is thrown as an exception. It can't be interpreted as `char*` means we need to go with a default catch that is why the **Default Exception** is printed.

## Question 177 | Exceptions

What will print this code ?

```
#include <iostream>

class B {};
class D: public B {};

int main() {
    D d;
    try {
        throw d;
    } catch (B b) {
        std::cout << "Caught Base Exception" << std::endl;
    } catch (D d) {
        std::cout << "Caught Derived Exception" << std::endl;
    }
    return 0;
}
```

### Answer

Caught Base Exception

If both base and derived classes are caught as exceptions, then the catch block of the derived class must appear before the base class because if we put the base class first then the derived class catch block will never be reached.

# Design Patterns and Architectural Concepts

---

## Question 178 | Design Patterns and Architectural Concepts

What is the rule of zero in C++ ?

### **Answer**

**The Rule of Zero** in C++ is a programming guideline that states that you should avoid defining any special member functions (constructors, destructors, copy constructors, move constructors, and assignment operators) for your classes unless you have a good reason to do so. The compiler will provide default implementations for these functions if you don't define them explicitly.

The rationale behind the Rule of Zero is that it makes your code simpler and easier to understand. When you don't have to worry about defining special member functions, you can focus on the essential functionality of your class. This can also lead to more efficient code, as the compiler can optimize the default implementations of these functions.

**There are a few cases where you might need to define special member functions.**

1. *If your class owns resources that need to be managed, such as memory or file handles.*
2. *If your class needs to be compatible with a legacy API that requires custom implementations of these functions.*
3. *If you need to implement special behavior for these functions, such as preventing object copies or moving objects between different memory contexts.*

In general, you should follow the Rule of Zero and only define special member functions when you have a good reason to do so. This will make your code simpler, easier to understand, and more efficient.

## Question 179 | Design Patterns and Architectural Concepts

What is the rule of three in C++ ?

### **Answer**

The [Rule of Three](#) in C++ is a programming guideline that states that if you define one of the following special member functions

A copy constructor

A copy assignment operator

A destructor

then you should also define the other two.

The rationale behind the Rule of Three is that these three functions are closely related and should be implemented together to ensure that they work correctly.

For example, if you define a copy constructor but not a copy assignment operator, then the compiler will generate a default copy assignment operator that will simply copy the address of the source object to the destination object. This can lead to unexpected behavior, such as two objects sharing the same data.

The Rule of Three is not always a hard and fast rule. There are cases where it is okay to define one of these functions without defining the others.

It is generally a good idea to follow the Rule of Three to avoid potential problems.

**Here is an example of a class that follows the Rule of Three.**

```
class Rectangle {  
public:  
    int width;  
    int height;  
  
    Rectangle (int width, int height) :  
        width (width), height (height) {}  
  
    Rectangle (const Rectangle& other) :  
        width (other.width), height (other.height) {}  
  
    Rectangle& operator=(const Rectangle& other) {  
        width = other.width;  
        height = other.height;  
        return *this;  
    }  
  
    ~Rectangle() {}  
};
```

This class defines all three special member functions, a constructor, a copy constructor, and a copy assignment operator. This is necessary because the Rectangle class owns resources that need to be managed.

## Question 180 | Design Patterns and Architectural Concepts

What is the rule of five in C++ ?

### Answer

The **Rule of Five** in C++ is a guideline that states that if you define one of the following special member functions

A copy constructor

A copy assignment operator

A move constructor

A move assignment operator

A destructor

then you should also define the other four.

The rationale behind the Rule of Five is that these five functions are closely related and should be implemented together to ensure that they work correctly.

For example, if you define a copy constructor but not a copy assignment operator, then the compiler will generate a default copy assignment operator that will simply copy the address of the source object to the destination object. This can lead to unexpected behavior, such as two objects sharing the same data.

The Rule of Five is not always a hard and fast rule. There are cases where it is okay to define one of these functions without defining the others.

It is generally a good idea to follow the Rule of Five to avoid potential problems.

Here is an example of a class that violates the Rule of Five.

```
class Point {  
  
public:  
    int x;  
    int y;  
  
    Point (int x, int y) : x (x), y (y) {}  
  
    ~Point() {}  
};
```

This class defines a constructor and a destructor, but not a copy constructor, copy assignment operator, move constructor, or move assignment operator. This is okay because the Point class does not own any resources that need to be managed.

## Question 181 | Design Patterns and Architectural Concepts

### What is Copy and Swap Idiom in C++ ?

#### Answer

The **copy-and-swap** idiom lets us implement the copy assignment operator with a strong exception safety guarantee.

We finally have only a copy of the new data. So, the copy-and-swap idiom needs three things as are a **copy-constructor**, a **destructor**, and a **swap function**.

A swap function is a non-throwing function that swaps two objects of a class.

```
class Dummy {  
  
public:  
    Dummy& operator=(const Dummy& d) {  
        if (this != &d) {  
            // Copy-constructor and non-throwing swap  
            Dummy(d).swap(*this);  
        }  
        return *this;  
    }  
  
    // Non-throwing swap  
    void swap(Dummy& d) noexcept {  
        std::swap(this->str, d.str);  
    }  
  
private:  
    char* str;  
};
```

## Question 182 | Design Patterns and Architectural Concepts

What is CRTP ?

### Answer

CRTP stands for **Curiously Recurring Template Pattern**. It is an advanced C++ template technique used to achieve static polymorphism by inheriting from a template class. This pattern allows you to provide a form of polymorphism at compile time, similar to what you achieve with virtual functions and runtime polymorphism, but without the overhead of virtual function calls.

The CRTP involves creating a base class template with a derived class that inherits from the base class template and provides itself as the template argument.

Here is an example of achieving CRTP.

```
#include <iostream>

template <typename Derived>
class Base {
public:
    void interface() {
        static_cast<Derived*>(this)→implementation();
    }
};

class DerivedA : public Base<DerivedA> {
public:
    void implementation() {
        std::cout << "DerivedA implementation" << std::endl;
    }
};

class DerivedB : public Base<DerivedB> {
public:
    void implementation() {
        std::cout << "DerivedB implementation" << std::endl;
    }
};
```

```
int main() {  
    DerivedA a;  
    DerivedB b;  
  
    a.interface(); // Calls DerivedA implementation  
    b.interface(); // Calls DerivedB implementation  
  
    return 0;  
}
```

## Question 183 | Design Patterns and Architectural Concepts

**What is a virtual friend function ?**

### *Answer*

Friend functions can not be declared virtual and therefore no dynamic binding of friend functions is possible. Applying a friend function to an entire hierarchy of classes becomes awkward if an overloaded friend function is needed for every class in the hierarchy.

This lack of support for dynamic binding makes it hard to justify that friend functions are in fact an extension of the class's interface.

Virtual friend function idiom addresses this concern elegantly. Virtual friend function idiom makes use of an extra indirection to achieve the desired effect of dynamic binding for friend functions.

### *Example*

```
class Base {
public:
    friend ostream& operator << (ostream& o, const Base& b);

protected:
    virtual void print(ostream& o) const {
        // Print something...
    }
};

// Put this function into the header file
inline std::ostream& operator<<(std::ostream& o, const Base& b) {
    // Delegate the work to a polymorphic member function
    b.print(o);
    return o;
}

class Derived : public Base {
protected:
    virtual void print(ostream& o) const {
        // Print something...
    }
};
```

## Question 184 | Design Patterns and Architectural Concepts

What is policy based design ?

### Answer

**Policy-based design** is a C++ design approach that involves creating flexible and customizable classes by composing them from smaller, independent policies or components. These policies encapsulate specific behaviors, algorithms, or features, and they can be combined to create a class with various combinations of functionalities. In policy-based design, the class behavior is determined at compile time based on the policies selected.

This approach promotes code reuse, modularity, and separation of concerns.

### Example

```
#include <iostream>

// Policy Logger
template <typename Logger>
class DataManager : public Logger {
public:
    void processData() {
        Logger::log("Processing data...");
        // Perform data processing...
    }
};

struct ConsoleLogger {
    void log(const std::string& message) {
        std::cout << "Console: " << message << std::endl;
    }
};

struct FileLogger {
    void log(const std::string& message) {
        std::cout << "Writing to file... " << std::endl;
    }
};
```

```
int main() {  
  
    DataManager<ConsoleLogger> consoleDataManager;  
    consoleDataManager.processData();  
  
    DataManager<FileLogger> fileDataManager;  
    fileDataManager.processData();  
  
    return 0;  
}
```

## Question 185 | Design Patterns and Architectural Concepts

What is RAII (Resource Acquisition Is Initialization), and how is it useful in C++ ?

### ***Answer***

[RAII](#) is a C++ programming technique where the lifetime of a resource is tied to the scope of an object. When the object is created, the resource is acquired, and when the object goes out of scope (e.g., at the end of a block or when an exception is thrown, the resource is automatically released in the object's destructor.

RAII is useful because it ensures that resources are correctly and automatically managed without the need for explicit cleanup code. It helps prevent resource leaks and makes C++ code more robust and exception-safe.

# Performance Optimization

---

## Question 186 | Performance Optimization

**What is the difference between compile time and runtime optimization ?**

### ***Answer***

Compile time optimization focuses on improving code before it's executed. It includes inlining functions, loop unrolling, and constant propagation.

Runtime optimization involves techniques like just-in-time (JIT) compilation, adaptive optimization, and profiling-guided optimization.

While compile time optimizations are more predictable, runtime optimizations adapt to specific execution scenarios.

## Question 187 | Performance Optimization

**What is Small String Optimization (SSO) in C++ and how does it optimize string handling ?**

### ***Answer***

SSO is a technique used in C++ standard library `std::string` to optimize the storage of short strings. Instead of always allocating dynamic memory for strings, SSO implementations reserve a small buffer within the string object itself. Short strings are stored directly in this buffer, avoiding the need for dynamic memory allocation and deallocation.

SSO improves performance for common small strings, as it reduces memory allocation overhead.

## Example

```
#include <iostream>
#include <string>

void* operator new(std::size_t size) {
    std::cout << "Custom new called for size: "
               << size << std::endl;
    return malloc(size);
}

void operator delete(void* ptr) noexcept {
    std::cout << "Custom delete called" << std::endl;
    free(ptr);
}

int main() {

    std::string shortStr = "Hello, SS0!";
    std::string longStr = "This is longer.."; // 16 chars

    std::cout << "Short String: " << shortStr << std::endl;
    std::cout << "Long String: " << longStr << std::endl;

    return 0;
}
```

The **new** and **free** operators will be called for allocating space in the heap for longStr variable. For the shortStr variable it will not call new or free.

## Question 188 | Performance Optimization

**What is copy-on-write (COW) optimization in C++ strings, and how does it work ?**

### ***Answer***

**Copy-on-write** is a technique used in C++ string implementations to optimize memory usage.

When a copy of a string is created, the actual copying data is **deferred** until a modification is attempted. If the copied string is not modified, both the original and the copy share the same memory, reducing memory overhead.

## Question 189 | Performance Optimization

**What is Empty Base Class Optimization (EBCO) ?**

### ***Answer***

In some cases, when a class is derived from an empty base class, the size of the derived class might not include additional space for the empty base class. This is known as the Empty Base Class Optimization (EBCO), and it's an optimization technique employed by some compilers to save memory.

## Question 190 | Performance Optimization

**Explain Copy Elision in C++ and how it optimizes code.**

### ***Answer***

Copy elision is an optimization technique where the compiler eliminates unnecessary copy operations during object construction. It is often associated with Return Value Optimization (RVO) and Named Return Value Optimization (NRVO).

Copy elision improves code performance by avoiding the creation of temporary objects and redundant copy/move operations, especially in return statements.

## Question 191 | Performance Optimization

How can you optimize memory usage in C++ for containers like `std::vector` ?

### **Answer**

Optimizing memory usage in `std::vector` and other containers involves

Using `reserve()` to preallocate memory.

Using `shrink_to_fit()` to reduce excess capacity.

Avoiding excessive reallocations by estimating container size.

Using `emplace_back()` for efficient element insertion.

Minimizing unnecessary copies of elements.

## Question 192 | Performance Optimization

What is the const-correctness principle, and how can it aid code optimization ?

### **Answer**

**Const-correctness** is a coding practice that involves using the `const` qualifier to indicate when objects and functions do not modify data. It aids optimization by allowing the compiler to make assumptions about the immutability of data, enabling better optimization opportunities. It also provides safety guarantees by preventing unintentional modifications of data.

## Question 193 | Performance Optimization

What is loop unrolling, and how can it optimize code ?

### *Answer*

**Loop unrolling** is an optimization technique where the compiler or programmer manually transforms a loop with a fixed number of iterations into a series of repeated instructions. This reduces loop control overhead and can improve performance by allowing the compiler to better utilize instruction pipelines and registers.

## Question 194 | Performance Optimization

What is Loop Fusion, and how can it optimize code ?

### *Answer*

**Loop fusion** is an optimization technique that combines multiple loops into a single loop. It can optimize code by reducing loop overhead and improving **cache locality**.

Loop fusion reduces the number of loop iterations and memory accesses, resulting in better performance. However, it should be used carefully to avoid making code less readable.

## Question 195 | Performance Optimization

**Explain the concept of branch prediction and how it can affect code optimization.**

### ***Answer***

**Branch prediction** is a crucial feature in modern CPU architectures designed to improve instruction-level parallelism and execution efficiency. It primarily deals with conditional branch instructions (e.g., if statements) in program code.

These branches represent points in the code where the execution can follow one of two or more possible paths, depending on a condition.

### **Here is how Branch Prediction Works**

#### **1. Prediction**

When a branch instruction is encountered, the CPU's branch predictor makes an educated guess about which path the program is likely to take. It predicts whether the condition will be true or false based on historical behavior and the program's execution context.

#### **2. Speculative Execution**

The CPU speculatively executes instructions along the predicted path. This means it starts executing instructions before the condition's actual result is known, assuming the prediction is correct.

This speculative execution aims to keep the CPU pipeline filled with instructions and prevent stalls.

#### **3. Confirmation or Misprediction**

At some point, the actual condition result becomes available, and it may confirm or contradict the prediction. If the prediction is correct, the CPU continues executing the speculatively chosen path without interruption. If the prediction is incorrect, the CPU discards the speculative work done along the incorrect path, flushes the pipeline, and resumes execution from the correct path.

## Question 196 | Performance Optimization

What is SIMD ?

### Answer

**SIMD** (Single Instruction, Multiple Data) is a parallel computing technique that performs the same operation on multiple data elements simultaneously.

It's an essential concept in modern CPU architectures and programming, particularly in areas such as multimedia processing, scientific computing, and data intensive applications.

#### 1. Single Instruction

SIMD instructions enable a processor to execute a single instruction on multiple data elements in parallel. This instruction is applied to all the data elements in a single clock cycle.

#### 2. Multiple Data

SIMD operates on data in a vectorized or packed format. Instead of processing individual data elements one at a time, SIMD units can process multiple data elements together in a single operation.

## Question 197 | Performance Optimization

What is loop tiling (loop blocking) and how does it enhance cache efficiency ?

### Answer

**Loop tiling**, also known as **loop blocking**, is an optimization technique that divides a large loop into smaller, nested loops. These smaller loops operate on chunks of data that fit into the CPU's cache. This enhances cache efficiency by reducing cache misses and improving data locality.

## Example

```
// Size of the matrices
const int N = 1000;
// Block size as per cache size
const int blockSize = 4;
float A[N][N], B[N][N], C[N][N];

for (int i = 0; i < N; i += blockSize) {
    for (int j = 0; j < N; j += blockSize) {
        for (int k = 0; k < N; k += blockSize) {
            // Multiply the tiny matrix
            for (int i2 = i; i2 < min(i + blockSize, N); ++i2)
                for (int j2 = j; j2 < min(j + blockSize, N); ++j2)
                    for (int k2 = k; k2 < min(k + blockSize, N); ++k2)
                        C[i2][j2] += A[i2][k2] * B[k2][j2];
        }
    }
}
```

In this example, we're multiplying two matrices **A** and **B** to get a result matrix **C**. Instead of the naive approach of three nested loops each running from 0 to N, we are breaking the computation down into smaller blocks of size **blockSize**.

This ensures that the subset of the matrices we're working with in each iteration of the outer loops fits into the cache, reducing cache misses and improving performance.

## Question 198 | Performance Optimization

**Explain the concept of data prefetching in optimization. How does it improve memory access patterns ?**

### Answer

Data prefetching is a technique where the processor predicts which memory data will be needed in the future and fetches it into a cache before it's actually used. This reduces memory latency and improves memory access patterns, enhancing overall performance. It's particularly useful for loops with regular memory access patterns.

## Question 199 | Performance Optimization

What is `__restrict` in C++ ?

### Answer

Restrict says that two pointers cannot point to overlapping memory regions. The most common usage is for function arguments. This restricts how the function can be called, but allows for more compiler optimizations. If the caller does not follow the restrict contract, undefined behavior can occur.

The primary purpose of the restrict keyword is to enable optimizations that may not be possible without it. It tells the compiler that, within a given scope, the pointer declared as restrict will not be accessed through any other pointer that may point to the same memory location.

This allows the compiler to make optimizations like reordering memory operations and prefetching data.

```
void compute (int* __restrict__ a,
              int* __restrict__ b, int size) {
    for (int i = 0; i < size; ++i) {
        a[i] = b[i] * 2;
    }
}
```

## Question 200 | Performance Optimization

What is the strict aliasing rule in C++ ?

### Answer

The **strict aliasing rule** is a set of rules defined by the C++ standard that dictate how pointers of different types can be used to access objects in memory.

The rule is intended to allow the compiler to make certain optimizations by assuming that pointers of different types do not alias (point to the same memory location) unless there is explicit type-based aliasing.

## Question 201 | Performance Optimization

What is narrowing type punning, and when is it typically used ?

### Answer

Narrowing type punning is a type of type punning where you access an object through a pointer of a smaller type, like a `char*`.

It's primarily used for tasks like serialization, memory manipulation, and bit manipulation. Accessing an object through a smaller type, like a `char*`, is generally safe and is allowed by the strict aliasing rule.

## Question 202 | Performance Optimization

What is the use of `std::memcpy` in C++ ?

### Answer

`std::memcpy` is a C++ standard library function that is used for copying a block of memory from one location to another.

### Memory Copying

`std::memcpy` is designed to efficiently copy a block of memory from a source location to a destination location. It's a byte-by-byte copy operation, which means it doesn't interpret the data being copied; it just moves raw bytes.

```
void* memcpy(void* dest, const void* src, size_t count);
```

**dest** - A pointer to the destination location where the data will be copied.

**src** - A pointer to the source location from which data will be copied.

**count** - The number of bytes to copy.

# Modern C++ Standards C++17, C++20

---

## Question 203 | Modern C++ Standards C++17, C++20

What is `std::optional` and what is its usage ?

### ***Answer***

`std::optional` is a C++17 standard library class template that provides a way to represent an optional value, meaning a value that may or may not be present. It's designed to handle scenarios where a value may be absent without resorting to using null pointers or sentinel values.

`std::optional` is particularly useful when dealing with functions that might return a value or nothing, avoiding the need for error-prone error codes or null pointers.

## Question 204 | Modern C++ Standards C++17, C++20

What is a fold expression in C++17 ?

### Answer

A fold expression is a feature introduced in C++17 that simplifies the application of binary operators to a sequence of template parameters. It allows you to **fold** the values in the sequence using a binary operator in a concise and readable way.

For example, you can use fold expressions to compute the sum or product of a parameter pack.

### Example

```
#include <iostream>

template<typename... Args>
auto sum(Args... args) {
    return (args + ...);
}

int main() {
    int result = sum(1, 2, 3, 4);
    std::cout << "Result: " << result << std::endl;
    return 0;
}
```

And the output will be **10** as a sum of values.

## Question 205 | Modern C++ Standards C++17, C++20

What is structured binding in C++17 ?

### Answer

**Structured binding** is a feature introduced in C++17 that allows you to unpack the members of a class or tuple-like object directly into variables. It provides a more concise way to access the elements of a complex data structure without explicitly specifying each member.

### Example

```
#include <tuple>
#include <string>
#include <iostream>

int main() {
    std::tuple<int, double, std::string> data =
        std::make_tuple(42, 3.14, "hello");
    auto [num, value, text] = data; // Structured binding
    std::cout << "num: " << num << ", value: "
        << value << ", text: " << text << std::endl;
    return 0;
}
```

## Question 206 | Modern C++ Standards C++17, C++20

**What is `std::variant` in C++17 and how is it used ?**

`std::variant` is a C++17 standard library class template that represents a type-safe union. It can store a value of one of its alternative types, and the active type is explicitly specified at runtime. Unlike `std::tuple` or `std::pair`, which store values of all types simultaneously, `std::variant` holds a value of

### Multiple Types

`std::variant` is useful when you have a set of possible types, and you want to store an instance of one of these types.

### Type Safe Unions

It provides type safety, as only one type is active at a time.

### Visitors

When you need to perform different operations based on the type stored in the variant, you can use visitors to perform type-specific actions of only one type.

## Question 207 | Modern C++ Standards C++17, C++20

What is the C++17 `std::string_view` and how does it differ from `std::string` ?

### *Answer*

`std::string_view` is a lightweight non-owning view of a string's character data. It allows you to work with strings without copying the data.

Unlike `std::string`, `std::string_view` doesn't own the underlying memory, making it more suitable for cases where you want to avoid unnecessary memory allocations and copies.

## Question 208 | Modern C++ Standards C++17, C++20

What is Class Template Argument Deduction introduced in C++17 ?

### *Answer*

Class Template Argument Deduction is a C++17 feature that allows you to omit template arguments when creating instances of class templates if the compiler can deduce the arguments from the constructor arguments. This feature reduces the verbosity of template instantiation and makes code more concise.

## Question 209 | Modern C++ Standards C++17, C++20

### What is `std::clamp` ?

`std::clamp` is a C++ Standard Library function introduced in C++17. It is used to constrain or clamp a value within a specified range. This function ensures that a given value remains within the specified lower and upper bounds.

Here's what each parameter represents.

```
template <typename T>
constexpr const T& clamp(const T& value,
                        const T& lower, const T& upper);
```

value - The value you want to clamp or constrain.

lower - The lower bound of the allowed range.

upper - The upper bound of the allowed range.

The `std::clamp` function returns a reference to the clamped value.

### Example

```
#include <iostream>
#include <algorithm>

int main() {
    int value = 42;
    int lower = 22;
    int upper = 43;

    // Clamp the value within the specified range
    int clampedValue = std::clamp(value, lower, upper);

    std::cout << "Original value: " << value << std::endl;
    std::cout << "Clamped value: " << clampedValue
              << std::endl;
    return 0;
}
```

## Question 210 | Modern C++ Standards C++17, C++20

What is alternative operator representation in C++ ?

### Answer

The naming **and** is used as an alternative to **&&** and **or** as an alternative to **||** for logical operations. These alternatives are provided to make the code more readable and can be particularly useful in contexts where you want to express logical conditions more naturally.

```
bool cdt1 = true;
bool cdt2 = false;

if(cdt1 and cdt2) {
    std::cout << "Both conditions are true." << std::endl;
} else {
    std::cout << "At least one condition is false."
               << std::endl;
}
```

## Question 211 | Modern C++ Standards C++17, C++20

What is the execution policy in C++17 ?

### Answer

The concept of **execution policies** was introduced to specify the level of parallelism or execution constraints when using algorithms from the C++ Standard Library. Execution policies allow you to control how standard algorithms, like `std::for_each`, `std::transform`, and `std::reduce`, perform their operations.

The primary purpose is to enable parallelism in a controlled and standardized way.

There are three execution policies defined in C++17

1. `std::execution::seq` (Sequenced Execution)

This policy specifies that the algorithm must be executed sequentially, meaning that it should not use parallelism. It provides a guarantee that the algorithm will not run concurrently, making it suitable for algorithms that may have data dependencies or must maintain a specific order of execution.

2. `std::execution::par` (Parallel Execution)

This policy allows the algorithm to execute in parallel, taking advantage of multiple threads or processing units if available. It does not guarantee parallel execution, but it enables it when possible. You can think of it as a hint to the compiler to parallelize the operation when it makes sense.

3. `std::execution::par_unseq` (Parallel Unsequenced Execution)

This policy enables parallel execution and allows for vectorization or other optimizations that may not strictly maintain the sequential order of execution.

It provides more freedom to the compiler to optimize the operation aggressively.

*Example*

```
std::for_each(std::execution::par, data.begin(),
              data.end(), [](int& value) {
    // Process each element in parallel
});
```

# Computer Architecture and Advanced Memory Management

---

## Question 212 | Computer Architecture and Advanced MM

What is memory fragmentation ?

### ***Answer***

Memory fragmentation refers to the phenomenon where the memory space in a computer system becomes divided into small, non-contiguous chunks, making it challenging to allocate large blocks of contiguous memory even when the total available memory is sufficient. Memory fragmentation can occur in both physical (RAM) and virtual (address space) memory.

### **External Fragmentation**

External fragmentation occurs when free memory blocks are scattered throughout the memory space, leaving gaps between them. These gaps are too small to satisfy large memory allocation requests, even if the total free memory is sufficient. This type of fragmentation is more common in systems with dynamic memory allocation and deallocation, such as those using heap memory management.

### **Internal Fragmentation**

Internal fragmentation happens when allocated memory blocks are larger than the requested data size. In other words, it is the wasted space within allocated memory blocks. This type of fragmentation is more common in systems with fixed size memory allocation, such as those using stack memory management or block based memory allocation.

## Question 213 | Computer Architecture and Advanced MM

What is virtual address space ?

### **Answer**

**Virtual address space** is a fundamental concept in modern computer systems, especially those with memory management units (MMUs) like in most contemporary operating systems. It is a memory addressing scheme that abstracts the underlying physical memory and provides several benefits for program execution and system stability.

Here are key points about virtual address space

#### **1. Abstraction of Physical Memory**

Virtual address space abstracts the physical memory (RAM) available in a computer system. Instead of directly interacting with physical memory addresses, programs work with virtual addresses.

#### **2. Isolation**

Each process running on the system has its own private virtual address space. This isolation prevents one process from accessing or modifying the memory of another process, enhancing system security and stability.

#### **3. Efficient Utilization**

Virtual address space allows efficient memory utilization. A process can allocate more memory than physically available (known as overcommitting), relying on techniques like demand paging and swapping to bring data in and out of physical memory as needed.

#### **4. Memory Protection**

Virtual memory systems offer memory protection mechanisms. Unauthorized access to memory regions can result in exceptions or segmentation faults, preventing rogue processes from compromising system integrity.

## Question 214 | Computer Architecture and Advanced MM

What is the difference between physical and virtual addresses ?

### **Answer**

#### 1. Abstraction

Virtual addresses abstract the underlying physical memory. They provide processes with a consistent and isolated view of memory, allowing each process to use addresses as if it were the only process running.

#### 2. Isolation

Virtual addresses enable process isolation. Each process has its own virtual address space, preventing unauthorized access to or interference with the memory of other processes.

#### 3. Flexibility

Virtual addresses allow for flexible memory management. Memory can be allocated, deallocated, and moved within a process's virtual address space without requiring changes to physical memory locations.

## Question 215 | Computer Architecture and Advanced MM

What is ASLR ?

### **Answer**

#### Address Space Layout Randomization (ASLR)

ASLR is a security feature that randomizes the base address of a process's executable and libraries within its virtual address space. This makes it harder for attackers to predict memory locations for exploits.

## Example

```

#include <iostream>

int globalVariable;

int main() {

    int localVariable;
    int* heapVariable = new int;

    std::cout << "Address of main: " << (void*)&main << '\n';
    std::cout << "Address of global variable: "
                << &globalVariable << '\n';
    std::cout << "Address of local variable: "
                << &localVariable << '\n';
    std::cout << "Address of heap variable: "
                << heapVariable << '\n';

    delete heapVariable;

    return 0;
}

```

This program prints out the addresses of a function (main), a global variable, a local variable (in stack), and a dynamically allocated variable (in heap).

If we run this program multiple times on a system with ASLR enabled, we will see that the addresses of these items change between runs, demonstrating the effect of ASLR.

## Question 216 | Computer Architecture and Advanced MM

**What is cache coherency ?**

### Answer

**Cache coherency** is a concept in computer architecture, particularly in multiprocessor systems, where multiple CPU cores or processors share access to the same main memory (RAM) through separate caches.

Cache coherency mechanisms ensure that all processors observe a consistent and up-to-date view of memory, even though they may have their own local caches.

## Question 217 | Computer Architecture and Advanced MM

How many cache coherency protocols are there, list some.

### Answer

There are different types of cache coherency protocols, each with its own design and mechanisms to maintain cache consistency in multiprocessor systems.

Some of the well known cache coherency protocols include

1. **MESI Protocol** (Modified, Exclusive, Shared, Invalid)

MESI is one of the most commonly used cache coherency protocols. It defines four states for each cache line: Modified (M), Exclusive (E), Shared (S), and Invalid (I).

2. **MOESI Protocol** (Modified, Owner, Exclusive, Shared, Invalid)

MOESI is an extension of the MESI protocol and includes an additional **Owned** state, which helps in managing caches in systems with more than two processors.

3. **MESIF Protocol** (Modified, Exclusive, Shared, Invalid, Forward)

MESIF is an extension of the MESI protocol that includes a **Forward** state. This state allows a processor to forward modified data to another cache without writing it back to memory.

## Question 218 | Computer Architecture and Advanced MM

What is the MOESI protocol, and how does it differ from MESI ?

### Answer

MOESI stands for Modified, Owned, Exclusive, Shared, and Invalid. It is an extension of the MESI protocol, with an additional **Owned** state. In the MOESI protocol, when a cache has data exclusively, it also indicates that it's the owner of that data.

This allows for optimizations when another cache wants to acquire exclusive ownership.

## Question 219 | Computer Architecture and Advanced MM

What is the purpose of the Modified state in MESI and MOESI protocols ?

### **Answer**

The **Modified** state indicates that the cache line has been modified compared to the main memory, and it is only present in the cache. Any write operation to this line requires an update to the main memory, making it more up-to-date than the **Exclusive** state.

## Question 220 | Computer Architecture and Advanced MM

How does the Exclusive state differ from the Modified state in MESI and MOESI ?

### **Answer**

The **Exclusive** state indicates that the cache line is present exclusively in a single cache, but it matches the main memory. It is read only in this cache.

In contrast, the **Modified** state indicates that the cache line is also exclusive but has been modified and needs to be written back to main memory.

## Question 221 | Computer Architecture and Advanced MM

Why is cache coherency important in multiprocessor systems ?

### **Answer**

Cache coherency ensures that all processors in a multiprocessor system have a consistent view of memory. It prevents issues like data corruption or race conditions that can arise when multiple processors have different copies of the same memory location.

## Question 222 | Computer Architecture and Advanced MM

**When does a cache line transition from Shared to Invalid in the MESI protocol ?**

### ***Answer***

In the MESI protocol, a cache line transitions from **Shared** to **Invalid** when a processor writes to that cache line. This transition ensures that other caches are aware that the data is no longer valid.

## Question 223 | Computer Architecture and Advanced MM

**What is a cache line or cache block ?**

### ***Answer***

A cache line, also known as a cache block, is the smallest unit of data that can be stored in a cache. When data is fetched from memory into a cache, it is loaded into one or more cache lines.

## Question 224 | Computer Architecture and Advanced MM

**Explain the concept of cache hit and cache miss.**

### ***Answer***

A cache hit occurs when the CPU requests data that is already present in the cache, resulting in a faster data access. A cache miss occurs when the requested data is not in the cache, requiring the system to fetch it from a slower memory level.

## Question 225 | Computer Architecture and Advanced MM

What is cache associativity, and how does it affect cache performance ?

### **Answer**

Cache associativity refers to how cache lines are mapped to specific locations in the cache. It affects cache performance by determining how easily data can be replaced or accessed when the cache is full.

Common associativity levels include **direct mapped**, **set associative**, and **fully associative** caches.

## Question 226 | Computer Architecture and Advanced MM

What is a write-through cache and a write-back cache ?

### **Answer**

In a **write-through** cache, data is written to both the cache and the main memory simultaneously.

In a **write-back** cache, data is initially written only to the cache, and it is written to the main memory later when the cache line is replaced.

## Question 227 | Computer Architecture and Advanced MM

How does cache prefetching improve cache performance ?

### **Answer**

**Cache prefetching** is a technique where the CPU predicts what data it will need next and fetches it into the cache ahead of time. This can reduce cache misses and improve overall performance by minimizing data access latency.

## Question 228 | Computer Architecture and Advanced MM

What is cache eviction, and how is it managed in a cache system ?

### ***Answer***

**Cache eviction** is the process of selecting which cache lines to replace when new data needs to be loaded into the cache. Various cache replacement policies, such as **LRU** (Least Recently Used) or **FIFO** (First-In, First-Out), are used to manage cache eviction.

## Question 229 | Computer Architecture and Advanced MM

What is cache line size, and how does it affect cache performance ?

### ***Answer***

Cache line size refers to the number of bytes stored in a cache line. A larger cache line size can reduce the number of cache misses for data accessed sequentially but may lead to wasted memory if only a portion of the line is used.

## Question 230 | Computer Architecture and Advanced MM

### What are Direct-mapping, Associative Mapping & Set-Associative Mappings?

#### Answer

Direct-mapping, Associative Mapping, and Set-Associative Mapping are three different cache mapping techniques used in computer architecture.

#### Direct-Mapping

In a **direct-mapped** cache, each block of main memory can be mapped to exactly one specific cache line. This mapping is achieved by using a modulus function, which means that a block in main memory is placed in a specific cache line based on the block's address modulo the number of cache lines. Direct-mapped caches are simple and efficient for small caches but can suffer from conflicts where multiple blocks in main memory map to the same cache line.

#### Associative Mapping

In an **associative-mapped** cache, any block of main memory can be placed in any cache line. There is no strict one-to-one mapping, and the cache controller can place a block in any available cache line. Associative caches are highly flexible and can reduce cache conflicts, but they are more complex to implement than direct-mapped caches.

#### Set-Associative Mapping

**Set-associative** mapping is a compromise between direct-mapped and fully associative caches.

The cache is divided into a set of groups or **ways** and each way contains multiple cache lines. A block from main memory can be placed in any cache line within a specific set.

The number of ways in a set-associative cache is referred to as the associativity level (e.g., 2-way, 4-way, 8-way). Set-associative caches aim to balance the simplicity of direct-mapped caches with the flexibility of fully associative caches, reducing conflicts while keeping some structure.

## Question 231 | Computer Architecture and Advanced MM

### What is False Sharing ?

#### Answer

**False Sharing** is a performance problem that can occur in multithreaded or multiprocessor systems, especially in the context of shared memory parallelism. It happens when multiple threads or processors unintentionally access and modify data that resides on the same cache line, even though they are working on different variables within that line.

This can lead to unnecessary cache synchronization and contention, causing a significant performance bottleneck.

#### Example

```
#include <thread>
#include <iostream>
#include <chrono>
#include <new>

#ifdef __cpp_lib_hardware_interference_size
    using std::hardware_constructive_interference_size;
    using std::hardware_destructive_interference_size;
#else
    // 64 bytes on x86-64 | L1_CACHE_BYTES | L1_CACHE_SHIFT |
    // __cacheline_aligned
    constexpr std::size_t hardware_constructive_interference_size = 64;
    constexpr std::size_t hardware_destructive_interference_size = 64;
#endif

struct alignas(hardware_destructive_interference_size) AlignedInt {
    int value = 0;
};

/*
 * Uncomment and use unaligned struct below,
 * for comparison / time testing
 */
//struct NonAlignedInt {
//    int value = 0;
//};

AlignedInt counters[2];
```

```

void increment(int index, int num_iterations) {
    for (int i = 0; i < num_iterations; ++i) {
        ++counters[index].value;
    }
}

int main() {
    const int num_iterations = 100000000;

    auto start = std::chrono::high_resolution_clock::now();

    std::thread t1(increment, 0, num_iterations);
    std::thread t2(increment, 1, num_iterations);

    t1.join();
    t2.join();

    auto end = std::chrono::high_resolution_clock::now();
    std::chrono::duration<double> diff = end - start;

    std::cout << "Time taken = " << diff.count() << " seconds\n";
    std::cout << "Counters = " << counters[0].value
                << ", " << counters[1].value << "\n";

    return 0;
}

```

## Question 232 | Computer Architecture and Advanced MM

What is cache contention, and how does it relate to false sharing ?

### Answer

**Cache contention** refers to the competition for access to a shared cache resource among multiple threads or processors. False sharing exacerbates cache contention by causing threads to contend for the same cache line, leading to unnecessary cache line invalidation and synchronization. It can result in performance bottlenecks.

## Question 233 | Computer Architecture and Advanced MM

**Explain cache line alignment and its importance in optimizing cache usage.**

### ***Answer***

**Cache line alignment** involves arranging data structures and memory accesses in a way that aligns with the cache line size. Proper alignment minimizes cache wastage, reduces false sharing, and improves cache utilization efficiency.

## Question 234 | Computer Architecture and Advanced MM

**Explain cache coloring and how it can be used to mitigate cache coherency issues.**

### ***Answer***

**Cache coloring** is a technique used to control the placement of data in cache sets.

By assigning specific colors to memory regions, developers can influence how data is distributed across cache sets.

Properly coloring data can reduce cache contention and mitigate cache coherency issues.

## Question 235 | Computer Architecture and Advanced MM

What is cache thrashing ?

### ***Answer***

**Cache thrashing** happens when a cache repeatedly loads and evicts data, and as a result, no thread or process can efficiently use the cache. This typically occurs when multiple threads are competing for access to the same cache lines, and these lines are being invalidated frequently.

For example when several threads trying to access the same memory locations, but the cache is not large enough to accommodate all the data they need simultaneously.

As a result, cache lines are constantly being loaded and evicted. This constant churning of cache data slows down the execution of threads because they spend more time waiting for data to be loaded into the cache than actually processing it.

## Question 236 | Computer Architecture and Advanced MM

What is pipelining in computer architecture ?

### ***Answer***

**Pipelining** is a technique used in computer architecture to improve CPU performance by breaking down the execution of instructions into several stages, each performed by a different segment of the CPU pipeline.

This allows multiple instructions to be in various stages of execution simultaneously, resulting in faster and more efficient processing.

## Question 237 | Computer Architecture and Advanced MM

**Explain the concept of out-of-order execution.**

### ***Answer***

**Out-of-order execution** is a CPU design technique where the processor executes instructions not in the order they appear in the program but based on the availability of input data and execution resources. This technique aims to maximize CPU utilization and performance by executing independent instructions as soon as possible, even if they are out of order from the program's original sequence.

## Question 238 | Computer Architecture and Advanced MM

**What is the purpose of a hardware interrupt in a CPU, and how does it differ from a software interrupt ?**

### ***Answer***

A hardware interrupt is a signal sent by external hardware devices to request the CPU's attention. When a hardware interrupt occurs, the CPU stops executing its current program and switches to an interrupt handler routine.

In contrast, a software interrupt is triggered by software instructions and is used to request specific services from the operating system.

## Question 239 | Computer Architecture and Advanced MM

**Explain the advantages and disadvantages of symmetric multiprocessing (SMP) and asymmetric multiprocessing (AMP) in multi-core processors.**

### ***Answer***

**SMP** provides equal access to all CPU cores, making it suitable for mission critical tasks that require load balancing. **AMP** dedicates specific cores to specific tasks, which can enhance determinism and isolation.

## Question 240 | Computer Architecture and Advanced MM

**What is the role of Instruction-Level Parallelism (ILP) in processor design and how can ILP be leveraged for mission critical applications ?**

### ***Answer***

**ILP** allows multiple instructions to be executed concurrently within a single processor core, which can lead to significant performance gains.

For mission critical applications, compilers and processors can exploit ILP to achieve high throughput and responsiveness.

## Question 241 | Computer Architecture and Advanced MM

**What is NUMA architecture ?**

### **Answer**

**NUMA** (Non-Uniform Memory Access) architecture is a type of shared memory multiprocessing architecture used in multiprocessor systems. In NUMA systems, the physical memory is divided into multiple memory nodes, and each node is associated with one or more processors.

### **The key characteristics of NUMA architecture**

#### **Memory Nodes**

The system's memory is partitioned into distinct memory nodes, with each node connected to a subset of CPU cores. Each memory node includes a portion of the total RAM, and it is associated with the CPUs that can access it most efficiently.

#### **Local and Remote Memory Access**

In NUMA systems, memory access times are not uniform across all nodes. Accessing memory from the local node is faster and has lower latency compared to accessing memory from remote nodes. Remote memory access typically incurs higher latency.

#### **Interconnect**

NUMA systems use a high speed interconnect, often referred to as the NUMA interconnect or system interconnect, to enable communication between nodes. This interconnect ensures that processors can access memory from remote nodes when necessary.

#### **Scalability**

NUMA architectures are highly scalable and can accommodate a large number of processors. This makes them suitable for servers, data centers, and high performance computing environments.

#### **Load Balancing**

NUMA systems often employ load balancing mechanisms to optimize performance. These mechanisms aim to distribute processes and threads across processors and memory nodes efficiently to minimize remote memory access.

## Cache Coherency

NUMA architectures implement cache coherency protocols to maintain data consistency between the local caches of different processors.

These protocols ensure that changes made by one processor to shared data are visible to other processors.

NUMA architectures are designed to address the limitations of traditional UMA (Uniform Memory Access) architectures, where all processors have equal and uniform access to memory.

In NUMA systems, the goal is to reduce memory contention and improve memory access efficiency by emphasizing local memory access.

## Question 242 | Computer Architecture and Advanced MM

Can you explain the concept of memory affinity in NUMA systems ?

### **Answer**

**Memory affinity** in NUMA systems refers to the tendency of a processor to preferentially access memory from its local memory node. When a thread or process exhibits memory affinity, it tries to allocate and access memory that is physically closer to the processor, minimizing remote memory access. This behavior is essential for optimizing memory access patterns in NUMA architectures.

# Operating System

---

## Question 243 | Operating System

What are page tables ?

### ***Answer***

**Page tables** are data structures used in computer operating systems and memory management units (MMUs) to manage the mapping between **virtual addresses** and **physical addresses** in a computer's memory system.

They play a crucial role in the implementation of virtual memory, which allows processes to use more memory than is physically available and provides memory protection and isolation.

## Question 244 | Operating System

Explain how page tables work.

### ***Answer***

#### **Virtual Address Space**

Each process in a multitasking operating system operates within its own virtual address space. This address space is divided into fixed-size blocks called pages.

A virtual address consists of two parts: **a page number** and an **offset** within the page.

#### **Physical Memory**

The computer's physical memory (RAM) is also divided into fixed-size blocks, corresponding to the pages in the virtual address space.

These physical memory blocks are called page frames or simply frames.

### **Mapping**

Page tables store the mapping between virtual pages and physical frames. Each entry in the page table contains the correspondence between a virtual page number and a physical frame number. When a program accesses a virtual address, the page table is consulted to determine the corresponding physical address.

### **Translation**

When a program accesses a virtual address, the hardware's Memory Management Unit (MMU) translates the virtual address into a physical address by looking up the page table. The offset within the page remains the same, while the page number is replaced with the corresponding physical frame number.

Page tables are hierarchical data structures due to their size. In modern computer systems with large virtual address spaces, maintaining a single, flat page table would be impractical. Instead, they are typically organized as a multilevel hierarchy.

Common levels include.

#### **Page Directory**

The top level page table, which maps a portion of the virtual address space to a page table.

#### **Page Table**

Intermediate level tables that map a smaller portion of the address space to frames.

#### **Page Table Entry**

Entries in the page tables that map individual pages to frames.

## Question 245 | Operating System

**What is TLB ?**

**Answer**

**TLB** (Translation Lookaside Buffer) is a small, high speed cache used in computer processors to store frequently accessed virtual-to-physical memory address translations.

It accelerates the translation of virtual addresses used by programs to physical addresses in RAM, reducing memory access latency and improving overall system performance.

When a program accesses memory, the TLB is checked first for the translation; if it is found (a TLB hit), the memory access is faster. If the translation isn't in the TLB (a TLB miss), the system fetches it from the main memory.

The TLB helps manage the efficient use of virtual memory and is an integral part of modern computer architectures.

## Question 246 | Operating System

**What is the difference between a one-level and a two-level Page Table, and what are the advantages of using a two-level Page Table ?**

**Answer**

**One-Level Page Table**

In a one-level Page Table, a single table contains mappings for the entire virtual address space. This approach can be inefficient for systems with large address spaces because it requires a substantial amount of memory for the Page Table.

**Two-Level Page Table**

In a two-level Page Table, the Page Table is divided into two levels. The first level (also called the Page Directory) points to second-level tables. This two-level structure allows for more efficient use of memory by allocating Page Table entries only when needed, saving memory in sparse address spaces.

## Question 247 | Operating System

**Discuss the concept of TLB miss and the steps taken by the CPU when a TLB miss occurs.**

### ***Answer***

A **TLB miss** occurs when the CPU cannot find a **virtual-to-physical** address translation in the TLB. When this happens, the CPU must retrieve the translation from the Page Table, which involves the following steps.

The CPU generates a page table lookup request using the virtual address.  
The memory management unit (MMU) searches the Page Table to find the translation.  
If the translation is found, it is loaded into the TLB for future use.

The CPU retries the memory access with the corrected translation.

## Question 248 | Operating System

**What are Huge Pages in Linux, and when might they be advantageous ?**

### ***Answer***

**Huge Pages** in Linux are larger sized memory pages compared to the standard page size. They can be beneficial in scenarios where a process requires large contiguous blocks of memory. Huge Pages reduce the overhead associated with managing a large number of smaller pages, such as TLB entries and Page Table entries.

They are often used in database systems and high performance computing applications.

## Question 249 | Operating System

**Can you explain the difference between the TLB and the Page Cache in Linux ?**

### ***Answer***

**TLB (Translation Lookaside Buffer)**

The **TLB** is a hardware cache that stores virtual-to-physical address translations to speed up memory access for the CPU. It is small in size and primarily used for reducing translation latency.

**Page Cache**

The **Page Cache** is a software managed cache in the Linux kernel use to store copies of file data and metadata from disk. It helps speed up file I/O operations by reducing the need to access data directly from disk.

## Question 250 | Operating System

**Explain TLB shutdown and its significance in multi-processor or multi-core systems.**

### ***Answer***

**TLB shutdown** is a mechanism used in multi-processor or multi-core systems to ensure the consistency of TLB entries across all cores. When one core changes a page table entry (e.g. due to memory mapping changes or page removal), it notifies other cores to invalidate or update their corresponding TLB entries.

## Question 251 | Operating System

**Discuss the benefits and challenges of using Transparent Huge Pages (THP) in Linux memory management.**

### ***Answer***

#### **Benefits**

Transparent Huge Pages combine smaller memory pages into larger, more efficient pages, reducing memory overhead and TLB pressure. This can improve performance in memory-intensive applications.

#### **Challenges**

THP can introduce memory fragmentation and may not be suitable for all workloads. It may also impact memory usage visibility and make it challenging to determine memory allocations accurately.

## Question 252 | Operating System

**What is TLB preloading, and how does it help improve memory access performance in certain scenarios ?**

### ***Answer***

**TLB preloading** involves prefetching TLB entries before they are actually needed by the CPU.

This technique can be beneficial in scenarios where memory access patterns are predictable and repetitive, as it reduces TLB miss latency by ensuring that required translations are already in the TLB when needed.

## Question 253 | Operating System

**How does Linux implement demand paging, and what role do Page Tables play in this memory management strategy ?**

### ***Answer***

**Demand paging** in Linux involves loading pages into physical memory only when they are accessed.

Page Tables help implement demand paging by maintaining mappings between virtual addresses and physical frames. When a page is accessed but not yet in memory, the Page Table indicates its location on secondary storage, and the kernel fetches it as needed.

# Compiler and Linker Internals

---

## Question 254 | Compiler and Linker Internals

Explain the compilation process in C++ and the roles of the preprocessor, compiler, assembler, and linker.

### *Answer*

The compilation process involves several stages

1. **Preprocessor**

The preprocessor handles directives like `#include` and `#define`, and perform text substitution. It generates a preprocessed source file.

2. **Compiler**

The compiler translates the preprocessed source code into assembly code.

3. **Assembler**

The assembler converts assembly code into machine code (object code).

4. **Linker**

The linker combines object files, resolves external references, and generates an executable program.

## Question 255 | Compiler and Linker Internals

**What is the purpose of symbol resolution in the linking phase, and how does the linker accomplish it ?**

### ***Answer***

**Symbol resolution** involves matching references to external symbols in one object file with definitions in other object files or libraries.

The linker accomplishes this by creating a symbol table and resolving references by linking to the appropriate symbol definition, either within the same program or in external libraries.

## Question 256 | Compiler and Linker Internals

**Explain the differences between static linking and dynamic linking. What are the advantages and disadvantages of each ?**

### ***Answer***

#### **Static Linking**

In static linking, all library code used by a program is copied into the final executable. Advantages include ease of distribution and no external dependencies.

However, it can lead to larger executable sizes and increased memory usage.

#### **Dynamic Linking**

In dynamic linking, a program uses shared libraries (.dll, .so) at runtime. Advantages include smaller executables and the ability to update libraries without recompiling the program.

Disadvantages include potential versioning issues and runtime dependency management.

## Question 257 | Compiler and Linker Internals

What is name mangling in C++, and why is it used by the compiler ?

### **Answer**

**Name mangling** is a technique used by the compiler to encode information about a function's arguments and return type into its name. It helps distinguish between overloaded functions and allows for proper linkage when functions have the same name but different signatures.

## Question 258 | Compiler and Linker Internals

What are the common techniques used by the compiler to optimize C++ code during the compilation process ?

### **Answer**

Common optimization techniques employed by the compiler include

#### **Constant Folding**

Evaluating constant expressions at compile time.

#### **Dead Code Elimination**

Removing code that cannot be executed.

#### **Inlining**

Replacing function calls with the actual code to eliminate call overhead.

#### **Loop Optimization**

Transforming loops for better performance.

#### **Register Allocation**

Efficiently utilizing CPU registers for variables.

#### **Code Scheduling**

Reordering instructions to minimize pipeline stalls.

## Question 259 | Compiler and Linker Internals

**Explain the concept of relocation in the context of linking, and how the linker resolves relocations.**

### **Answer**

**Relocation** is the process of adjusting memory addresses or offsets in compiled code to account for the actual location of symbols in memory. The linker resolves relocations by calculating the correct addresses for references to external symbols and updating the code or data accordingly during the linking process.

## Question 260 | Compiler and Linker Internals

**What is profile-guided optimization (PGO) in G++ ?**

### **Answer**

Profile-guided optimization (PGO) is a technique where the compiler uses profiling data collected during program execution to make informed optimization decisions. g++ supports PGO using flags `-fprofile-generate` and `-fprofile-use`.

## Question 261 | Compiler and Linker Internals

**What is the One Translation Unit (OTU) Rule, and how does it relate to C++ code compilation and linking process ?**

### **Answer**

The **One Translation Unit Rule** states that a variable or function with external linkage must have a single definition within the entire program.

Violating this rule can result in linker errors.

The rule emphasizes the importance of declaring variables and functions in header files and defining them in a single source file to avoid multiple definitions.

## Credits

### > **Photographic Credits: Solar eclipse progression by Kalan**

> [https://commons.wikimedia.org/wiki/File:2008-08-01\\_Solar\\_eclipse\\_progression.jpg](https://commons.wikimedia.org/wiki/File:2008-08-01_Solar_eclipse_progression.jpg)

We Value Your Feedback! 🌟

Dear Reader,

Thank you for choosing to read this book.

Your support means the world to us!

If you found this book helpful in your interview preparation or if it sparked your interest in new topics, we would love to hear from you.

Please consider leaving a review on Amazon.

Your feedback not only helps us improve but also assists other readers in finding valuable resources.

Your honest review will make a significant difference.

Thank you for your time and support!

(Amazon Link | <https://www.amazon.com/dp/B0DV6LB2RK>)

Warm regards,

Gleamixie