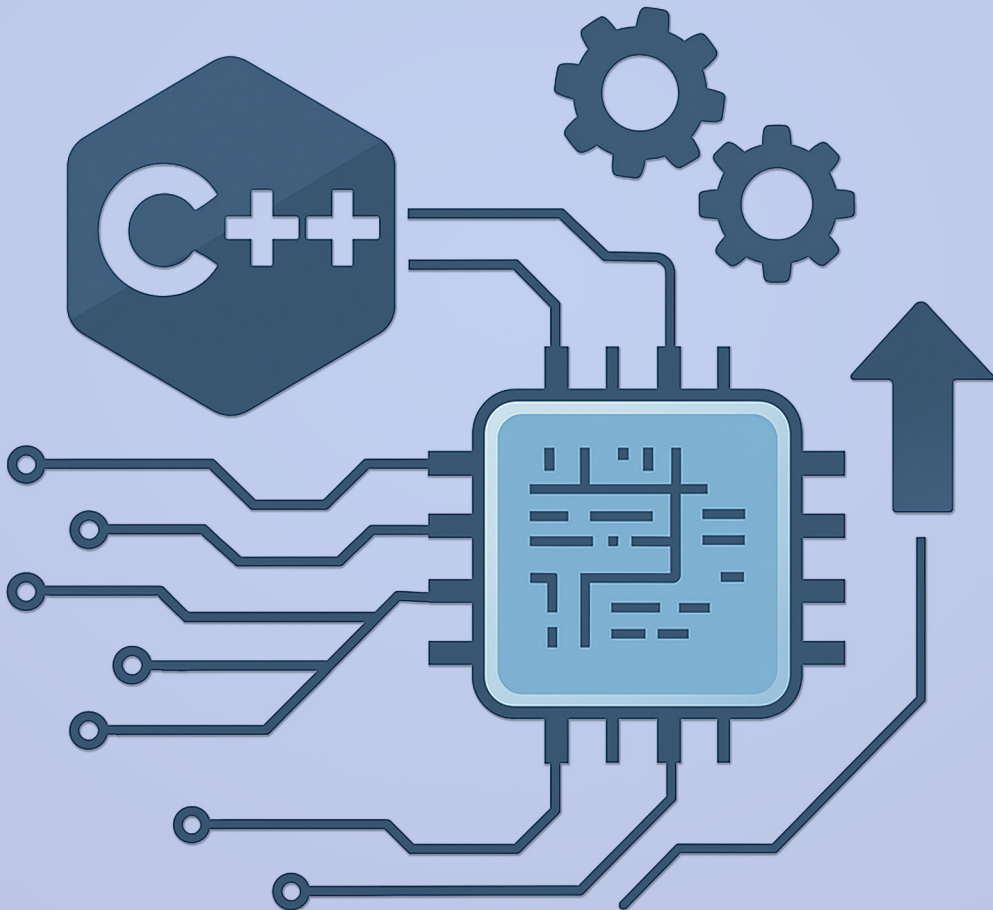


C++ Low-Level Optimization

Performance Secrets in a Concise Guide



C++ Low-Level Optimization

Performance Secrets in a Concise Guide

Prepared by Ayman Alheraki

simplifycpp.org

November 2025

Contents

Contents	2
Author’s Preface	13
1 Introduction to Low-Level Optimization	15
1.1 What Is Low-Level Optimization?	15
1.1.1 Beyond Big-O Notation	16
1.1.2 The Real Cost Model	16
1.1.3 Key Optimization Dimensions	17
1.1.4 A Guiding Principle	18
1.2 When to Optimize	18
1.2.1 The Wisdom Behind “Premature Optimization”	18
1.2.2 A Practical Optimization Workflow	18
1.2.3 Where Optimization Matters Most	19
1.2.4 When Not To Optimize	19
1.3 The Performance Mindset	20
1.3.1 Measurement First	20
1.3.2 Understanding Cost Models	20
1.3.3 Data-Centric Thinking	21
1.3.4 Simplicity and Predictability	21

1.3.5	The Power of Incremental Gains	22
1.3.6	Knowing When to Stop	22
	Exercise (Chapter 1)	23
2	Modern CPU and Memory Architecture	24
2.1	A Simplified CPU Model	24
2.1.1	CPU Cores and Execution Units	25
2.1.2	Pipelines	25
2.1.3	Out-of-Order Execution	26
2.1.4	Registers: The Fastest Storage	27
2.1.5	Memory: The Great Bottleneck	27
2.2	Cache Hierarchy and Locality	28
2.2.1	Spatial and Temporal Locality	28
2.2.2	Cache Lines	28
2.2.3	Direct-Mapped, Set-Associative, and Fully Associative Caches . . .	29
2.2.4	Illustrative Example: Vector vs List Traversal	29
2.3	The Translation Lookaside Buffer (TLB)	30
2.3.1	Why TLB Misses Matter	30
2.4	NUMA (Non-Uniform Memory Access)	31
2.5	Hardware Prefetching	31
2.6	Cache Coherence and False Sharing	32
2.6.1	False Sharing: A Devastating Performance Killer	32
2.6.2	Example	32
2.6.3	How to Avoid False Sharing	33
2.7	Putting It All Together	33
	Exercise (Chapter 2)	34

3	Compiler Optimizations and Build Modes	35
3.1	Optimization Levels	35
3.1.1	-O0: No Optimizations	36
3.1.2	-O1: Basic Optimizations	36
3.1.3	-O2: Recommended Default	37
3.1.4	-O3: Aggressive Optimizations	38
3.1.5	-Ofast: Breaking the Rules	39
3.1.6	Practical Guidelines	39
3.2	Link-Time Optimization (LTO)	39
3.2.1	How LTO Works	40
3.2.2	Benefits of LTO	40
3.2.3	CMake Integration	41
3.2.4	ThinLTO	41
3.3	Profile-Guided Optimization (PGO)	41
3.3.1	How PGO Works	42
3.3.2	What PGO Optimizes	42
3.3.3	Huge Performance Gains	43
3.4	Target-Specific Optimizations	43
3.4.1	-march=native	44
3.4.2	-mtune=	44
3.4.3	Caveat: Portability	44
3.5	Debug vs Release Builds	44
3.5.1	Debug Builds (-O0, -g)	44
3.5.2	Release Builds (-O2, -O3)	45
3.6	Debugging Optimized Code	45
3.7	Interprocedural Optimization (IPO)	45
3.8	Inlining: Friend and Enemy	46

3.8.1	Benefits of Inlining	46
3.8.2	Risks of Over-Inlining	46
3.9	Escape Analysis and SROA	47
3.10	Vectorization and Loop Optimizations	48
3.10.1	Using restrict / <code>__restrict</code>	48
	Exercise (Chapter 3)	49
4	Data Layout and Cache-Friendly Design	50
4.1	Array of Structures vs Structure of Arrays	50
4.1.1	Array of Structures (AoS)	51
4.1.2	Structure of Arrays (SoA)	51
4.1.3	Performance Implications	52
4.1.4	Hybrid Layouts: AoSoA and Tiles	53
4.2	Alignment and Padding	53
4.2.1	Alignment Basics	54
4.2.2	Automatic Padding	54
4.2.3	Cache Line Alignment	55
4.3	Hot/Cold Data Splitting	55
4.3.1	Example	56
4.4	Data Locality and Cache-Friendly Iteration	57
4.4.1	Sequential Access	57
4.4.2	Strided Access	57
4.4.3	Random Access	58
4.5	Small Object Optimization and In-Place Storage	58
4.5.1	Use <code>std::vector</code> Instead of Many <code>new</code>	59
4.5.2	Small String Optimization (SSO)	59
4.5.3	Small Buffer Optimization (SBO) for Custom Types	59
4.5.4	Memory Pools and Custom Allocators	60

4.6	Object Contiguity and Pointer Chasing	60
4.6.1	Why Contiguity Matters	60
4.6.2	Linked Lists: The Performance Villain	61
4.6.3	Better Alternatives	61
4.7	False Sharing in Data Layout	61
4.8	Case Study: Physics Update Loop	62
4.8.1	AoS Performance Characteristics	62
4.8.2	SoA Version	62
	Exercise (Chapter 4)	63
5	Control Flow, Branch Prediction, and Branchless Code	64
5.1	Branch Prediction Basics	64
5.1.1	Why Branch Prediction Exists	65
5.1.2	Prediction-Friendly Patterns	65
5.1.3	Difficult-to-Predict Patterns	66
5.2	Deep Dive: How Branch Predictors Work	67
5.2.1	Static Predictors	67
5.2.2	Dynamic Predictors	67
5.2.3	Branch Target Buffer (BTB)	67
5.2.4	Speculative Execution	68
5.3	The Cost of Misprediction	68
5.4	Reducing Branches	68
5.4.1	Fast-Path / Slow-Path Structure	68
5.4.2	Reordering Code	69
5.4.3	Avoiding Unpredictable Branches in Hot Loops	70
5.5	Branchless Techniques	70
5.5.1	Clamp Example	70
5.5.2	Conditional Select Instructions	71

5.5.3	Multiplication Trick	72
5.5.4	Lookup Table (LUT) Based Branch Removal	72
5.5.5	SIMD Predication	73
5.6	When Branchless Code Is Not Better	73
5.7	Practical Patterns for Predictable Branches	74
5.7.1	Sort Data to Improve Predictability	74
5.7.2	Group Data by Likelihood	74
5.7.3	Use Sentinel Values	75
5.8	Case Study: Counting Elements Above a Threshold	75
	Exercise (Chapter 5)	76
6	Cost of Abstractions: Inlining, Functions, and Virtual Calls	77
6.1	Function Calls and Inlining	77
6.1.1	The Mechanics of a Function Call	78
6.1.2	Inlining: Eliminating Call Overhead	79
6.1.3	What inline Really Means	79
6.1.4	Inlining Heuristics	80
6.1.5	Dangers of Over-Inlining	80
6.2	Virtual Functions	81
6.2.1	Mechanics of a Virtual Call	81
6.2.2	Devirtualization	81
6.2.3	When Virtual Functions Hurt Performance	82
6.3	Alternatives to Virtual Functions	83
6.3.1	Templates and Static Polymorphism	83
6.3.2	CRTP (Curiously Recurring Template Pattern)	83
6.3.3	Variant-Based Polymorphism	83
6.3.4	Function Objects and Lambdas	84
6.4	Zero-Cost Abstractions	84

6.4.1	Cost-Free RAII	84
6.4.2	Views Instead of Copies	85
6.4.3	Zero-Cost Lambdas	85
6.4.4	The Importance of Transparency	85
6.5	Type Erasure and Its Costs	85
6.5.1	Costs of <code>std::function</code>	86
6.5.2	Better Alternatives	86
6.6	Code Bloat and Its Performance Impact	86
6.7	Static vs Dynamic Polymorphism: Quantifying Costs	87
6.8	Case Study: Three Implementations of a Hot Loop	88
6.8.1	Raw Pointer Loop	88
6.8.2	Helper Function	88
6.8.3	Template Functor	88
	Exercise (Chapter 6)	89
7	Move Semantics, Value Categories, and Copies	90
7.1	The Cost of Copies	90
7.2	Value Categories and Why They Matter	92
7.3	Move Semantics in Practice	93
7.3.1	Move Constructor and Move Assignment	93
7.3.2	When Does Moving Happen?	94
7.3.3	<code>std::move</code> Is Not a Move	94
7.4	RVO, NRVO, and Copy Elision	94
7.4.1	Why RVO Matters	95
7.5	Perfect Forwarding and References	96
7.6	Minimizing Copies with Views	96
7.6.1	Other Useful Views	97
7.7	Optimizing Function Interfaces	98

7.8	Cheaply Movable Types	98
7.9	Common Pitfalls That Prevent Moves	99
7.9.1	Accidental Copies	99
7.9.2	Using Moved-From Objects Unexpectedly	100
7.9.3	Overloading Mistakes	100
7.10	Case Study: Passing by Value vs View	100
7.10.1	Problems	100
7.10.2	Refactor	101
	Exercise (Chapter 7)	101
8	Dynamic Memory, Allocators, and Ownership	103
8.1	The Hidden Cost of Allocation	103
8.1.1	Costs Inside a Memory Allocator	104
8.1.2	Common Allocation Anti-Patterns	105
8.1.3	Practical Advice	105
8.2	Memory Fragmentation and Growth	105
8.3	Allocator Design in C++	106
8.3.1	Types of Allocators	107
8.4	Monotonic Buffer Resource	107
8.5	Memory Pools and Free-Lists	108
8.6	NUMA Awareness (Advanced)	109
8.7	Ownership and RAII	110
8.7.1	Smart Pointers	110
8.7.2	RAII Wrappers and Performance	111
8.8	Avoiding Overuse of <code>shared_ptr</code>	111
8.9	Case Study: Eliminating Allocations in a Hot Loop	112
	Exercise (Chapter 8)	113

9	SIMD and Auto-Vectorization	114
9.1	What Is SIMD?	114
9.1.1	SIMD Across Architectures	115
9.2	Why SIMD Matters	116
9.3	Enabling Auto-Vectorization	116
9.3.1	When Auto-Vectorization Fails	117
9.3.2	Fixing Alias Problems	117
9.3.3	Dealing With Data Dependencies	118
9.4	Loop Transformations for Vectorization	118
9.4.1	Hoist Invariants	118
9.4.2	Remove Conditionals from Hot Loops	118
9.4.3	Align Data	119
9.5	Explicit SIMD via Intrinsics	119
9.5.1	When to Use Intrinsics	120
9.6	High-Level SIMD Abstractions	120
9.7	Execution Policies and Vectorization	121
9.8	Gather, Scatter, and Masked Operations	122
9.9	When SIMD Fails or Underperforms	122
9.10	Case Study: Scalar vs SIMD Loop	123
	Exercise (Chapter 9)	124
10	Concurrency, Synchronization, and False Sharing	125
10.1	The Cost of Synchronization	125
10.1.1	Why Synchronization Is Expensive	126
10.1.2	Avoid Synchronization in Hot Paths	126
10.1.3	Coarse-Grained vs Fine-Grained Locks	126
10.2	CPU Cache Coherence and Synchronization Costs	127
10.3	Mutexes and Lock Contention	128

10.3.1	Spinlocks vs Blocking Mutexes	128
10.4	Atomics and Memory Order	129
10.4.1	Relaxed Atomics	129
10.4.2	Acquire/Release Semantics	130
10.4.3	Sequential Consistency	130
10.5	Avoiding Contention	130
10.5.1	Sharding	130
10.5.2	Per-Thread Buffers	130
10.5.3	Read-Mostly Data	131
10.5.4	Combining (Software Combining Tree)	131
10.5.5	Avoiding False Sharing	131
10.5.6	Recognizing False Sharing	132
10.6	Thread Scheduling and Oversubscription	132
10.7	Lock-Free and Wait-Free Algorithms	133
10.8	Case Study: Multi-Threaded Counter	133
10.8.1	Version 1: Mutex	133
10.8.2	Version 2: Atomic Counter	134
10.8.3	Version 3: Thread-Local Counters	135
	Exercise (Chapter 10)	135
11	Measurement, Benchmarking, and Profiling	136
11.1	Micro-Benchmarks	136
11.1.1	Warm-Up Phase	137
11.1.2	Avoid Dead-Code Elimination	137
11.1.3	Avoid Measuring I/O	138
11.1.4	A Minimal Benchmark Helper	138
11.1.5	Limitations of Micro-Benchmarks	138
11.2	Macro-Benchmarks	139

11.3	System Profilers	140
11.3.1	Linux Profilers	141
11.3.2	Windows Profilers	141
11.3.3	macOS Profilers	142
11.4	Interpreting Profiling Results	142
11.4.1	Look For:	143
11.4.2	Common Patterns and Their Causes	143
11.5	Benchmarking Best Practices	144
11.6	Automated Benchmarking Tools	145
11.6.1	Google Benchmark	145
11.6.2	Celero	146
11.6.3	hayai	146
	Exercise (Chapter 11)	146
12	Common Patterns, Anti-Patterns, and Checklists	148
12.1	Good Patterns	148
12.2	Anti-Patterns	151
12.3	Practical Optimization Checklist	153
	Exercise (Chapter 12)	154
	Final Conclusion	155
	References	160

Author's Preface

Modern C++ (from C++11 up to C++23 and beyond) has transformed how we write high-performance systems. It is no longer just about raw pointers and hand-written loops: the language and standard library now contain powerful abstractions that are, when used correctly, zero-cost at runtime. At the same time, modern hardware has grown extremely complex: multi-core, multi-level caches, sophisticated branch predictors, out-of-order execution, SIMD units, and often NUMA memory.

This mini-booklet — C++ Low-Level Optimization – Performance Secrets in a Concise Guide — aims to bridge the gap between high-level C++ code and the low-level reality of how that code executes on the machine. It targets developers who already know C++ syntax and semantics, but want to:

- Understand how compilers turn C++ into machine code.
- Reason about costs: function calls, branches, memory allocation, cache misses, synchronization, etc.
- Write efficient, cache-conscious, branch-predictor-friendly code.
- Use modern features (move semantics, constexpr, templates, ranges) without sacrificing performance.
- Apply profiling tools and measurement techniques to prove performance improvements.

This is not a collection of micro-tricks. Instead, it is structured as a practical roadmap:

1. Start from the hardware: caches, pipelines, branches, and memory.
2. See how the compiler optimizes (and sometimes fails to optimize) your code.
3. Learn to shape your data and algorithms around locality, predictability, and vectorization.
4. Learn to measure, profile, and validate optimizations.

You do not need to be an assembly expert to benefit from this guide, but reading compiler output (or at least recognizing patterns) will help you greatly. Throughout the chapters, you will find code examples, explanations of subtle pitfalls, and practical checklists you can apply to your own projects.

I hope this concise guide becomes a useful reference on your journey to mastering C++ low-level performance.

Ayman Alheraki

Chapter 1

Introduction to Low-Level Optimization

Low-level optimization in C++ is an advanced discipline that focuses on aligning your source code with the underlying behavior of modern CPUs, memory subsystems, and compilers. While C++ gives developers high-level expressive power, its real strength lies in the fact that it compiles directly to native machine code. This chapter lays the conceptual foundation for the rest of the booklet: what low-level optimization means, why it matters, when it should be applied, and how to develop the right mindset for doing it effectively.

1.1 What Is Low-Level Optimization?

Low-level optimization describes a family of techniques that aim to reduce the cost of program execution by making code friendly to the CPU, caches, memory hierarchy, and compiler optimizer. Unlike high-level optimization (which focuses mostly on algorithmic complexity or design-level improvements), low-level optimization deals with the actual runtime behavior of your program on the hardware.

This involves understanding:

- The way instructions flow through pipelines.
- How and when the CPU fetches data from memory.
- How branches impact pipeline performance.
- How compilers transform your C++ into machine code.
- How data layout influences cache locality.

1.1.1 Beyond Big-O Notation

Big-O notation gives you a theoretical upper bound on computational complexity, but it does not reflect:

- How your data is accessed in memory.
- How operations map to CPU instructions.
- How predictable your control flow is.
- What the cache, TLB, or prefetcher are doing.
- How SIMD instructions could speed up loops.

For example, two $O(n)$ algorithms can differ by 10x or even 100x in actual runtime depending on locality, branch predictability, and vectorization.

1.1.2 The Real Cost Model

Low-level optimization requires a mental cost model of operations. Approximate latencies (varies by architecture):

- Integer add: ~ 1 cycle

- Floating multiply: 3–5 cycles
- L1 cache hit: 3–5 cycles
- L2 cache hit: 10–15 cycles
- L3 cache hit: 30–50 cycles
- DRAM access: 150–300 cycles
- Branch misprediction: 10–20 cycles (pipeline flush)

Thus, the cost of a single cache miss can be equivalent to hundreds of arithmetic operations. This is why low-level optimization often focuses on memory patterns rather than arithmetic.

1.1.3 Key Optimization Dimensions

Low-level optimization spans several interconnected dimensions:

Instruction Efficiency Reducing number of CPU instructions or improving ILP (Instruction-Level Parallelism).

Data Locality Arranging data contiguously to minimize cache misses.

Branch Predictability Restructuring logic to avoid unpredictable branches.

Parallelism Exploiting SIMD and multi-core parallelism.

Memory Efficiency Reducing allocations, avoiding fragmentation, improving reuse.

Compiler Assistance Writing code in a way that enables compiler optimizations.

1.1.4 A Guiding Principle

Make it easy for the compiler and hardware to do the right thing.

The compiler is powerful—but only if your code is structured symmetrically, predictably, and transparently.

1.2 When to Optimize

Low-level optimization is a powerful but expensive tool. It requires time, expertise, and often reduces code clarity. Therefore, it must be applied selectively.

1.2.1 The Wisdom Behind “Premature Optimization”

Donald Knuth’s famous quote:

“Premature optimization is the root of all evil.”

is often misunderstood. It does not mean “do not optimize.” It means:

- Do not optimize before measuring.
- Do not optimize code that does not matter.
- Do not guess about performance characteristics.

1.2.2 A Practical Optimization Workflow

Low-level optimization should follow a structured process:

1. Write correct, readable code.

2. Profile the program using real workloads.
3. Identify hotspots—the 1–5% of code where optimization matters.
4. Analyze problem areas: memory access? branches? allocations?
5. Choose targeted optimizations based on evidence.
6. Re-measure to validate improvements.

1.2.3 Where Optimization Matters Most

Some parts of the codebase naturally demand low-level optimization:

- Inner loops running inside numerical computations or physics engines.
- Low-latency systems (HFT, embedded, real-time audio, games, signal processing).
- High-throughput servers handling millions of requests.
- Core data structures and algorithms used throughout the program.
- Serialization, parsing, hashing, sorting routines.

These areas must be extremely efficient because:

Even a tiny improvement scales to massive overall savings.

1.2.4 When Not To Optimize

Avoid low-level optimization in:

- Rarely used code paths.
- Initialization routines that run once.

- Code where clarity is more valuable than speed.
- Features likely to change soon.
- Parts of the system where performance does not matter.

Optimization is a trade-off. Maintainability is also a performance factor in long-lived systems.

1.3 The Performance Mindset

Low-level optimization is not about “making code faster.” It is about learning to think like the machine, developing an engineering approach that balances clarity, correctness, and efficiency.

1.3.1 Measurement First

Never rely solely on intuition. Humans are terrible at predicting performance in modern architectures. Small changes in data layout can outperform algorithmic improvements. Always measure:

- Use micro-benchmarks for isolated behaviors.
- Use system-level profilers (perf, VTune, Instruments).
- Benchmark representative workloads.

1.3.2 Understanding Cost Models

A performance engineer must know how expensive different operations are. Two examples:

- An unpredictable branch can cost more than a dozen arithmetic operations.
- A cache miss can cost more than a tight loop of 200 simple operations.

Thus:

Optimizing memory access patterns often yields larger gains than optimizing math.

1.3.3 Data-Centric Thinking

Many developers focus on algorithms but ignore how data flows through the system. Key principle of modern optimization:

Data layout is more important than code layout.

Memory-friendly code is:

- contiguous,
- coherent,
- cache-aligned,
- predictable,
- and reuses recent data.

1.3.4 Simplicity and Predictability

Simple code is easier for the compiler to optimize. Complicated branch structures, unpredictable data dependencies, and dynamic behavior can confuse the optimizer. Prefer:

- Flat loops over nested branches.
- Clear ownership semantics.
- Consistent patterns of memory access.
- Templates for static polymorphism rather than dynamic dispatch in hot loops.

1.3.5 The Power of Incremental Gains

In low-level optimization, performance gains compound:

- 5% faster memory access + 10% fewer branch misses + 7% better vectorization results in far more than a 22% speed improvement.

Optimizing several layers—data layout, branches, instructions, vectorization—yields multiplicative benefits.

1.3.6 Knowing When to Stop

Optimization can continue indefinitely. A disciplined engineer knows:

- When diminishing returns begin.
- When code becomes too complex to justify small gains.
- When hardware upgrades solve more problems than extreme optimization.

Performance is important—but so is engineering time.

Exercise

Pick one of your existing C++ projects. Identify one function you believe might be performance-critical.

1. Run it through a profiler such as perf, VTune, Instruments, or your IDE's profiler.
2. Compare the results with your intuition.
3. Determine whether the function is part of a hot code path.
4. If it is hot: Examine it for possible inefficiencies (branches, allocations, poor locality).
5. If it is not hot: Consider how your intuition misled you—and why measurement matters.

Document your findings. This exercise trains you to think like a performance engineer.

Chapter 2

Modern CPU and Memory Architecture

Modern performance begins with understanding how the CPU executes instructions and how memory is delivered to it. While C++ compiles to extremely efficient native machine code, the speed of that machine code depends heavily on what the hardware can do, how it schedules instructions, and how fast it can move data through its hierarchy.

This chapter provides a detailed but practical model of modern CPUs and memory systems, focusing specifically on performance characteristics that matter to C++ programmers designing low-level, high-performance software.

2.1 A Simplified CPU Model

Real CPUs (Intel, AMD, ARM, Apple Silicon) are astonishingly complex machines containing billions of transistors, multiple cores, deep pipelines, speculative execution engines, reorder buffers, and sophisticated cache hierarchies. Fortunately, we do not need to understand every detail—the goal is to develop a useful mental model to guide optimization decisions.

Below is a simplified but accurate representation of what happens when you run C++ code.

2.1.1 CPU Cores and Execution Units

Each CPU consists of several cores. Each core has:

- A front-end for fetching and decoding instructions.
- A scheduler that reorganizes instructions for parallel execution.
- Multiple ALUs (integer arithmetic units).
- FPU (floating-point units).
- Load/store units for memory access.
- SIMD/vector units (SSE, AVX, NEON, SVE).

Modern CPUs are superscalar: they can execute multiple instructions per cycle (often 4–6 or more) as long as dependencies allow it. This is known as instruction-level parallelism (ILP).

2.1.2 Pipelines

Instructions pass through several pipeline stages:

- Fetch
- Decode
- Rename (register renaming)
- Dispatch

- Execute
- Retire

Pipelines allow high throughput, but they also introduce hazards:

- Data hazards: instruction B depends on the result of instruction A.
- Control hazards: branch instructions disrupt flow.
- Structural hazards: too few execution units for competing instructions.

Long pipelines (15–20 stages or more) make branch mispredictions and cache misses very expensive.

2.1.3 Out-of-Order Execution

Modern CPUs do not execute instructions strictly in the order written. Instead, they use:

- Reorder buffers (ROB)
- Reservation stations
- Scoreboards

These mechanisms allow the CPU to look ahead, find instructions whose dependencies are ready, and execute them early. This can greatly speed up code—but only if the code has enough independent work to expose ILP.

2.1.4 Registers: The Fastest Storage

Registers are extremely fast (single-cycle access). Compilers try to keep hot variables in registers whenever possible. Performance collapses when the compiler spills variables to memory, usually due to:

- Too many live variables.
- Complex expressions.
- Heavy use of large structs or classes.
- Deep inlined code.

As a C++ developer, favor:

- Simpler loop bodies.
- Small fixed-size data structures.
- Reducing unnecessary temporaries.

These keep register pressure low.

2.1.5 Memory: The Great Bottleneck

Even the fastest DRAM is hundreds of cycles slower than L1 cache. A single round-trip to memory can take as long as 200–300 CPU cycles. That is effectively an eternity compared to a single-cycle arithmetic instruction.

Thus:

Modern performance is defined more by memory behavior than by arithmetic.

2.2 Cache Hierarchy and Locality

CPUs compensate for slow main memory by using a hierarchy of caches. A modern CPU has:

- L1 cache: ultra-fast, tiny (32–64 KB)
- L2 cache: fast, medium-sized (256–1 024 KB)
- L3 cache: slower but large (8–64 MB shared among cores)

Each level is slower but larger than the previous. When data is not found in a given level, the CPU searches the next level. A miss in all levels results in a trip to DRAM.

2.2.1 Spatial and Temporal Locality

To make the best use of caches, programs must exhibit:

Spatial Locality Accessing memory addresses that are close together. Because caches operate on blocks (cache lines) of 64 bytes, accessing `v[i]` likely loads `v[i+1]`, `v[i+2]`, etc. into L1 automatically.

Temporal Locality Reusing data soon after it was last accessed. If a variable is reused quickly, it remains in L1 or L2 cache.

2.2.2 Cache Lines

The fundamental unit of caching is the cache line—a 64-byte chunk (on most x86-64 systems). Even if you only read one int, the entire 64-byte region is loaded.

This has important consequences:

- Strided access wastes bandwidth.
- Linked lists destroy spatial locality.
- AoS vs SoA decisions matter.

2.2.3 Direct-Mapped, Set-Associative, and Fully Associative Caches

Caches are not simple arrays. They use one of three strategies:

1. Direct-mapped: each address maps to exactly one slot. Fast but prone to conflicts.
2. N-way set-associative: most common. Each address maps to a set with N choices.
3. Fully associative: any address can go anywhere. Used only in tiny caches like TLBs.

Most modern L1 and L2 caches are 8-way associative; L3 caches may be up to 16–20-way.

2.2.4 Illustrative Example: Vector vs List Traversal

Revisiting the example:

```
#include <vector>
#include <list>

void sum_vector(const std::vector<int>& v, long long& res) {
    long long s = 0;
    for (int x : v) { s += x; }
    res = s;
}
```

```
void sum_list(const std::list<int>& l, long long& res) {  
    long long s = 0;  
    for (int x : l) { s += x; }  
    res = s;  
}
```

- Vector: contiguous elements → excellent spatial locality → CPU prefetchers work well → auto-vectorization possible.
- List: scattered nodes → random memory access → TLB misses → cache thrashing → no vectorization.

Even though both are $O(n)$, the real-world difference can be 5x–20x.

2.3 The Translation Lookaside Buffer (TLB)

The TLB is a small, extremely fast cache that stores virtual-to-physical page translations. Because modern processes use virtual memory, every memory access must first look up a page mapping.

If the mapping is not in the TLB, a TLB miss can be as expensive as a cache miss.

2.3.1 Why TLB Misses Matter

Because the TLB is small:

- Frequent random access over large arrays overwhelms it.
- Linked lists often cause TLB misses.
- Large data structures benefit from huge pages (2 MB or 1 GB pages).

If your memory footprint exceeds the TLB entries, performance suffers dramatically.

2.4 NUMA (Non-Uniform Memory Access)

On multi-socket systems (servers, HPC clusters), memory is divided into NUMA nodes. Each CPU socket has its own attached memory. Accessing local memory is fast; accessing remote memory is slower.

Rules for optimal NUMA performance:

- Threads should allocate and access memory on their own NUMA node.
- Thread migration can cause performance collapse.
- Bind threads using OS APIs or thread affinity.

While NUMA is less of a concern on laptops and desktops, it is crucial for scalable servers and multi-socket systems.

2.5 Hardware Prefetching

Modern CPUs include hardware prefetchers that attempt to guess which memory locations will be accessed soon and load them into cache early.

Common patterns recognized:

- Linear forward iteration (for (i=0; i<n; i++))
- Small fixed stride access
- Certain predictable pointer-chasing patterns

Patterns that defeat prefetchers:

- Random access (hash tables, linked lists)

- Indirection chains (like `A[B[i]]`)
- Data-dependent branches leading to unpredictable memory paths

Good prefetcher alignment can substantially boost performance. Poor patterns can completely stall pipelines.

2.6 Cache Coherence and False Sharing

In multi-core CPUs, caches must agree on shared data. This is enforced through a cache coherence protocol (e.g., MESI, MESIF, MOESI), ensuring all cores see consistent memory values.

2.6.1 False Sharing: A Devastating Performance Killer

False sharing happens when:

- Two threads access different variables
- But those variables reside on the same cache line

Even if they never touch each other's data, the cache line is repeatedly invalidated between cores.

2.6.2 Example

```
struct Counters {  
    alignas(64) long long a;  
    long long b; // shares cache line with 'a'  
};
```

If thread 1 updates `a` and thread 2 updates `b`, performance collapses.

2.6.3 How to Avoid False Sharing

- Align variables to cache lines using `alignas(64)`.
- Insert padding to separate variables.
- Group thread-local data together.

False sharing is often invisible in code, but devastating in practice.

2.7 Putting It All Together

Modern CPUs reward:

- contiguous memory,
- predictable control flow,
- sequential iteration,
- high ILP,
- and reduced synchronization.

They punish:

- pointer-chasing,
- random access,
- small allocations,
- unpredictable branches,

- false sharing,
- and TLB/cache thrashing.

Understanding this model is key to designing high-performance C++.

Exercise

Measure access time of a large `std::vector<int>` using two patterns:

1. Sequential access: `v[i]`
2. Strided access: `v[i + 16]`, `v[i + 32]`, etc.

Use a high-resolution clock to time both patterns.

- Plot or record total time.
- Observe how stride length impacts performance.
- Explain the results in terms of cache lines, TLB behavior, and prefetching.

This exercise builds intuition for memory locality.

Chapter 3

Compiler Optimizations and Build Modes

Modern C++ performance relies not only on your code but also on the compiler's ability to transform that code into high-quality machine instructions. Compilers today (GCC, Clang/LLVM, MSVC) contain extremely sophisticated optimizers that apply hundreds of transformations: constant folding, inlining, dead code elimination, loop unrolling, vectorization, escape analysis, and auto-parallelization.

This chapter provides a comprehensive understanding of compiler optimization levels, link-time and profile-guided optimizations, and how to use build modes effectively to extract maximum performance from modern hardware.

3.1 Optimization Levels

Every major compiler provides several optimization levels, each balancing:

- build time,
- debugging friendliness,
- generated code size,

- and runtime performance.

Below is a detailed description of the most commonly used levels in GCC and Clang.

3.1.1 -O0: No Optimizations

Characteristics:

- Most direct translation of C++ to machine code.
- Keeps variables in memory, not registers.
- No inlining.
- No loop transformations.
- No dead-code elimination.
- Step-by-step debugging works perfectly.

Suitable for:

- initial debugging,
- teaching and experimentation,
- inspecting compiler behavior.

3.1.2 -O1: Basic Optimizations

The compiler applies:

- simple inlining,
- constant propagation,

- some local dead-code elimination,
- basic loop optimizations,
- better register allocation.

Useful when you want some performance but still maintain readable debugging traces.

3.1.3 -O2: Recommended Default

-O2 enables the majority of optimizations that do not increase code size excessively or risk violating strict standards compliance. It typically includes:

- loop unrolling and fusion,
- common subexpression elimination,
- full dead-code elimination,
- improved inlining heuristics,
- better vectorization attempts,
- branch prediction hints,
- basic alias analysis,
- interprocedural optimizations within the compilation unit.

This is the standard production setting for compiled C++ code.

3.1.4 -O3: Aggressive Optimizations

Enables all of -O2 plus aggressive transformations intended for compute-heavy workloads:

- larger inlining budgets,
- full loop unrolling,
- loop vectorization (outer-loop vectorization),
- instruction scheduling for ILP,
- speculative execution,
- enhanced alias analysis,
- more aggressive register promotion.

Potential downsides:

- Larger binaries (code bloat).
- More inlining → worse instruction cache pressure.
- Debugging becomes nearly impossible.
- Performance may drop for some workloads.

Best for:

- numerical kernels,
- high-performance computing,
- tight loops and SIMD-heavy tasks.

3.1.5 -Ofast: Breaking the Rules

-Ofast applies everything from -O3 and additionally:

- disables strict IEEE floating-point compliance,
- removes checks for NaN and Inf,
- assumes associative math is safe,
- allows unsafe math transformations,
- enables -ffast-math.

This can yield substantial speedups in simulation, graphics, signal processing, and machine-learning workloads—but beware: results may differ from standard-compliant outputs.

3.1.6 Practical Guidelines

- Use -O2 as your default release-level optimization.
- Use -O3 for performance-critical algorithms.
- Use -Ofast only when you fully understand the consequences.
- Always benchmark after changing optimization levels.

3.2 Link-Time Optimization (LTO)

Traditional compilation separates translation units (.cpp files). Each is optimized independently. This prevents the compiler from:

- inlining across translation units,
- removing dead code across units,
- performing global alias analysis,
- optimizing interdependent templates,
- applying global constant propagation.

LTO removes these restrictions.

3.2.1 How LTO Works

During normal builds:

- .cpp files → optimized .o files.

With LTO:

- The compiler produces IR (LLVM IR or GCC GIMPLE).
- The IR is linked together.
- Full-program optimization is applied.
- Only then is machine code generated.

3.2.2 Benefits of LTO

- Cross-module inlining.
- Dead code elimination across libraries.
- More effective constant propagation.

- Fewer indirect calls.
- Better alias analysis.
- Smaller final binaries.

3.2.3 CMake Integration

```
set(CMAKE_INTERPROCEDURAL_OPTIMIZATION ON)
```

This automatically enables LTO if supported by your toolchain.

3.2.4 ThinLTO

Clang/LLVM's ThinLTO is a scalable version of LTO:

- parallel optimization across files,
- faster incremental builds,
- smaller memory usage,
- similar performance to full LTO.

Recommended for large C++ codebases.

3.3 Profile-Guided Optimization (PGO)

PGO (also known as FDO—Feedback-Directed Optimization) uses runtime information to guide compiler decisions. It allows the compiler to learn:

- Which branches are hot or cold.

- Which functions are most frequently executed.
- Which indirect calls usually resolve to which targets.
- Cache-line and alignment-sensitive behaviors.
- Typical loop trip counts.
- Real-world data distributions.

This enables optimizations that static analysis cannot infer.

3.3.1 How PGO Works

Workflow:

1. Build instrumented binaries:

```
g++ -O2 -fprofile-generate main.cpp -o app
```

2. Run the program with representative inputs:

```
./app <real_workload>
```

This produces profile data files (e.g., .gcda, .profrac).

3. Rebuild using the profile:

```
g++ -O2 -fprofile-use main.cpp -o optimized_app
```

3.3.2 What PGO Optimizes

Branch Prediction The compiler knows actual branch probabilities.

Inlining Decisions Hot functions are inlined aggressively; cold functions remain out-of-line.

Indirect Call Devirtualization The compiler sees which virtual calls resolve to which derived types most of the time.

Code Layout The compiler rearranges code so frequently used functions are placed close together in memory, improving I-cache locality.

Loop Optimizations Typical loop bounds can be inferred from runtime behavior.

3.3.3 Huge Performance Gains

PGO can yield:

- 10–30% gains in large applications,
- 50–100% in tight loops or branch-heavy code,
- dramatic improvements in parsers, database engines, and CPU-bound libraries.

Many industry compilers (Chromium, Firefox, Clang, MSVC standard library) depend heavily on PGO for their performance.

3.4 Target-Specific Optimizations

Compilers can generate code for different architectures, using `-march=` and `-mtune=` flags:

3.4.1 -march=native

Generates code tuned for your CPU:

- enables AVX/AVX2/AVX512 automatically,
- uses fast scalar instructions,
- uses specialized instructions (BMI2, POPCNT, LZCNT),
- improves register allocation and unrolling decisions.

This can give massive performance boosts on modern CPUs.

3.4.2 -mtune=

Adjusts scheduling and microarchitectural tuning—useful when distributing a binary but wanting some CPU-specific boosts.

3.4.3 Caveat: Portability

-march=native is not suitable for binaries distributed across unknown systems.

3.5 Debug vs Release Builds

Optimization levels dramatically change application behavior under debugging.

3.5.1 Debug Builds (-O0, -g)

- Variables preserved exactly.
- Execution order follows source code.

- Values visible in debugger.
- Slower by 10x–100x compared to Release builds.

3.5.2 Release Builds (-O2, -O3)

- Code reordered, inlined, eliminated.
- Variables placed in registers (or optimized away).
- Stepping in debugger becomes nonlinear.
- Realistic performance profiling is only valid here.

3.6 Debugging Optimized Code

Even in Release mode, using `-g` preserves debug symbols. However:

- The compiler may remove variables.
- Breakpoints may hit unexpected places.
- Stepping line-by-line may jump unpredictably.

This is normal behavior in highly optimized code.

3.7 Interprocedural Optimization (IPO)

IPO is a set of optimizations applied across function boundaries:

- Cross-function constant propagation.

- Better alias analysis.
- Global inlining decisions.
- Dead code elimination across modules.
- Function cloning for specialization.

IPO is most effective when combined with LTO and PGO.

3.8 Inlining: Friend and Enemy

Inlining removes call overhead and allows additional optimizations—but must be used carefully.

3.8.1 Benefits of Inlining

- Eliminates overhead of function calls.
- Unlocks constant propagation.
- Enables vectorization and unrolling.
- Allows better register allocation.

3.8.2 Risks of Over-Inlining

- Code bloat → instruction cache misses.
- Slower performance if hot loops expand beyond I-cache capacity.
- Longer compilation time.

- Harder debugging.

Compilers attempt to balance this automatically, but PGO and LTO improve decisions substantially.

3.9 Escape Analysis and SROA

Modern compilers use:

- Escape Analysis: determines whether an object escapes a function's scope.
- Scalar Replacement of Aggregates (SROA): breaks structs/objects into individual scalars.

Benefits:

- Eliminates heap allocations.
- Promotes fields into registers.
- Enables more aggressive constant folding and dead-field elimination.

Example transformation:

```
struct P { int x, y; };
```

```
P foo() {  
    P p{1, 2};  
    return p;  
}
```

This may compile as if P never existed—just two registers.

3.10 Vectorization and Loop Optimizations

Compilers attempt auto-vectorization when:

- memory accesses are regular,
- loop bounds are known,
- dependencies are simple,
- pointers do not alias.

Optimizations include:

- loop unrolling,
- loop fusion,
- loop interchange,
- strength reduction,
- pointer alias analysis,
- array bounds elimination.

3.10.1 Using restrict / __restrict__

This qualifier informs the compiler that pointers do not alias, greatly improving vectorization possibilities.

Example:

```
void scale(float* __restrict__ a, float* __restrict__ b, size_t n);
```

Exercise

Create a small benchmark that sums a large `std::vector<int>`:

1. Compile and run with `-O0`.
2. Compile and run with `-O2`.
3. Compile and run with `-O3`.
 - Compare timings.
 - Inspect generated assembly (Compiler Explorer recommended).
 - Note differences in inlining, unrolling, and vectorization.
 - Observe how binary size changes.
 - Write a brief explanation of why `-O3` may or may not be faster.

Chapter 4

Data Layout and Cache-Friendly Design

Data layout is one of the strongest influences on real-world performance in modern systems. Although algorithms determine asymptotic complexity, data layout determines how efficiently the CPU can feed those algorithms with data. In practice, poor cache behavior can degrade performance by a factor of 5–50 compared to the theoretical peak. This chapter explores how memory organization, structure layouts, container choice, and alignment decisions affect performance. We focus on the CPU cache hierarchy, prefetching behavior, vectorization opportunities, and memory allocation patterns in modern C++.

4.1 Array of Structures vs Structure of Arrays

One of the most fundamental design choices is how to group related data: either as Array of Structures (AoS) or Structure of Arrays (SoA).

4.1.1 Array of Structures (AoS)

```
struct ParticleAoS {  
    float x, y, z;  
    float vx, vy, vz;  
};
```

```
std::vector<ParticleAoS> particles;
```

Characteristics:

- Natural, intuitive representation of objects.
- Excellent when operations require access to all fields at once (e.g., physics integrations).
- Each particle is contiguous, but its fields are interleaved with fields of other particles.

Memory layout (small excerpt):

```
| x y z vx vy vz | x y z vx vy vz | x y z vx vy vz | ...
```

4.1.2 Structure of Arrays (SoA)

```
struct ParticleSoA {  
    std::vector<float> x, y, z;  
    std::vector<float> vx, vy, vz;  
};
```

Characteristics:

- Fields of the same type are stored contiguously.

- Excellent for vectorization and bulk operations.
- Ideal when processing one property across many elements.
- Better for SIMD (SSE/AVX/NEON).

Memory layout:

```
x: | x x x x x x x ... |  
y: | y y y y y y y ... |  
z: | z z z z z z z ... |  
vx: | ...                |  
...
```

4.1.3 Performance Implications

AoS Advantages

- More intuitive API.
- Good for per-object operations where all fields are required.
- Coherent object representation.

AoS Disadvantages

- Poor spatial locality when processing only some fields.
- Compiler may struggle to autovectorize due to interleaving.
- Larger structures cause cache pollution.

SoA Advantages

- Maximizes spatial locality for single-field operations.
- Enables tight vectorized loops.
- Eliminates unnecessary loads.

SoA Disadvantages

- Less intuitive API.
- Difficult when frequent random access to whole objects is needed.
- More complex ownership semantics.

4.1.4 Hybrid Layouts: AoSoA and Tiles

Some engines (HPX, ECS frameworks, rendering engines) use AoSoA—Array of Structures of Arrays—splitting data into blocks sized for cache lines or SIMD registers.

Example: 16 particles per block (matching AVX-512 width on floats).

This reduces cache misses and aligns data with the hardware's vector width.

4.2 Alignment and Padding

Alignment specifies how memory addresses should be aligned for optimal access. The CPU often requires certain instructions (e.g., SIMD loads) to operate on memory aligned by 16, 32, or 64 bytes.

4.2.1 Alignment Basics

```
struct alignas(16) Vec4 {  
    float x, y, z, w;  
};
```

Effects of alignment:

- Ensures the object begins at a 16-byte boundary.
- Faster loads/stores on SIMD-capable architectures.
- Avoids penalties from unaligned memory access.

4.2.2 Automatic Padding

The compiler inserts padding between structure members to satisfy alignment requirements.

Example (poor ordering):

```
struct Bad {  
    char a;  
    double b;  
    int c;  
};
```

Memory layout:

| a | pad | pad | pad | pad | pad | pad | pad | b | c | pad |

Better ordering:

```
struct Good {  
    double b;  
    int c;  
    char a;  
};
```

Padding reduced significantly.

4.2.3 Cache Line Alignment

For performance-critical structs used in multithreaded systems, align objects to the cache line size (usually 64 bytes):

```
struct alignas(64) Node {  
    int value;  
    // padding automatically added  
};
```

Benefits:

- Avoids false sharing.
- Keeps hot data independent.

4.3 Hot/Cold Data Splitting

Not all fields of an object are accessed with the same frequency. A classic high-performance design technique is hot/cold splitting—separating frequently accessed data from rarely used data.

4.3.1 Example

```
struct Entity {  
    // Hot  
    float x, y;  
    int hp;  
  
    // Cold  
    std::string name;  
    std::vector<int> inventory;  
};
```

Instead:

```
struct EntityHot {  
    float x, y;  
    int hp;  
};  
  
struct EntityCold {  
    std::string name;  
    std::vector<int> inventory;  
};  
  
std::vector<EntityHot> hotData;  
std::vector<EntityCold> coldData;
```

Advantages:

- Hot loops process only small, cache-resident data.
- Cold data accessed infrequently avoids polluting L1.

This technique is heavily used in game engines, simulation frameworks, and rendering pipelines.

4.4 Data Locality and Cache-Friendly Iteration

To exploit the cache hierarchy, iterate in a predictable, sequential manner.

4.4.1 Sequential Access

This is ideal:

```
for (size_t i = 0; i < N; ++i)
    sum += a[i];
```

Because:

- Hardware prefetchers detect the pattern.
- Adjacent elements share cache lines.
- TLB usage is efficient.

4.4.2 Strided Access

This is less ideal:

```
for (size_t i = 0; i < N; i += 16)
    sum += a[i];
```

Causes:

- Reduced spatial locality.
- Prefetchers may not detect the pattern.
- Each access hits a different cache line.

4.4.3 Random Access

Worst case:

```
for (auto i : indices)
    sum += a[i];
```

This often destroys performance because:

- Every access is unpredictable.
- TLB misses are common.
- Prefetching is nearly impossible.

4.5 Small Object Optimization and In-Place Storage

Allocating many small objects on the heap introduces several costs:

- allocator overhead,
- TLB pressure,
- pointer chasing,
- cache misses due to scattered memory,
- frequent system calls when fragmentation occurs.

Strategies to mitigate:

4.5.1 Use `std::vector` Instead of Many new

`std::vector` provides:

- contiguous storage,
- amortized constant-time push back,
- excellent spatial locality,
- fast iteration.

4.5.2 Small String Optimization (SSO)

Many C++ standard library implementations store short strings inline within the `std::string` object:

- avoids heap allocation for small strings,
- reduces fragmentation,
- improves cache locality,
- makes passing strings cheaper.

4.5.3 Small Buffer Optimization (SBO) for Custom Types

You can embed a fixed-size buffer inside your class to avoid allocations:

```
class SmallVec {  
    std::size_t size = 0;  
    std::size_t cap = 8;  
    int data[8]; // small buffer optimization  
};
```

4.5.4 Memory Pools and Custom Allocators

Memory pools allocate large blocks of memory and serve small allocations from them.

Benefits:

- lower fragmentation,
- faster allocation/deallocation,
- predictable performance,
- fewer system calls.

Example use:

```
std::pmr::monotonic_buffer_resource pool;  
std::pmr::vector<int> v(&pool);
```

Polymorphic memory resources (PMR) in C++17+ provide an excellent framework for high-performance allocation strategies.

4.6 Object Contiguity and Pointer Chasing

4.6.1 Why Contiguity Matters

Contiguous data reduces:

- cache misses,
- TLB misses,
- pointer following,
- unpredictable control flow.

4.6.2 Linked Lists: The Performance Villain

- Nodes allocated separately.
- Randomized memory layout.
- Next pointer requires dereference.
- Impossible to vectorize.

Even for $O(n)$ tasks, a linked list is orders of magnitude slower than a vector.

4.6.3 Better Alternatives

- `std::vector<T>`
- `boost::static_vector`
- array-based queues/rings
- ECS tables (entity-component systems)

4.7 False Sharing in Data Layout

False sharing occurs when two threads modify data within the same cache line.

Mitigation:

- use `alignas(64)` on thread-heavy structures,
- group per-thread state into arrays (each index belongs to one thread),
- pad structures to cache-line boundaries.

Example:

```
struct alignas(64) Counter {  
    std::atomic<long long> value;  
};
```

4.8 Case Study: Physics Update Loop

Consider updating positions based on velocities:

```
for (size_t i = 0; i < n; ++i) {  
    particles[i].x += particles[i].vx * dt;  
    particles[i].y += particles[i].vy * dt;  
    particles[i].z += particles[i].vz * dt;  
}
```

4.8.1 AoS Performance Characteristics

- Each particle is contiguous.
- Fields are grouped → good spatial locality.
- Compiler can autovectorize the loop.

4.8.2 SoA Version

```
for (size_t i = 0; i < n; ++i) {  
    x[i] += vx[i] * dt;  
    y[i] += vy[i] * dt;  
    z[i] += vz[i] * dt;  
}
```

Advantages:

- Perfect alignment for SIMD loads.

- No unused fields fetched.
- Faster for very large particle sets.

Exercise

Design two versions of a simple entity system:

1. Using AoS:

```
struct Entity { float x, y; int hp; };  
std::vector<Entity> entities;
```

2. Using SoA:

```
std::vector<float> posX, posY;  
std::vector<int> hp;
```

Tasks:

- Write functions that update only x and y.
- Profile both versions under heavy iteration.
- Compare cache behavior using a profiler (Perf/VTune).
- Explain differences in terms of locality and vectorization.

Chapter 5

Control Flow, Branch Prediction, and Branchless Code

Branches are essential to programming, but they create challenges for modern CPU pipelines. Efficient handling of control flow is critical for high-performance C++, especially in tight loops or data-processing routines. This chapter provides a deep exploration of branch behavior, prediction algorithms, speculative execution, branchless transformations, and practical techniques for writing predictable and efficient code.

5.1 Branch Prediction Basics

Modern CPUs are deeply pipelined machines—some pipelines exceed 15–25 stages (sometimes more in high-frequency cores). If the CPU had to wait for every branch to resolve before proceeding, performance would collapse. To mitigate this, CPUs use sophisticated branch predictors and speculative execution.

5.1.1 Why Branch Prediction Exists

When reaching a conditional branch, the CPU has two choices:

- The branch is taken (jump to another address).
- The branch is not taken (continue sequentially).

Waiting for this decision requires evaluating the condition, which can take many cycles. Instead, the CPU guesses the outcome and continues executing instructions from the predicted path.

If the prediction is correct:

- execution continues with no penalty.

If the prediction is wrong:

- the pipeline is flushed,
- speculative instructions are discarded,
- and execution restarts at the correct path.

A misprediction can cost 10–25 cycles or more.

5.1.2 Prediction-Friendly Patterns

Predictors learn patterns over time. Easily predictable branches include:

- Always taken (e.g., loop back edges).
- Always not taken.
- Monotonic conditions (e.g., $i < \text{limit}$).

- Simple alternation (some CPUs learn 1-bit and 2-bit patterns).
- Skewed distributions (e.g., condition true 99% of the time).

Loops are excellent for branch prediction because their pattern is stable: the branch is taken for many iterations and then not taken once.

5.1.3 Difficult-to-Predict Patterns

Hard-to-predict branches include:

- Data-dependent conditions on unpredictable input.
- Branches based on:
 - hash values,
 - cryptographic computations,
 - random numbers,
 - user-generated input,
 - complex predicates.
- Switch statements with many cases and uniform distribution.

Example:

```
if (hash(x) & 1) { /* ... */ }
```

Since the hash output is effectively random, prediction accuracy is low, and mispredictions accumulate.

5.2 Deep Dive: How Branch Predictors Work

Modern CPUs use multiple predictors:

5.2.1 Static Predictors

Basic rules applied before runtime learning:

- Backward branches (loops) → predict taken.
- Forward branches → predict not taken.

5.2.2 Dynamic Predictors

Dynamic predictors learn from history:

1-Bit Predictor Predicts the last outcome; flips on misprediction.

2-Bit Saturating Counter Better stability—requires two mispredictions to switch direction.

Global History Predictor Uses a global register storing recent branch outcomes.

Local History Predictor Tracks per-branch history.

Tournament Predictors Chooses the most accurate predictor based on recent performance. Used by modern x86-64 and ARM CPUs.

5.2.3 Branch Target Buffer (BTB)

Caches the destination addresses of taken branches to speed up fetching.

5.2.4 Speculative Execution

The CPU executes instructions along the predicted path before knowing if it is correct. This dramatically boosts performance but introduces side effects (e.g., Spectre).

5.3 The Cost of Misprediction

A mispredicted branch causes:

- pipeline flush,
- invalidation of speculative work,
- instruction cache disruption,
- wasted energy and cycles.

On modern Intel/AMD CPUs:

- Cost of misprediction: 15–25 cycles.

On tight loops, even a few mispredictions can dominate runtime.

5.4 Reducing Branches

Some branches are unavoidable. Others are unnecessary or can be restructured.

5.4.1 Fast-Path / Slow-Path Structure

```
if (likely(x > 0)) {  
    process_positive(x);  
}
```

```
} else {  
    process__nonpositive(x);  
}
```

likely and unlikely are hints:

- GCC/Clang: `__builtin_expect`
- MSVC: has `[[likely]]` and `[[unlikely]]` in C++20

Hints do not force the CPU—they reorganize code for better layout and I-cache locality.

5.4.2 Reordering Code

Sometimes branches can be written in a more predictable pattern.

Example:

```
if (rare_error_condition()) {  
    handle_error();  
}  
process_mainpath();
```

Better structure for performance:

```
if (!rare_error_condition()) {  
    process_mainpath();  
} else {  
    handle_error();  
}
```

5.4.3 Avoiding Unpredictable Branches in Hot Loops

```
for (int x : v) {  
    if (x == special) { ++count; }  
}
```

Better:

```
int special_val = special;  
for (int x : v) {  
    count += (x == special_val);  
}
```

The comparison becomes a predicated instruction, eliminating the branch entirely.

5.5 Branchless Techniques

Branchless code replaces conditionals with arithmetic, bitwise, or intrinsic operations that avoid control flow.

5.5.1 Clamp Example

Original:

```
int clamp(int x, int lo, int hi) {  
    if (x < lo) x = lo;  
    else if (x > hi) x = hi;  
    return x;  
}
```

Branchless version:

```
int clamp_branchless(int x, int lo, int hi) {  
    x = std::max(x, lo);  
    x = std::min(x, hi);  
    return x;  
}
```

This typically compiles to:

```
x = max(x, lo)  
x = min(x, hi)
```

No branches.

5.5.2 Conditional Select Instructions

Many CPUs support conditional moves:

- x86: `cmov`
- ARM: predicated instructions

Example:

```
int abs_branchless(int x) {  
    int mask = x >> 31; // all 1s for negative, 0 for positive  
    return (x ^ mask) - mask;  
}
```

No branch.

5.5.3 Multiplication Trick

Transforming:

```
if (cond) x = y;
```

to:

```
x = cond * y + (!cond) * x;
```

Works for integer/boolean types, though compilers often optimize even cleaner.

5.5.4 Lookup Table (LUT) Based Branch Removal

Instead of:

```
if (state == 0) a = f();  
else if (state == 1) a = g();  
else a = h();
```

Use:

```
using F = int(*)();  
static F table[] = {f, g, h};  
a = table[state]();
```

Warning: This replaces a branch with an indirect jump, which is sometimes slower if branch misprediction is rare.

Important: measure both ways.

5.5.5 SIMD Predication

Vector instructions perform operations element-wise without branches.

Example: count positive values

```
for (size_t i = 0; i < n; ++i)
    count += (v[i] > 0);
```

Auto-vectorization produces:

```
cmp > 0 (vector)
mask
count bits
```

Branch-free and extremely fast.

5.6 When Branchless Code Is Not Better

Branchless code is not always superior. Pitfalls include:

- Higher register pressure.
- More instructions than the branching version.
- Extra arithmetic or bitwise operations.
- Difficult-to-read logic.
- Worse performance if branches are highly predictable.

Rule of thumb:

Only convert unpredictable branches into branchless code.

5.7 Practical Patterns for Predictable Branches

5.7.1 Sort Data to Improve Predictability

Sorted arrays often perform significantly faster than random arrays in branch-heavy code.

Example:

```
for (int x : sorted_vector) {  
    if (x > threshold) count++;  
}
```

Predictor sees a long run of false, then a long run of true. Perfect for prediction.

5.7.2 Group Data by Likelihood

Instead of mixing hot and cold cases:

```
for (auto& e : events) {  
    if (e.type == HOT) handle_hot(e);  
    else handle_cold(e);  
}
```

Group before processing:

```
std::partition(events.begin(), events.end(),  
               [](auto& e){ return e.type == HOT; });  
  
for (auto& e : events) handle(e);
```

5.7.3 Use Sentinel Values

Instead of:

```
if (ptr) use(ptr);
```

consider initializing `ptr` to a valid sentinel object to remove null checks.

5.8 Case Study: Counting Elements Above a Threshold

Original version (branching):

```
int count_gt_branching(const std::vector<int>& v, int thr) {  
    int count = 0;  
    for (int x : v)  
        if (x > thr) ++count;  
    return count;  
}
```

Branchless version:

```
int count_gt_branchless(const std::vector<int>& v, int thr) {  
    int count = 0;  
    for (int x : v)  
        count += (x > thr);  
    return count;  
}
```

Performance Dependence on Data Distribution

- Sorted data → branching version is nearly perfect.
- Random data → branchless version wins dramatically.

Exercise

Implement:

```
int count_gt(const std::vector<int>& v, int threshold);
```

Tasks:

1. Benchmark on sorted data.
2. Benchmark on reverse-sorted data.
3. Benchmark on random data.

Questions to answer:

- Which dataset yields best performance for branching?
- Which dataset yields best performance for branchless?
- How many mispredictions per iteration occur (use perf stat)?
- How does ILP differ between versions?

Chapter 6

Cost of Abstractions: Inlining, Functions, and Virtual Calls

Abstractions are essential to writing maintainable and expressive C++ code. They allow developers to structure large systems in manageable pieces. However, abstractions come with performance implications—sometimes small, sometimes dramatic—depending on how the compiler lowers them into machine code.

This chapter explores the true performance cost of various C++ abstractions, the mechanics of function calls, compiler inlining strategies, dynamic vs static polymorphism, devirtualization opportunities, and the concept of zero-cost abstractions that Modern C++ strives to achieve.

6.1 Function Calls and Inlining

Every function call in C++ has a non-zero overhead, although it is typically small. Understanding this overhead is crucial when writing tight loops or performance-critical inner kernels.

6.1.1 The Mechanics of a Function Call

Under most ABIs (e.g., System V x86-64 ABI), a function call involves:

- placing arguments in registers or on the stack,
- potentially spilling registers to maintain calling convention,
- jumping to the callee (an indirect branch),
- executing the callee's prologue,
- executing the function body,
- restoring registers,
- returning to the caller (another indirect branch).

Costs include:

- control-flow disruption (a jump),
- potential I-cache pressure (callee not near caller),
- branch predictor interaction,
- stack manipulation,
- register spills and reloads.

For most code, this cost is irrelevant. But in tight loops—executed millions or billions of times—it matters.

6.1.2 Inlining: Eliminating Call Overhead

Inlining replaces a function call with the function's body.

Example:

```
inline int add(int a, int b) {  
    return a + b;  
}
```

Inlining benefits:

- Removes call/return overhead.
- Enables constant propagation.
- Allows loop unrolling.
- Facilitates auto-vectorization.
- Removes stack usage.
- Exposes optimization opportunities across boundaries.

6.1.3 What inline Really Means

The inline keyword:

- is primarily an ODR (One Definition Rule) mechanism,
- does not force inlining,
- gives permission for multiple identical function definitions in headers.

Modern compilers use sophisticated heuristics to decide when to inline.

6.1.4 Inlining Heuristics

Compilers inline functions based on:

- function size,
- frequency of call sites,
- impact on code size,
- expected gains from constant folding,
- ability to optimize surrounding code,
- profile data (with PGO),
- LTO-based whole-program analysis.

Compiler flags such as `-O3` increase inlining aggressiveness. PGO greatly improves inlining decisions.

6.1.5 Dangers of Over-Inlining

Excessive inlining can degrade performance:

- Code bloat increases I-cache misses.
- Larger working sets reduce effectiveness of L1/L2 caches.
- Less predictable instruction layout harms branch prediction.

A frequently overlooked truth:

Inlining increases throughput only if it reduces instruction count or exposes new optimizations.

Otherwise, it may slow the program.

6.2 Virtual Functions

Dynamic polymorphism is a powerful tool but comes with performance implications because virtual calls are indirect.

6.2.1 Mechanics of a Virtual Call

A virtual function call involves:

1. Loading the vtable pointer from the object.
2. Looking up the function pointer in the vtable.
3. Performing an indirect jump to the method.

Costs compared to a direct call:

- Indirection $\rightarrow$ harder for branch predictor.
- Prevents inlining in most cases.
- Impedes optimization across call boundaries.

Latency: 5–30 cycles depending on CPU microarchitecture.

6.2.2 Devirtualization

Compilers can sometimes optimize virtual calls into direct calls when:

- The dynamic type is known at compile time.
- Only one derived type is possible in context.

- LTO and PGO provide class usage information.

Example:

```
Base* p = new Derived;  
p->do_work(); // compilers often devirtualize this
```

If p's dynamic type is provably Derived, the compiler replaces the virtual call with a non-virtual call.

6.2.3 When Virtual Functions Hurt Performance

Virtual functions degrade runtime when:

- The call occurs in a tight loop.
- The dispatch pattern is unpredictable.
- The call prevents vectorization.
- Cache misses occur due to scattered class layouts.

Common in:

- scene graphs,
- entity-component systems (except SoA designs),
- AST walkers,
- dynamic dispatch-heavy subsystems.

6.3 Alternatives to Virtual Functions

6.3.1 Templates and Static Polymorphism

Using templates (or CRTP) allows polymorphism resolved at compile time. Dynamic branching disappears; the optimizer sees the actual code and can inline aggressively.

```
template <class T>
void process(T& obj) {
    obj.do_work();
}
```

6.3.2 CRTP (Curiously Recurring Template Pattern)

```
template <typename Derived>
struct BaseCRTP {
    void do_work() {
        static_cast<Derived*>(this)->do_work_impl();
    }
};
```

Advantages:

- Zero runtime overhead.
- Full inlining.
- No vtable usage.

6.3.3 Variant-Based Polymorphism

`std::variant` and `std::visit` provide type-safe alternatives to virtual dispatch.

- No vtable.
- Decisions resolved via static dispatch in `std::visit`.
- Often faster than virtual dispatch for small variant sets.

6.3.4 Function Objects and Lambdas

Lambdas:

- Inline naturally.
- Enable customization without virtual dispatch.
- Are optimized aggressively under `-O2/-O3`.

6.4 Zero-Cost Abstractions

C++ aims for abstractions that compile to the same (or nearly the same) machine code as hand-written low-level code.

Examples:

- Range-based for loops: compile to simple pointer increments.
- `std::span`: wrapper around `pointer+size`, trivial to inline.
- Iterators: compile away when optimized.
- Templates: allow compile-time polymorphism with no overhead.

6.4.1 Cost-Free RAI

RAII objects (locks, file handles, wrappers) often inline destructor calls, leaving no function-call overhead.

6.4.2 Views Instead of Copies

Using:

- `std::string_view`
- `std::span<T>`
- `gsl::span`

reduces copying and allocation dramatically.

6.4.3 Zero-Cost Lambdas

Lambdas with no captures convert to function pointers. Capturing lambdas often inline fully under `-O2`.

6.4.4 The Importance of Transparency

To enable zero-cost abstractions:

- Types must be trivially small.
- All operations must be inlineable.
- Templates must expose all logic to the optimizer.

Opaque constructs defeat optimization.

6.5 Type Erasure and Its Costs

Type erasure (e.g., `std::function`, custom wrappers) hides concrete types behind uniform interfaces.

6.5.1 Costs of `std::function`

- Allocations (if target does not fit in small buffer).
- Indirect function calls (via vtable-like dispatch).
- Prevents inlining.

6.5.2 Better Alternatives

- Inlineable functors (templated callables).
- `auto` parameters with generic lambdas.
- Custom small function wrappers using SBO.

6.6 Code Bloat and Its Performance Impact

Inlining and templates can explode binary size.

Consequences of code bloat:

- Increased instruction cache (I-cache) misses.
- Lower effective throughput.
- Higher memory usage.
- Slower cold-start performance.
- Poorer branch predictor accuracy.

Techniques to mitigate bloat:

- Avoid over-templating.

- Limit excessive inlining.
- Move rarely used code to separate translation units.
- Use PGO + LTO for smart global optimization.

6.7 Static vs Dynamic Polymorphism: Quantifying Costs

Static Polymorphism (Templates)

- Resolved at compile time.
- Zero dispatch overhead.
- Potentially infinite inlining.
- Potential code bloat.

Dynamic Polymorphism (Virtual Functions)

- Dispatch cost: 5–30 cycles.
- No inlining in hot loops (usually).
- Smaller binary size (common implementation shared).
- Good for dynamic dispatching.

Rule of thumb:

Use virtual functions when flexibility matters. Use templates when performance matters.

6.8 Case Study: Three Implementations of a Hot Loop

Task: multiply each element of an array by a constant factor.

6.8.1 Raw Pointer Loop

```
void scale_raw(float* a, size_t n, float factor) {  
    for (size_t i = 0; i < n; ++i)  
        a[i] *= factor;  
}
```

Fastest; compiles to tight SIMD loop.

6.8.2 Helper Function

```
inline void scale_one(float& x, float f) { x *= f; }  
  
void scale_func(float* a, size_t n, float factor) {  
    for (size_t i = 0; i < n; ++i)  
        scale_one(a[i], factor);  
}
```

Depending on heuristics, compiler may inline `scale_one`, producing identical assembly.

6.8.3 Template Functor

```
struct Scaler {  
    float f;  
    void operator()(float& x) const { x *= f; }  
};  
  
template <class Op>
```

```
void apply(float* a, size_t n, Op op) {  
    for (size_t i = 0; i < n; ++i)  
        op(a[i]);  
}
```

Under -O2, compiler fully inlines everything; assembly matches the raw loop.

Result:

Template-based code and raw loops produce identical machine code. Thus: zero-cost abstraction.

Exercise

Write three versions of a small computation:

1. Raw loop with pointers
2. Function that calls a small helper function in a tight loop
3. Template-based version likely to be inlined

Tasks:

- Compile with -O0, -O2, and -O3.
- Compare generated assembly using Compiler Explorer.
- Measure performance on large arrays.
- Determine which versions are inlined and why.
- Explain how aliasing, inlining, and vectorization differ between cases.

Chapter 7

Move Semantics, Value Categories, and Copies

Modern C++ (C++11 and beyond) fundamentally changed how objects are created, transferred, and optimized. Move semantics and value-category rules are at the heart of these improvements. They allow C++ to eliminate unnecessary copies, reduce allocations, and optimize the performance of complex programs without sacrificing expressiveness.

This chapter provides a deep exploration of how copying works, why moves are cheaper, how compilers optimize returns, how views eliminate unnecessary transfers, and how value categories affect overload resolution and performance.

7.1 The Cost of Copies

Copying objects may seem harmless, but depending on the type, the cost can be substantial. The biggest offenders are:

- large `std::string` objects,
- containers such as `std::vector`, `std::deque`, `std::map`,
- heavy classes with dynamic memory,
- large aggregates (structs with many double-precision fields),
- classes with resource ownership (file handles, sockets, mutexes).

A naive copy of a large `std::vector<T>` requires:

- allocating new memory,
- copying all elements (possibly millions),
- destroying old data (if temporary),
- updating source/destination metadata.

This can cost microseconds or milliseconds inside tight loops—catastrophic for performance.

Modern C++ helps you avoid these unnecessary copies by providing:

- move semantics: transfer ownership cheaply,
- return value optimization (RVO): eliminate copies entirely,
- views: read-only or non-owning references to data,
- reference passing: avoid ownership transfer when possible.

7.2 Value Categories and Why They Matter

Understanding move semantics requires understanding C++ value categories. C++ has the following:

lvalue has a persistent location (named objects).

xvalue “expiring value”—an object eligible for moving.

prvalue “pure rvalue”—a temporary without location.

C++ groups them into:

- glvalue = lvalue or xvalue,
- rvalue = xvalue or prvalue.

Relevance to optimization:

- lvalues prefer copying,
- rvalues prefer moving,
- xvalues often trigger moves,
- prvalues often trigger RVO.

Example:

```
std::string s = "hello";  
std::string t = s;      // lvalue -> copy  
std::string u = std::move(s); // xvalue -> move
```

7.3 Move Semantics in Practice

Move semantics transfer resources instead of duplicating them.

Classic example:

```
std::vector<int> make_large_vector() {  
    std::vector<int> v(1'000'000, 42);  
    return v; // moved or elided  
}
```

If the return value is moved:

- No data is copied.
- Only metadata (pointer, size, capacity) is moved.
- The source vector becomes empty.

This is extremely cheap.

7.3.1 Move Constructor and Move Assignment

A typical move constructor of `std::vector`:

```
pointer p    -> moved  
size s      -> moved  
capacity c  -> moved  
source = {nullptr, 0, 0}
```

No elements are copied or moved individually.

7.3.2 When Does Moving Happen?

- when binding to an rvalue,
- when calling `std::move`,
- inside containers during reallocation,
- when returning local objects under NRVO or RVO rules,
- when function parameters are passed by value and the caller passes an rvalue.

7.3.3 `std::move` Is Not a Move

`std::move` casts an lvalue to an xvalue.

It does not move anything by itself—it only makes the object eligible for moving.

```
auto v = make_large_vector(); // RVO
consume(std::move(v));        // calls move ctor
```

Never `std::move` a variable twice—moved-from state is valid but unspecified.

7.4 RVO, NRVO, and Copy Elision

Return Value Optimization removes copies entirely. Instead of copying or moving, the compiler constructs the return value directly in the caller's storage.

Types of elision:

- RVO: when returning a temporary.
- NRVO: when returning a named local variable.

- Mandatory elision (C++17): many temporary constructions do not create lifetime-bound objects.

Example:

```
std::vector<int> make() {  
    return std::vector<int>(1000); // RVO: no move, no copy  
}
```

Even:

```
return v;
```

where `v` is a local non-volatile variable allows NRVO.

7.4.1 Why RVO Matters

It eliminates:

- copy,
- move,
- temporary lifetime overhead,
- destructor calls.

RVO is one of the biggest contributors to C++ performance.

7.5 Perfect Forwarding and References

Templates can preserve value categories using forwarding references:

```
template <class T>
void wrapper(T&& x) {
    process(std::forward<T>(x));
}
```

`std::forward` ensures:

- rvalues remain rvalues (moved),
- lvalues remain lvalues (copied or referenced),
- overload resolution remains optimal.

This makes generic code efficient without repeating overloads manually.

7.6 Minimizing Copies with Views

Sometimes you do not need ownership. You only need to look at data.

For strings:

```
void process(std::string_view sv);
```

- No allocation.
- No copying.
- Constant-time construction.
- Useful for parsing, logging, validation.

For contiguous ranges:

```
void operate(std::span<int> sp);
```

`std::span<T>` provides:

- pointer + size,
- no copy,
- no ownership transfer,
- lightweight slicing.

7.6.1 Other Useful Views

- `std::basic_string_view`
- `std::ranges::subrange`
- `gsl::span`
- `mdspan` (C++23)

Using views avoids:

- heap allocations,
- large object copies,
- unnecessary reference counting.

7.7 Optimizing Function Interfaces

Guidelines:

- Pass large read-only objects as `const T&` or `std::string_view`.
- Pass large mutable objects as `T&`.
- Pass small trivially copyable objects (`int`, `double`, small structs) by value.
- Return by value if using RVO (preferred).
- Avoid passing complex objects by value unless move-only or small.

Example of an interface improving performance:

Before:

```
void analyze(std::string data);
```

After:

```
void analyze(std::string_view data);
```

Massive reduction in allocations and copies.

7.8 Cheaply Movable Types

For user-defined types, ensure movable semantics are cheap:

- Move constructor should steal resources, not allocate.
- Avoid heavy copying in move constructors.

- Ensure invariant: moved-from objects remain valid.

Example:

```
struct Buffer {
    size_t size = 0;
    char* data = nullptr;

    Buffer(size_t n) : size(n), data(new char[n]) { }

    // Move constructor
    Buffer(Buffer&& other) noexcept
        : size(other.size), data(other.data)
    {
        other.size = 0;
        other.data = nullptr;
    }

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

7.9 Common Pitfalls That Prevent Moves

7.9.1 Accidental Copies

```
void foo(std::vector<int> v); // takes by value
foo(myvec);                  // copies!
```

Fix:

```
void foo(const std::vector<int>& v); // no copy
```

7.9.2 Using Moved-From Objects Unexpectedly

```
auto s2 = std::move(s1);  
use(s1); // allowed, but s1 is empty
```

7.9.3 Overloading Mistakes

If you write:

```
void set(std::string s);  
void set(const std::string& s);
```

Passing an rvalue selects the `std::string` by value overload, causing an extra copy.

7.10 Case Study: Passing by Value vs View

Suppose you frequently call:

```
void process(std::string s);
```

And the call occurs thousands of times per second with 4 KB strings.

7.10.1 Problems

- Each call copies 4 KB.
- Significant memory bandwidth usage.
- Cache pollution.
- Increased heap pressure.

7.10.2 Refactor

```
void process(std::string_view s);
```

Now:

- No copies.
- No heap allocations.
- Faster calling convention.

In real systems, this often results in $3\times$ – $20\times$ speedups.

Exercise

Take a function in your codebase that currently takes a `std::string` by value but only reads it.

Tasks:

- Change the signature to accept `std::string_view`.
- Measure total runtime across thousands of calls.
- Measure total allocations using a tool such as:
 - `valgrind --tool=massif`
 - `perf stat -e malloc,free`
 - `heaptrack`
- Compare:

- allocation count,
- CPU time,
- memory usage,
- peak resident set size,
- cache misses (L1/L2/L3).

Chapter 8

Dynamic Memory, Allocators, and Ownership

Dynamic memory is one of the most expensive operations in many real-world applications. Although modern hardware is fast, the act of requesting memory from the operating system or global allocator introduces synchronization, metadata management, fragmentation, and unpredictable latencies. Understanding these costs—and how to minimize them—is essential for writing high-performance C++.

This chapter explores allocation overheads, standard and custom allocator models, memory pools, NUMA characteristics, PMR (Polymorphic Memory Resources), and modern ownership patterns using smart pointers and RAII.

8.1 The Hidden Cost of Allocation

On the surface, `new`, `delete`, `malloc`, and `free` appear simple. But under the hood, they are complex operations involving:

- Synchronization across threads,
- Interaction with OS-level memory management,
- Maintenance of metadata structures (free lists, trees, bins),
- Page allocation and deallocation,
- Cache invalidation,
- TLB (Translation Lookaside Buffer) pressure.

Even a single heap allocation can cost dozens to hundreds of CPU cycles—far more than most instructions.

8.1.1 Costs Inside a Memory Allocator

Common steps include:

- Taking a global or per-heap lock.
- Searching for a block of sufficient size (first-fit, best-fit, segregated lists).
- Splitting blocks when necessary.
- Updating the data structures representing free blocks.
- Zeroing memory (sometimes required by OS).

On multi-threaded systems, this is further complicated:

- Frequent contention on allocator locks.
- False sharing on allocator metadata.
- Unpredictable latency for allocation/free operations.

8.1.2 Common Allocation Anti-Patterns

- Allocating inside tight loops.
- Repeatedly resizing containers without `reserve()`.
- Creating large temporary objects repeatedly.
- Using `std::string` copies in log loops.
- Allocating per-element objects instead of contiguous blocks.

8.1.3 Practical Advice

- Reserve container capacity: `reserve()` eliminates most future reallocations.
- Use stack allocation: faster, no fragmentation, no locks.
- Reuse buffers: keep scratch memory around instead of freeing and reallocating.
- Move semantics: ensure your types are cheaply movable, reducing allocations during container growth.
- Avoid heap allocations in per-frame or per-request hot paths: preallocate instead.

8.2 Memory Fragmentation and Growth

Fragmentation occurs when memory is allocated and freed in patterns that leave many small gaps. Over time:

- heaps grow,
- allocation latency increases,

- cache locality degrades,
- system memory usage inflates.

Strategies to combat fragmentation:

- Use arenas/monotonic allocators: allocate many objects together, free them at once.
- Keep objects of similar size together (pool allocators).
- Avoid alternating large and small object allocations.
- Use vector-like contiguous containers instead of fragmented node-based structures.

8.3 Allocator Design in C++

The C++ allocator model allows customizing how containers obtain memory.

A standard allocator must provide:

- `allocate(n)`: allocate memory for `n` objects,
- `deallocate(p, n)`: return memory to the allocator,
- propagate traits such as rebinding for other types.

Modern C++17 introduced PMR (Polymorphic Memory Resources), which simplify the use of custom allocators significantly.

8.3.1 Types of Allocators

Monotonic Resource A simple allocator that allocates linearly from a buffer and never frees until the buffer is destroyed. Excellent for temporary object graphs, request-scoped work, parsing, ASTs, etc.

- Extremely fast.
- Zero fragmentation.
- Ideal for one-shot allocations.

Pool Allocators Designed for fixed-size objects. Commonly used for:

- Nodes of linked lists,
- Game entities,
- Small objects created and destroyed often.

Stack Allocators Allocate from a fixed buffer using bump-pointer techniques. Useful for temporary scratch buffers.

Custom Allocators For specific workload patterns: slab allocators, region-based allocators, lock-free allocators.

8.4 Monotonic Buffer Resource

`std::pmr::monotonic_buffer_resource` is incredibly useful for high-performance code:

- Allocations come from a contiguous region.
- No deallocation occurs until the resource is destroyed.

- Amortized constant-time allocation.
- Perfect for parsing tasks, AST creation, and large batch operations.

Example:

```
std::pmr::monotonic_buffer_resource pool;
std::pmr::vector<std::string> v(&pool);

v.emplace_back("hello");
v.emplace_back("world");
// no deallocate until pool goes out of scope
```

8.5 Memory Pools and Free-Lists

Memory pools preallocate many fixed-size blocks and serve them to callers:

- Faster than using the general-purpose allocator.
- Avoids fragmentation because all blocks are identical.
- Often implemented using linked free lists.

Example: a pool for 64-byte objects.

```
struct Node64 {
    Node64* next;
    char data[64 - sizeof(Node64*)];
};
```

Allocating from such a pool is essentially:

pop from free list

Deallocating is:

push to free list

Both operations take $O(1)$ time.

8.6 NUMA Awareness (Advanced)

On NUMA (Non-Uniform Memory Access) systems, memory is divided into nodes, each local to a subset of CPU cores.

Access patterns:

- Local memory access: fast.
- Remote memory access: slow (up to $2\times$ – $4\times$ slower).

High-performance systems should consider:

- NUMA-local allocations (via OS APIs),
- pinning threads to CPUs,
- avoiding memory migration,
- using NUMA-aware allocators.

This becomes essential in:

- high-frequency trading,
- scientific computing,
- large servers with many sockets,
- in-memory datastores.

8.7 Ownership and RAII

RAII (Resource Acquisition Is Initialization) guarantees that resources are:

- acquired on construction,
- released on destruction.

This applies to:

- memory,
- file handles,
- sockets,
- mutex locks,
- GPU resources,
- thread handles.

8.7.1 Smart Pointers

`std::unique_ptr` Exclusive ownership. Zero overhead compared to raw pointers. Perfect for owning dynamic memory.

`std::shared_ptr` Reference-counted ownership. More expensive due to:

- atomic increments/decrements,
- separate control blocks,
- potential false sharing.

Use sparingly.

`std::weak_ptr` Non-owning reference to `std::shared_ptr`. Avoids cycles.

8.7.2 RAII Wrappers and Performance

RAII is usually zero-cost or nearly zero-cost. For example:

- File RAII wrapper only calls `open()` in constructor and `close()` in destructor.
- Mutex lock wrappers (e.g., `std::lock_guard`) usually inline fully.
- Memory buffers in RAII remain fast because they allocate once and deallocate once.

RAII helps performance by avoiding leaks and ensuring deterministic cleanup (e.g., flushing buffers exactly once).

8.8 Avoiding Overuse of `shared_ptr`

`shared_ptr` should not be used as a default ownership mechanism.

Costs include:

- atomic refcount operations 20–80 cycles,
- extra memory allocation for control block,
- pointer aliasing inhibits compiler optimizations,
- increased risk of memory fragmentation.

Use `unique_ptr` unless shared ownership is truly required.

8.9 Case Study: Eliminating Allocations in a Hot Loop

Suppose you have:

```
for (int i = 0; i < N; ++i) {  
    std::string tmp;  
    tmp = format_message(i); // allocates each iteration  
    process(tmp);  
}
```

Problems:

- tmp grows repeatedly,
- allocation per iteration,
- potential heap fragmentation,
- excessive copying.

Refactor:

```
std::string tmp;  
tmp.reserve(max_len);  
  
for (int i = 0; i < N; ++i) {  
    tmp.clear();  
    append_formatted(tmp, i); // no new allocations  
    process(tmp);  
}
```

Or use a monotonic pool:

```
std::pmr::monotonic_buffer_resource pool;
std::pmr::string tmp{&pool};

for (int i = 0; i < N; ++i) {
    tmp = format(i); // very cheap
}
```

Performance often improves by $3\times$ – $50\times$ depending on workload.

Exercise

Identify a loop in your project that performs memory allocations on every iteration. Refactor it to:

1. Pre-reserve container capacity, or
2. Reuse a scratch buffer allocated once outside the loop, or
3. Replace dynamic allocations with a monotonic buffer resource.

Measure:

- total runtime,
- number of allocations (use valgrind, heaptrack, or perf),
- L1/L2 cache misses,
- RSS memory usage,
- fragmentation over time.

Chapter 9

SIMD and Auto-Vectorization

Modern CPUs contain specialized hardware units capable of performing the same operation on multiple elements at once. This is known as SIMD—Single Instruction, Multiple Data. SIMD is one of the most powerful ways to increase throughput in numeric and data-processing workloads, and modern C++ compilers can leverage SIMD automatically (auto-vectorization) or through explicit intrinsics and high-level APIs. This chapter provides a comprehensive overview of SIMD architecture, automatic vectorization by compilers, manual vectorization for performance-critical kernels, and modern C++ abstractions that make SIMD more accessible.

9.1 What Is SIMD?

SIMD instructions apply a single operation across multiple data elements concurrently. Examples:

- Add eight 32-bit floats in parallel.
- Multiply four 64-bit doubles at once.

- Compare 16 bytes simultaneously.

This works because SIMD registers are wide:

x86 SSE 128-bit registers (4 floats or 2 doubles)

x86 AVX2 256-bit registers (8 floats or 4 doubles)

x86 AVX-512 512-bit registers (16 floats or 8 doubles)

ARM NEON 128-bit registers (similar to SSE)

ARM SVE Scalable 128-2048-bit vectors (length determined at runtime)

RISC-V V spec Scalable vectors with flexible widths and lengths

These architectures enable massive throughput increases for arithmetic, comparison, filtering, and data transformation.

9.1.1 SIMD Across Architectures

x86-64 The primary SIMD extensions are:

- SSE / SSE2 / SSE3 / SSSE3 / SSE4
- AVX / AVX2
- AVX-512 (various subsets)

ARM64 NEON is available on all performance-oriented ARM CPUs.

ARM SVE / SVE2 Scalable vectors (length unknown at compile time) enabling portable vectorization.

RISC-V V Generalized scalable vector architecture with flexible register lengths.

9.2 Why SIMD Matters

Using SIMD, the CPU processes:

- 4–16× more elements per instruction,
- with nearly the same energy cost as a scalar operation,
- with high instruction-level parallelism (ILP),
- with reduced loop overhead,
- and improved throughput.

For large arrays, the performance gains can be massive (5×–30×).

9.3 Enabling Auto-Vectorization

Auto-vectorization is when the compiler transforms scalar loops into SIMD instructions automatically. This requires:

- A regular loop structure.
- No data dependencies across iterations.
- Alignable and non-aliasing pointers.
- Predictable memory access patterns.

Example scalar loop:

```
void scale(float* data, std::size_t n, float factor) {  
    for (std::size_t i = 0; i < n; ++i) {  
        data[i] *= factor;  
    }  
}
```

Compilers typically generate AVX/AVX2 vector instructions under -O3.

9.3.1 When Auto-Vectorization Fails

Cases that prevent vectorization include:

- Function calls inside the loop.
- Unknown aliasing between arrays.
- Complex control flow.
- Pointer arithmetic that the compiler cannot reason about.
- Early exits (e.g., break, return).

9.3.2 Fixing Alias Problems

Use restrict (or `__restrict` for GCC/Clang):

```
void add(float* __restrict a,  
         float* __restrict b,  
         float* __restrict c,  
         std::size_t n)  
{  
    for (std::size_t i = 0; i < n; ++i)  
        c[i] = a[i] + b[i];  
}
```

The compiler can now safely assume a, b, and c do not overlap.

9.3.3 Dealing With Data Dependencies

Loops like:

```
for (size_t i = 1; i < n; ++i)
    a[i] += a[i - 1];
```

cannot be vectorized because each iteration depends on the previous.

9.4 Loop Transformations for Vectorization

To help the compiler, transform loops into simpler, predictable forms.

9.4.1 Hoist Invariants

BAD:

```
for (size_t i = 0; i < n; ++i)
    a[i] *= expensive(i);
```

BETTER:

```
auto c = expensive_constant;
for (size_t i = 0; i < n; ++i)
    a[i] *= c;
```

9.4.2 Remove Conditionals from Hot Loops

BAD:

```
for (size_t i = 0; i < n; ++i)
    if (mask[i]) a[i] += b[i];
```

BETTER (branchless):

```
for (size_t i = 0; i < n; ++i)
    a[i] += b[i] * mask[i];
```

9.4.3 Align Data

Alignment helps vector loads avoid slow unaligned paths.

```
alignas(32) float data[1024];
```

9.5 Explicit SIMD via Intrinsics

For maximum control—particularly in mathematical kernels—use SIMD intrinsics.

x86 AVX Example

```
#include <immintrin.h>

void scale_avx(float* data, size_t n, float f) {
    __m256 factor = __mm256_set1_ps(f);

    size_t i = 0;
    for (; i + 8 <= n; i += 8) {
        __m256 v = __mm256_loadu_ps(&data[i]);
        v = __mm256_mul_ps(v, factor);
        __mm256_storeu_ps(&data[i], v);
    }

    for (; i < n; ++i)
        data[i] *= f; // remainder cleanup
}
```

This version multiplies 8 floats at once.

ARM NEON Example

```
#include <arm_neon.h>

void scale_neon(float* data, size_t n, float f) {
    float32x4_t factor = vdupq_n_f32(f);

    size_t i = 0;
    for (; i + 4 <= n; i += 4) {
        float32x4_t v = vld1q_f32(&data[i]);
        v = vmulq_f32(v, factor);
        vst1q_f32(&data[i], v);
    }
}
```

9.5.1 When to Use Intrinsics

Use intrinsics when:

- auto-vectorization fails,
- maximum throughput is required,
- operations must be precise (e.g., mask loads, gathers),
- instructions must use specific hardware capabilities (FMA, AVX-512).

9.6 High-Level SIMD Abstractions

Several libraries provide safe, portable SIMD abstractions:

- `std::experimental::simd` (C++23 vendor extensions)
- `Vc` library (SIMD vectors)

- Eigen (SIMD math)
- xtensor (vectorized operations)
- Boost.SIMD (SIMD wrappers)

Example using `std::experimental::simd`:

```
#include <experimental/simd>

void scale_simd(float* data, size_t n, float f) {
    namespace simd = std::experimental;

    using vec = simd::native_simd<float>;
    vec factor(f);

    size_t i = 0;
    for (; i + vec::size() <= n; i += vec::size()) {
        vec v = simd::load<vec>(&data[i]);
        v *= factor;
        v.copy_to(&data[i], simd::element_aligned);
    }
}
```

Provides portability across platforms.

9.7 Execution Policies and Vectorization

C++17 introduces parallel and unsequenced execution policies:

```
std::for_each(std::execution::par_unseq,
              v.begin(), v.end(),
              [](float& x) { x *= 1.5f; });
```

`par_unseq` encourages:

- multi-threading,
- vectorization,
- loop unrolling.

Compilers map this onto SIMD + thread parallelism where possible.

9.8 Gather, Scatter, and Masked Operations

Some SIMD architectures (AVX-512, SVE, RISC-V V) support:

- gather: load non-contiguous memory into a vector,
- scatter: write vector elements to scattered addresses,
- masked operations: selectively apply operations.

Example (AVX-512 mask):

```
__m512 a = ...  
__mmask16 mask = ...  
__m512 b = __mm512_mask_mul_ps(a, mask, a, factor);
```

This applies multiplication only to elements whose mask bits are set.

9.9 When SIMD Fails or Underperforms

SIMD is not always beneficial:

- Too-small data sets (overhead dominates).

- Data dependencies prevent vectorization.
- Misaligned or scattered memory.
- Branch-heavy code.
- Complex control flow.
- Memory bandwidth bottlenecks.

For small arrays (length < 16), scalar code may outperform SIMD.

9.10 Case Study: Scalar vs SIMD Loop

Scalar:

```
for (size_t i = 0; i < n; ++i)
    sum += a[i] * b[i];
```

SIMD (AVX2):

```
__m256 acc = _mm256_setzero_ps();

for (size_t i = 0; i < n; i += 8) {
    __m256 va = _mm256_loadu_ps(&a[i]);
    __m256 vb = _mm256_loadu_ps(&b[i]);
    acc = _mm256_fmadd_ps(va, vb, acc);
}
```

Performance improvement: often $5\times$ – $20\times$ on modern CPUs.

Exercise

Task 1: Basic Vectorization Check

Compile the following with and without -O3:

```
for (size_t i = 0; i < n; ++i)
    data[i] *= factor;
```

Check assembly: does the compiler emit SIMD instructions?

Task 2: Help the Compiler

Rewrite the loop to ensure:

- data is aligned,
- no pointer aliasing,
- no branches inside the loop.

Recompile and check again.

Task 3: Manual SIMD

Write a version using AVX2 intrinsics and compare:

- execution time,
- CPU cycles,
- L1/L2 bandwidth usage,
- instruction count.

Chapter 10

Concurrency, Synchronization, and False Sharing

Modern multi-core CPUs provide tremendous parallel computational power, but achieving high performance in concurrent C++ programs requires a deep understanding of synchronization, memory ordering, contention behavior, and subtle microarchitectural effects such as false sharing. This chapter covers essential concurrency concepts for writing efficient, scalable, and correct multi-threaded applications in Modern C++.

10.1 The Cost of Synchronization

Synchronization primitives—mutexes, atomics, and memory fences—are expensive operations. Although they enforce correctness in concurrent programs, they often become scalability bottlenecks.

10.1.1 Why Synchronization Is Expensive

Each synchronization event may cause:

- Cache line transfers: moving a cache line from one core's L1/L2 cache to another (via MESI/MOESI protocols).
- Pipeline stalls: instructions dependent on synchronization slow down the pipeline.
- Memory fencing: prevents reordering, limiting instruction-level parallelism (ILP).
- OS involvement: blocking mutexes may involve kernel scheduling.

The biggest cost is contention: when multiple threads access the same memory location.

10.1.2 Avoid Synchronization in Hot Paths

Use synchronization only when necessary. Prefer:

- Thread-local state: each thread owns its own data.
- Per-core data: similar to thread-local, but aligned to cache lines.
- Immutable state: share read-only structures freely.
- Bulk updates: aggregate results from threads after the main loop, not during computation.

10.1.3 Coarse-Grained vs Fine-Grained Locks

Fine-grained locks:

- Allow more concurrency,
- But create more locking events and complexity.

Coarse-grained locks:

- Simplify design,
- Often outperform fine-grained designs due to fewer synchronization events.

When in doubt, begin with coarse-grained locking, profile, then refine.

10.2 CPU Cache Coherence and Synchronization Costs

Understanding the hardware helps explain why synchronization is expensive.

Modern CPUs maintain cache coherence using protocols such as:

- MESI (Modified, Exclusive, Shared, Invalid),
- MOESI (adds Owned state).

When a thread writes to a shared variable:

1. Its cache line must be in the Modified state.
2. The line must be invalidated in other cores.
3. Other cores must refetch the cache line when reading.

This “cache line ping-pong” can destroy performance.

Even reading a shared atomic variable can generate cache traffic if the ownership state changes frequently.

10.3 Mutexes and Lock Contention

Mutexes provide mutual exclusion, ensuring only one thread executes a critical section at a time.

Costs include:

- Atomic operations on lock acquisition,
- Memory fencing,
- Kernel context switches if threads must block,
- Priority inversion risks,
- Scalability collapse under high contention.

10.3.1 Spinlocks vs Blocking Mutexes

Spinlocks:

- Busy-wait using CPU cycles.
- Effective when expected lock hold time is very short.

Blocking mutexes:

- Let OS put thread to sleep.
- Better for longer critical sections.

Spinlocks can saturate core resources and interfere with hyperthreading.

10.4 Atomics and Memory Order

`std::atomic` enables lock-free operations, but atomic instructions are still expensive.

Costs of atomic operations include:

- Full or partial memory fences,
- Cache coherence traffic,
- Serialization of dependent load/store operations.

Performance increases as memory order weakens:

`memory_order_seq_cst` strongest, slowest

`memory_order_acq_rel` mid-level

`memory_order_acquire` lighter

`memory_order_release` lighter

`memory_order_relaxed` weakest, fastest (no synchronization)

10.4.1 Relaxed Atomics

```
std::atomic<int> counter{0};

void worker() {
    counter.fetch_add(1, std::memory_order_relaxed);
}
```

Safe only if cross-thread ordering is not required. Great for statistics, counters, telemetry, etc.

10.4.2 Acquire/Release Semantics

```
flag.store(true, std::memory_order_release);  
while (!flag.load(std::memory_order_acquire)) { }
```

This synchronizes side effects across cores.

10.4.3 Sequential Consistency

Default ordering. Easiest to reason about. Slowest.

10.5 Avoiding Contention

Contention occurs when multiple threads access the same resource simultaneously.

10.5.1 Sharding

Partition data so each thread works on a separate shard:

```
std::vector<int> shards[num_threads];
```

Merge results at the end.

10.5.2 Per-Thread Buffers

Keep thread-local accumulators:

```
thread_local long long local_sum = 0;
```

At the end:

```
global_sum = sum(all local_sums);
```

10.5.3 Read-Mostly Data

Use `shared_mutex` or immutable structures:

```
std::shared_mutex m;
```

Readers acquire `shared_lock`, writers acquire `unique_lock`.

10.5.4 Combining (Software Combining Tree)

One thread combines updates from others:

- threads push requests into a queue,
- combiner thread processes requests,
- reduces contention dramatically.

This is similar to combining trees and flat-combining techniques.

10.5.5 Avoiding False Sharing

Padding

```
struct alignas(64) Counter {  
    long long value;  
    char pad[64 - sizeof(long long)];  
};
```

One Cache Line per Thread

```
alignas(64) long long counters[MAX_THREADS];
```

Structure-of-Arrays Layout Split hot fields across separate arrays so threads never touch adjacent elements.

10.5.6 Recognizing False Sharing

Symptoms:

- CPU utilization capped despite available cores,
- unpredictable latency spikes,
- performance scales poorly with thread count,
- heavy cache-coherence traffic in profiling tools.

Tools like `perf c2c`, Intel VTune, or AMD uProf can detect false sharing.

10.6 Thread Scheduling and Oversubscription

Thread oversubscription—spawning more threads than hardware cores—causes:

- context switching overhead,
- cache thrashing,
- less predictable performance,
- massive degradation under contention.

Guidelines:

- Limit thread count to hardware concurrency
(`std::thread::hardware_concurrency()`).

- Use thread pools instead of launching threads repeatedly.
- Pin critical threads to specific cores when required (advanced).

10.7 Lock-Free and Wait-Free Algorithms

Lock-free algorithms can improve throughput under contention. However, they are:

- extremely complex,
- difficult to reason about,
- sensitive to ABA problems,
- reliant on correct memory ordering,
- often slower in uncontended cases.

Use them only when justified (queues, ring buffers, stacks in extremely hot paths).

10.8 Case Study: Multi-Threaded Counter

10.8.1 Version 1: Mutex

```
std::mutex m;  
long long counter = 0;  
  
void work() {  
    for (int i = 0; i < N; ++i) {  
        std::lock_guard<std::mutex> lock(m);  
        ++counter;  
    }  
}
```

Issues:

- High contention on the mutex.
- Threads serialize around the lock.

10.8.2 Version 2: Atomic Counter

```
std::atomic<long long> counter{0};

void work() {
    for (int i = 0; i < N; ++i) {
        counter.fetch_add(1, std::memory_order_relaxed);
    }
}
```

Advantages:

- No mutex,
- No blocking,
- Relaxed ordering allows maximum throughput.

However:

- Counter is highly contended → may still suffer from cache line ping-pong.
- Performance does not scale linearly with cores.

10.8.3 Version 3: Thread-Local Counters

```
thread_local long long local;

void work() {
    for (int i = 0; i < N; ++i)
        ++local;
}
```

Combine:

```
long long global = 0;
for (auto& th_local : all_locals)
    global += th_local;
```

This version has the best scalability of all.

Exercise

Write two versions of a multi-threaded counter:

1. One using `std::mutex`.
2. One using `std::atomic<long long>`.

Measure:

- Execution time as thread count increases.
- CPU utilization.
- Cache coherency traffic (using `perf stat` or `VTune`).
- Scalability on 2, 4, 8, and 16 cores.

Then, implement a third version using thread-local counters and compare scalability.

Chapter 11

Measurement, Benchmarking, and Profiling

Modern performance engineering begins and ends with measurement. Optimizations have value only when backed by accurate, reproducible data. This chapter explains how to benchmark correctly, how to use micro and macro benchmarks, how to use system-level profilers, and how to interpret results in a meaningful way.

Incorrect measurement is the fastest way to waste days or weeks of engineering effort. Correct measurement is the foundation of all real low-level optimization.

11.1 Micro-Benchmarks

Micro-benchmarks measure small, isolated pieces of code such as:

- a single function,
- an arithmetic kernel,
- a loop that manipulates an array,
- a small algorithmic comparison.

They help confirm whether a small optimization idea improves performance. However, micro-benchmarks are sensitive to noise, compiler optimizations, CPU warm-up, and testing methodology.

11.1.1 Warm-Up Phase

Never measure cold performance. Warm the function first to ensure:

- caches are primed,
- branch predictors have learned patterns,
- memory allocators are initialized,
- CPU frequency has stabilized (TurboBoost / frequency scaling settles).

Run the function 10,000–1,000,000 times before timing.

11.1.2 Avoid Dead-Code Elimination

Compilers often optimize away code that produces no observable side effects. Example of incorrect micro-benchmarking:

```
for (int i = 0; i < n; ++i)
    compute();    // may be optimized away
```

Solutions:

- use volatile outputs,
- store results in global variables,
- print or accumulate results after the benchmark (not during),
- use `benchmark::DoNotOptimize` (Google Benchmark).

11.1.3 Avoid Measuring I/O

Printing or logging inside a benchmark destroys accuracy:

```
std::cout << compute(); // NEVER do inside benchmark loop
```

I/O is magnitudes slower than computation and will dominate timing.

11.1.4 A Minimal Benchmark Helper

```
#include <chrono>
#include <iostream>

template <typename F>
void benchmark(const char* name, F&& f, std::size_t iterations) {
    // Warm-up
    for (std::size_t i = 0; i < iterations / 10; ++i)
        f();

    // Measure
    auto start = std::chrono::high_resolution_clock::now();
    for (std::size_t i = 0; i < iterations; ++i)
        f();
    auto end = std::chrono::high_resolution_clock::now();

    std::chrono::duration<double> diff = end - start;
    std::cout << name << ": " << diff.count() << " s\n";
}
```

11.1.5 Limitations of Micro-Benchmarks

Micro-benchmarks often do not represent real-world performance because:

- code fits entirely in L1/L2 cache,

- tight loops get unrolled,
- branch predictors see ideal patterns,
- memory allocators behave differently under load,
- the OS may schedule differently in realistic workloads.

Micro-benchmarks are for relative comparisons (A vs B), not full program performance predictions.

11.2 Macro-Benchmarks

Macro-benchmarks evaluate real-world performance characteristics:

- latency,
- throughput,
- scalability,
- memory consumption,
- NUMA behavior,
- cache utilization under realistic conditions.

Examples:

- evaluating a game loop simulation,
- measuring HTTP request throughput,
- benchmarking a physics or ML simulation,

- testing database query engines.

For accurate macro-benchmarks:

- fix CPU frequency (disable scaling),
- isolate the machine (no background jobs),
- run each test multiple times,
- pin threads to CPU cores,
- collect median and variance.

11.3 System Profilers

Profilers reveal the true performance bottlenecks. They tell you:

- where time is spent,
- which functions are hot,
- which loops dominate execution,
- where cache misses occur,
- where branch mispredictions occur,
- where the CPU stalls,
- where threads block on locks.

Without a profiler, performance work is guesswork.

11.3.1 Linux Profilers

`perf` `perf` is the most powerful built-in profiler on Linux.

Basic usage:

```
perf record ./my_app
perf report
```

Useful commands:

- `perf stat` — cache misses, branches, cycles
- `perf record` — sampling profiler
- `perf annotate` — see annotated assembly
- `perf c2c` — detect false sharing and cache line contention

Other Linux tools:

- `heaptrack` — tracks heap allocations and memory leaks,
- `valgrind/callgrind` — slow but detailed call graph,
- `gprof` — outdated but may still help.

11.3.2 Windows Profilers

- Windows Performance Analyzer (WPA),
- Intel VTune Profiler (best for low-level CPU work),
- Visual Studio Profiler.

VTune provides:

- pipeline stall analysis,
- SIMD vectorization efficiency,
- NUMA locality breakdown,
- cache utilization metrics,
- HyperThreading/SMT performance insights.

11.3.3 macOS Profilers

Instruments (Xcode) includes:

- CPU Sampling,
- Time Profiler,
- Allocation tracking,
- Thread scheduling visualization,
- Energy impact profiling.

11.4 Interpreting Profiling Results

Profiling produces huge amounts of data. The goal is to identify the bottleneck that limits performance.

11.4.1 Look For:

- Hot functions — where total time is spent.
- Hot loops — usually the true hotspots.
- High cache miss rates — indicates poor locality.
- Many branch mispredictions — indicates unpredictable control flow.
- High memory bandwidth usage — indicates memory-bound kernels.
- Lock contention — indicates synchronization bottlenecks.
- TLB misses — indicates poor memory layout.

11.4.2 Common Patterns and Their Causes

Large Time in memcpy Indicates too much copying, poor move semantics, or unnecessary temporaries.

High L1/L2 Miss Rate Indicates:

- linked lists,
- fragmented memory,
- AoS layout instead of SoA,
- pointer chasing.

Branch Mispredictions Often from:

- unpredictable branches,
- random data,
- hash table lookups.

Lock Contention Typical causes:

- too many threads,
- false sharing,
- critical sections too large,
- inappropriate synchronization primitives.

11.5 Benchmarking Best Practices

- Benchmark on idle systems.
- Disable CPU scaling (set fixed frequency).
- Pin threads to CPU cores.
- Run tests multiple times and average results.
- Report standard deviation.
- Use large enough input sizes to reduce noise.
- Avoid HyperThreading when measuring single-thread performance.

For multi-threaded benchmarks:

- avoid oversubscription,
- warm up thread pools,
- measure steady-state performance,
- ensure consistent work distribution across threads.

11.6 Automated Benchmarking Tools

11.6.1 Google Benchmark

Features:

- automatic warm-up,
- prevents dead-code elimination,
- reports variance,
- scales input sizes,
- integrates with CMake.

Example:

```
static void BM_vec_sum(benchmark::State& state) {  
    std::vector<int> v(state.range(0), 1);  
    for (auto _ : state) {  
        long long sum = 0;  
        for (int x : v) sum += x;  
        benchmark::DoNotOptimize(sum);  
    }  
}
```

```
}  
}  
  
BENCHMARK(BM_vec_sum)->Range(1 << 10, 1 << 20);
```

11.6.2 Celero

Focuses on high statistical accuracy; includes percentile and distribution reporting.

11.6.3 hayai

Minimalistic single-header benchmark library.

Exercise

Profile a non-trivial application you wrote. Examples:

- a physics simulation,
- a rendering pipeline,
- a web server request handler,
- a numerical computation kernel,
- a parser or compiler stage.

For each of the top three hottest functions:

1. Explain why it is slow (memory-bound, branch-heavy, copy-heavy, lock contention, etc.).
2. Suggest one concrete optimization to reduce cost.

3. Measure again and compare results.

Chapter 12

Common Patterns, Anti-Patterns, and Checklists

High-performance C++ development is not only about individual optimizations, but about cultivating a mindset of efficient design. Patterns help guide correct and performant choices, while anti-patterns highlight common traps that silently destroy performance. This chapter expands on good practices, dangerous pitfalls, and practical checklists that you can apply during development, refactoring, and code reviews.

12.1 Good Patterns

Good performance comes from predictable access patterns, minimal overhead, and high-level clarity that allows the compiler to optimize aggressively. The following patterns consistently yield excellent performance across modern CPUs.

Use `std::vector` for contiguous data

- Provides excellent cache locality.

- Enables auto-vectorization.
- Allows amortized $O(1)$ growth.
- Ensures contiguous memory for tight loops.

Prefer `std::vector` over `std::list`, `std::deque`, or hand-rolled dynamic arrays except when semantics absolutely require other containers.

Use Standard Algorithms Prefer:

```
std::sort(v.begin(), v.end());
std::accumulate(v.begin(), v.end(), 0);
std::transform(...);
std::reduce(...);
```

Algorithms are:

- highly optimized,
- easier to inline and auto-vectorize,
- often parallelizable via execution policies,
- less error-prone than manual loops.

Reserve Container Capacity When size is known or can be approximated:

```
v.reserve(expected_size);
```

This avoids repeated reallocations, copying, and memory fragmentation. In performance-critical paths, calling `reserve()` often yields large improvements.

Use Non-Owning Views

- `std::string_view` for read-only substrings,
- `std::span<T>` for non-owning array and buffer views.

Advantages:

- avoid copying memory,
- avoid dynamic allocations,
- avoid unnecessary temporaries,
- enable cheap slicing and iteration.

Keep Hot Data Small and Contiguous Hot code paths should operate on:

- compact structures,
- tightly-packed arrays,
- small POD structs,
- SoA layouts if needed for SIMD.

Small, contiguous data maximizes:

- L1/L2 cache hit rates,
- prefetch efficiency,
- SIMD utilization,
- predictable control flow.

This pattern is essential for numeric kernels, simulations, games, and HPC workloads.

12.2 Anti-Patterns

Anti-patterns represent practices that silently destroy performance. Avoid these unless a very strong reason proves they are necessary.

Allocating Memory in Inner Loops Example of a dangerous anti-pattern:

```
for (int i = 0; i < n; ++i) {  
    std::vector<int> tmp;    // expensive allocation per iteration  
    tmp.push_back(i);  
}
```

Heap allocations are slow because they:

- take locks,
- update metadata,
- cause cache invalidation,
- fragment memory.

Solution: Preallocate, reuse buffers, or move allocations outside the loop.

Unpredictable Branches in Hot Loops Branches with random outcomes destroy branch predictor accuracy:

```
if (rand() % 2) process(a); else process(b);
```

This leads to:

- pipeline stalls,

- branch mispredictions,
- unpredictable latency.

Refactor to branchless code where possible or reorder data to become more predictable.

Overuse of Node-Based Containers `std::list`, `std::map`, `std::unordered_map`, and other node-based containers:

- fragment memory,
- break locality,
- prevent SIMD,
- cause pointer chasing,
- cause many cache misses.

Use them only when required by semantics (stable iterators, sorted keys, etc.).

Virtual Calls in Hot Loops Virtual dispatch prevents inlining and adds:

- vtable lookups,
- unpredictable branches,
- reduced optimization opportunities.

Replace with:

- templates (static polymorphism),
- CRTP,

- function objects,
- `std::variant`.

In extremely hot paths, avoid virtual dispatch entirely.

12.3 Practical Optimization Checklist

High-performance C++ development requires systematic thinking. Before attempting low-level tuning, walk through this checklist:

1. **Algorithmic Efficiency:** Is the algorithm optimal for your problem? A better algorithm can outperform months of micro-optimizations.
2. **Data Structure Selection:** Are containers designed for locality? Can data be stored contiguously?
3. **Minimizing Copies:** Are you eliminating unnecessary copying using move semantics, views, and references?
4. **Memory Allocation:** Are allocations minimized? Are you using `reserve()`, pools, or arenas?
5. **Branch Structure:** Are branches predictable? Can branchless techniques help?
6. **SIMD Awareness:** Does the code allow auto-vectorization? Are loops simple and dependency-free?
7. **Concurrency Strategy:** Are locks minimized? Is contention reduced? Are false-sharing risks addressed?
8. **Compiler Optimizations:** Are you compiling with `-O2` or `-O3`? Is LTO enabled? Is PGO appropriate?

9. Profiling: Have you used a profiler to find real hotspots? Are you optimizing code that actually matters?
10. Benchmarking: Have you ensured stable, reproducible measurements? Have you tested multiple data sizes and patterns?

This checklist is essential for code reviews, optimization cycles, and performance audits.

Exercise

Create a project-specific checklist of 10–15 performance rules based on patterns in your codebase. Examples:

- never allocate in hot loops,
- always reserve vectors when size is predictable,
- avoid virtual dispatch in compute kernels,
- use SoA for wide numeric datasets,
- avoid unpredictable branches on hot paths,
- use profiling before optimization,
- document assumptions about data locality,
- verify vectorization using compiler reports,
- avoid shared mutable data unless necessary,
- use `std::string_view` instead of copying strings.

Review all performance-critical code against your checklist during code reviews.

Final Conclusion

Low-level optimization in C++ is not a collection of isolated tricks, secret compiler switches, or obscure micro-optimizations. It is a disciplined way of thinking about software performance from the perspective of the machine itself. It is an engineering mindset rooted in measurement, clarity, and mechanical sympathy — a deep respect for how modern CPUs, memory hierarchies, compilers, and concurrency models actually behave.

At its core, low-level optimization means understanding and applying a few fundamental principles:

- Know your hardware: Modern CPUs are marvels of parallelism and sophistication — out-of-order execution, branch predictors, multi-level caches, and SIMD pipelines. Your code runs fast only when it cooperates with these features instead of fighting them.
- Value locality: Data layout determines performance far more than most developers realize. Spatial and temporal locality are not optional — they are the heart of fast software.
- Eliminate unnecessary work: Every copy, every branch, every allocation has a cost. Reducing redundancy often matters more than adding complexity.

- Express intent clearly: Modern C++ features such as move semantics, constexpr, templates, and views give you a language that is high-level, expressive, and highly optimizable.
- Measure relentlessly: Intuition fails more often than it succeeds. Profilers, benchmarks, and systematic experimentation are the only reliable guides.

Throughout this guide, one theme has been repeated: performance is mostly determined by big decisions, not minute tweaks. The most transformative improvements in real-world C++ systems usually come from:

- selecting the optimal algorithmic approach,
- choosing the right data structures for CPU and cache efficiency,
- avoiding unnecessary dynamic allocation and object churn,
- minimizing synchronization and contention in multi-threaded programs,
- shaping data layouts to exploit SIMD and parallelism,
- simplifying control flow to help branch prediction and instruction-level parallelism.

Once these fundamentals are correct, fine-grained optimization becomes both easier and more effective.

Modern C++ as a Performance Language

Modern C++ (from C++11 to C++26 and beyond) gives developers an unprecedented set of tools for writing fast, robust, and maintainable systems. Features like:

- `std::move`, perfect forwarding, and value categories,

- `std::span`, `std::string_view`, and non-owning references,
- `constexpr`, `constexpr`, and compile-time computation,
- concepts, ranges, and improved generics,
- parallel STL algorithms and high-level concurrency constructs,

allow you to write code that is both elegant and close to the metal. When used well, these features let the compiler generate code that is as efficient as hand-written C, assembly, or SIMD — without sacrificing type safety, abstraction, or maintainability.

The Role of the Programmer

Even with powerful compilers, the programmer still plays the decisive role. Only you can:

- design a data-oriented architecture,
- ensure predictable control flow,
- eliminate branches and allocations in the hot path,
- use profiling data to guide decisions,
- recognize when a problem is algorithmic, not mechanical,
- understand when simplicity outperforms cleverness.

Optimization is an iterative dialogue between you and the machine. Each run of the profiler teaches you something about how your program behaves in the real world. Each refactor aligns the code more closely with the hardware. Each improvement compounds the benefit of all the others.

A Philosophy of Sustainable Performance

High-performance C++ does not mean writing cryptic, fragile code that only one person can maintain. On the contrary, sustainable performance comes from:

- writing clear, simple, intention-revealing code,
- structuring programs so compilers can optimize them,
- building systems that scale predictably as data or cores increase,
- using abstractions responsibly — zero-cost whenever possible,
- favoring correctness and robustness first, tuning second.

Fast code that is fragile is not an achievement; fast code that is also clean, maintainable, and correct is the mark of a professional C++ engineer.

Looking Ahead

Computing hardware continues to evolve: more cores, wider SIMD, deeper pipelines, larger and more complex memory systems, and increasingly heterogeneous architectures. C++ will continue to evolve with it — offering even more tools for safely expressing high-performance intent across CPUs, GPUs, accelerators, and embedded systems. The skills developed in this guide will remain fundamental regardless of future hardware changes. They are not tricks tied to a specific CPU, but timeless habits of thought:

- think in terms of data,
- think in terms of cost,
- think in terms of locality and parallelism,
- think like the compiler,
- think like the CPU.

Final Thoughts

May this guide serve not only as a reference but as an inspiration—helping you see your C++ programs from the machine’s point of view, sharpen your intuition about performance, and encourage you to write software that achieves excellence in both speed and clarity.

With knowledge, measurement, and discipline, you can consistently build systems that are:

- fast, because they cooperate with hardware,
- reliable, because they avoid undefined behavior and fragile micro-optimizations,
- scalable, because they use concurrency intelligently,
- maintainable, because their structure is clean and intentional.

High performance in C++ is not magic. It is a craft — one you now have the tools to master.

May your future code run both beautifully and blazingly fast.

References

Performance engineering in C++ is a broad discipline that intersects with compiler theory, CPU architecture, memory systems, concurrent programming, numerical computing, and modern C++ language design. The following references provide authoritative, in-depth knowledge that complements the material in this guide. They are grouped by topic for easier navigation.

Core Modern C++ References

- Scott Meyers, *Effective Modern C++* A foundational text covering the best practices and idioms of C++11 and C++14. Its guidance on move semantics, auto type deduction, and modern abstractions directly supports writing efficient and expressive code.
- Herb Sutter, articles and presentations (pragmatic, optimization-oriented C++ guidance) Topics include:
 - Zero-cost abstractions,
 - Concurrency in C++,
 - C++ memory model,
 - Guidelines for modern coding style.

- Bjarne Stroustrup, *A Tour of C++ (2nd/3rd Editions)* A modern overview from the creator of C++, including language rules, philosophy, and performance design goals.
- Nicolai M. Josuttis, *C++17 — The Complete Guide*, *C++20 — The Complete Guide* Detailed explanations of modern C++ features from concepts to ranges, with insights on efficiency and implementational behavior.

Low-Level Optimization and Hardware Architecture

- Agner Fog, *Optimizing C++ Software*, *Instruction Tables*, *Microarchitecture Guides* Perhaps the most authoritative publicly available resource on:
 - instruction timings,
 - CPU pipelines,
 - superscalar execution,
 - ILP and vectorization,
 - branch prediction,
 - memory bandwidth,
 - micro-architectural bottlenecks.
- Ulrich Drepper, *What Every Programmer Should Know About Memory* Essential reading for understanding:
 - cache hierarchies,
 - NUMA architectures,
 - page tables and TLBs,

- cost of memory access,
 - CPU $\leftrightarrow$ memory interactions.
- Intel® 64 and IA-32 Architectures Optimization Reference Manual Covers hardware prefetching, vectorization, performance counters, and CPU design principles.
- AMD Optimization Manual A parallel to Intel’s guide but focused on AMD Zen architecture pipeline, front-end, and cache behavior.
- Daniel Lemire, blog and academic papers (SIMD, fast integer parsing, efficient algorithms) Practical research on making common operations dramatically faster using modern hardware.

Compilers and Toolchain References

- Matt Godbolt et al., Compiler Explorer An essential tool for:
 - inspecting compiler output,
 - learning how optimizers transform C++,
 - comparing GCC, Clang, and MSVC behavior,
 - studying inlining, vectorization, and RVO.
- GCC Optimization and Performance Tuning Manual Explains GCC-specific optimization flags, diagnostics, and code-generation internals.
- Clang/LLVM Documentation, especially:
 - optimization remarks (OptRemarks),
 - vectorization reports,

- LTO/ThinLTO guides.
- MSVC Optimization and Code Generation Guide Covers Visual Studio’s optimizer, PGO, auto-vectorization, and Windows-specific ABI details.
- LLVM Source Code and Compiler Design Papers For those wanting deep insight into:
 - SSA representation,
 - instruction selection,
 - loop optimization passes,
 - register allocation strategies.

Memory, Allocation, and Data Layout

- Jeff Preshing, blog posts on concurrency, atomics, and memory ordering Clear and practical explanations of:
 - lock-free programming,
 - memory fences,
 - cache behavior,
 - ABA problems,
 - atomic operations.
- Facebook Folly Library and Performance Notes Real-world engineering insights about:
 - scalable allocators,
 - cache-aware containers,

- lock-free and wait-free designs.
- Google’s TCMalloc and Abseil Documentation High-performance memory allocation strategies and design patterns.
- jemalloc Documentation Covers fragmentation, per-thread caches, and allocation tuning.

Parallelism, Concurrency, and Multithreading

- Anthony Williams, C++ Concurrency in Action A definitive reference on:
 - threads,
 - locks,
 - atomics,
 - memory models,
 - lock-free programming.
- Herb Sutter, Atomic Weapons, C++ Memory Model and Atomics Deep and accessible explanations of the C++ atomic model.
- Intel Threading Building Blocks (TBB) documentation High-performance task-based parallelism.
- OpenMP Specification Parallel loops and directives widely used in HPC.
- C++ Parallel STL, ISO papers Explains execution policies, parallel algorithms, and underlying threading models.

Profiling, Benchmarking, and Tools

- Google Benchmark Library Industry-standard micro-benchmarking framework for modern C++.
- perf (Linux) documentation Essential for sampling, hardware counters, cache misses, branch mispredicts, memory bandwidth.
- Intel VTune Profiler The gold standard for low-level CPU analysis: stalls, pipelines, ILP, vectorization, NUMA.
- Valgrind / Callgrind / Cachegrind Whole-program instrumentation for:
 - call graphs,
 - cache simulation,
 - cost-based analysis.
- Heaptrack Industry-grade memory profiling, allocation hot spot detection.
- Android / Linux Perfetto and Tracing Tools For latency profiling and system-wide tracing.

ISO C++ Standards and Committee Papers

- ISO C++ Standard (C++17, C++20, C++23, C++26 drafts) For authoritative definitions of:
 - object models,
 - constant evaluation,
 - memory models,

- concurrency rules,
 - new optimizable language features.
- WG21 Papers (wg21.link) Thousands of proposals, including:
 - Ranges,
 - Executors,
 - Low-level concurrency primitives,
 - Pattern matching,
 - Zero-cost abstractions,
 - SIMD extensions,
 - Performance-driven designs.

Academic and Research-Oriented Resources

- PLDI, PPOPP, CGO, MICRO, ASPLOS conferences Cutting-edge research on compiler optimization and CPU architecture.
- High-Performance Computing (HPC) courses Cover:
 - vectorization,
 - memory bandwidth models (Roofline),
 - data locality theory,
 - GPU/CPU hybrid workloads.
- MIT OCW – Performance Engineering of Software Systems Highly recommended free course.

Online Resources

- cppreference.com The most complete online reference for C++17–C++23.
- [Compiler Explorer \(godbolt.org\)](http://godbolt.org) Essential for studying generated assembly.
- [Reddit r/cpp](https://www.reddit.com/r/cpp) and [r/cpp_questions](https://www.reddit.com/r/cpp_questions) Community discussions, advice, and examples.
- [StackOverflow C++ tag](https://stackoverflow.com/questions/tagged/c++) High-quality Q&A on performance problems.
- [C++ Weekly \(Jason Turner\)](https://www.youtube.com/channel/UCWv3e333U333U333U333U333U) Detailed videos on optimization, constexpr, assembly, and modern idioms.
- [Andrzej's C++ articles \(Andrzej Krzemienski\)](https://ericniebler.com/) Deep dives into correctness and efficiency.

Purpose of This Reference List

This reference list is not intended to be exhaustive. Instead, it highlights the most influential, practical, and widely used resources for mastering low-level optimization in modern C++:

- authoritative books by leading experts,
- industrial-grade optimization manuals,
- academic research foundations,
- profiling and benchmarking frameworks,
- ISO papers driving the evolution of high-performance C++.

These materials provide a lifetime of learning and will serve as your roadmap as you deepen your understanding of performance engineering, CPU architecture, and advanced C++ design.