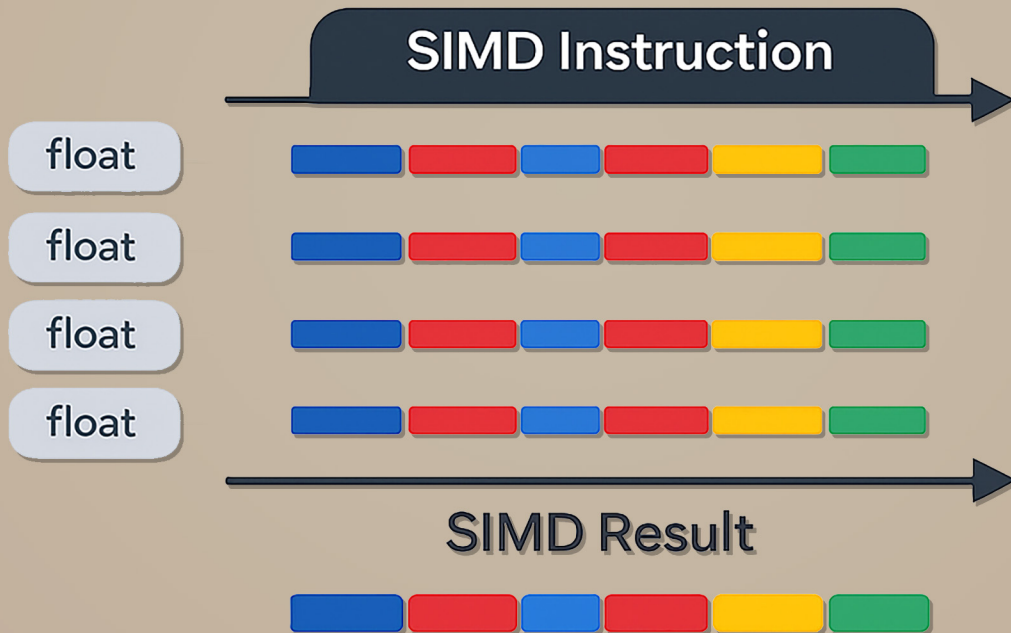


SIMD in Modern C++

A Shortcut Guide for Programmers (Without Complexity)



SIMD in Modern C++

A Shortcut Guide for Programmers (Without Complexity)

Prepared by Ayman Alheraki

simplifcpp.org

November 2025

Contents

Contents	2
Author's Introduction	5
1 Why SIMD Matters on Modern CPUs	9
1.1 What SIMD Does	9
1.2 Why SIMD Provides Real Performance Gains	11
1.2.1 Key Domains Where SIMD Dominates	11
1.3 Scalar vs Auto-Vectorized Code	12
1.3.1 Why Auto-Vectorization Works Well Today	13
2 Writing SIMD-Friendly Modern C++	15
2.1 Rules for Auto-Vectorization	15
2.2 Understanding Aliasing	17
2.3 Bad vs Good Example	17
2.3.1 Bad Example: Aliasing Prevents Vectorization	17
2.3.2 Good Example: Using <code>std::span</code> to Remove Aliasing Doubts	18
2.4 Using <code>std::transform</code> and C++ Algorithms	18
2.5 Additional Guidelines for SIMD-Friendly C++	19
2.5.1 Favor Data-Oriented Design	19
2.5.2 Minimize the Work Inside the Loop	19

2.5.3	Avoid Data-Dependent Branching	19
2.5.4	Use Compiler Vectorization Reports	20
3	Portable Explicit SIMD Using <code>std::experimental::simd</code>	22
3.1	Basic Example	22
3.2	Masked Operations	24
3.3	More Useful SIMD Operations	25
3.3.1	Elementwise Math	25
3.3.2	Blending and Conditional Assignment	25
3.3.3	Reduction	25
3.3.4	Gather/Scatter	25
3.4	Performance Notes	26
4	SIMD Intrinsics — When You Need Maximum Control	28
4.1	x86 SIMD Intrinsics Overview	29
4.2	x86 AVX2 Example	30
4.3	ARM NEON Intrinsics Overview	30
4.4	ARM NEON Example (2023+)	31
4.5	When to Use Intrinsics Instead of Portable SIMD	32
5	<code>xsimd</code> : A Modern Header-Only SIMD Library	33
5.1	<code>xsimd</code> Overview	34
5.2	<code>xsimd</code> Scaling Example	35
5.3	Vectorized Math with <code>xsimd</code>	36
5.4	Blending and Masking	36
5.5	Alignment Considerations	37
5.6	When to Choose <code>xsimd</code>	37
	Final Conclusion	40

Author's Introduction

Modern computing hardware has evolved dramatically over the past two decades. The era of relying purely on clock frequency is long over; architectural innovation now comes from deeper forms of parallelism. Today's CPUs — whether running on desktops, servers, embedded boards, or mobile devices — are highly sophisticated machines built around multiple layers of concurrency:

- Multi-core execution for task-level parallelism.
- Simultaneous multithreading (SMT / hyperthreading) for increased utilization.
- Out-of-order and speculative execution to maximize instruction throughput.
- Wide SIMD (Single Instruction, Multiple Data) units that operate on vectors of data in parallel.
- Dedicated functional units for fused operations, shuffles, masks, reductions, and data rearrangement.

Among these features, SIMD remains one of the most powerful yet underutilized capabilities available to C++ programmers. It is present in every modern CPU architecture — Intel and AMD (SSE, AVX, AVX2, AVX-512, AVX10), ARM (NEON, SVE, SVE2), and the rapidly growing RISC-V ecosystem (RVV). These vector engines

can perform between 4 and 64 operations per instruction, depending on width and datatype. Despite this, most software written today still processes data one element at a time, using scalar loops that leave vast computational potential unused.

Why does this happen? In practice, SIMD is often perceived as difficult or intimidating. Developers frequently associate SIMD with:

- low-level intrinsics filled with cryptic mnemonics,
- alignment constraints and manual data shuffling,
- platform-specific instruction sets,
- and complex compiler behavior that can make vectorization feel like trial-and-error.

As a result, many programmers treat SIMD as something only numerical libraries, game engines, physics engines, or compiler experts should touch. This is no longer the case.

Modern C++ has fundamentally changed how we access SIMD.

Since C++17, and especially in the C++20 - C++23 era, the language and its surrounding ecosystem have introduced tools that make data-parallel programming dramatically more accessible:

- Cleaner, safer data structures (`std::span`, `std::array`, `std::vector`) that help compilers auto-vectorize reliably.
- Execution policies such as `std::execution::par_unseq` that explicitly enable vectorization in algorithms.
- The Parallelism TS 2 and the `std::experimental::simd` library, which provide a portable, high-level abstraction over SIMD instructions.

- Upcoming standardization of `std::simd` — bringing portable SIMD into official ISO C++ for the first time.
- High-quality libraries like `xsimd`, `Eigen`, and the SIMD layers in popular ML frameworks.
- Compilers that have matured significantly, providing better diagnostics, improved vectorization heuristics, and broader ISA coverage.

This booklet is written with a clear mission: to demystify SIMD for professional and aspiring Modern C++ developers. You should not need to be a CPU microarchitect, compiler engineer, or linear algebra expert to benefit from SIMD. You simply need practical guidance — and that is what this booklet provides.

What you will learn in this booklet

1. How to write C++ code that compilers naturally and consistently auto-vectorize.
2. How to reason about performance, vector widths, alignment, memory access, and data layout.
3. How to use the portable, expressive API of `std::experimental::simd` and the upcoming `std::simd`.
4. When and how to fall back to intrinsics for maximum control without drowning in complexity.
5. How to use modern libraries (2020–2025) that provide clean abstractions while producing highly optimized vector code.
6. How to build a mental model for how SIMD works on x86, ARM, and RISC-V architectures.

Each chapter is structured around:

- Clear explanations of concepts without assuming deep mathematical or hardware background.
- Modern C++20/23/26 code examples that compile cleanly on today's compilers.
- Practical performance considerations rather than theoretical microbenchmark tricks.
- Exercises at the end of each chapter designed to reinforce understanding.
- Complete, precise solutions that illustrate correct patterns and best practices.

By progressing through these chapters, you will be able to integrate SIMD naturally into your everyday C++ development workflow. You will understand how to write loops, algorithms, and data processing code that scale automatically across the entire SIMD landscape — from compact NEON on ARM phones to massive 512-bit vectors on high-end servers.

More importantly, you will acquire a mindset:

Performance is not a late-stage optimization; it is a design philosophy.

SIMD is one of the strongest tools you can add to your modern performance toolbox. My goal in this booklet is to show you that SIMD is not a mysterious, low-level specialty - but a natural, powerful extension of clean, modern, expressive C++.

By the end of this booklet, SIMD will no longer feel like a hidden world inside your CPU. It will feel accessible, predictable, and deeply empowering — enabling you to produce faster, more efficient, and more scalable C++ software across all modern platforms.

Ayman Alheraki

Chapter 1

Why SIMD Matters on Modern CPUs

1.1 What SIMD Does

Modern processors are designed to extract maximum throughput from parallel data operations. One of the core mechanisms enabling this high throughput is SIMD — Single Instruction, Multiple Data. It describes a model where a single machine instruction performs the same operation on multiple values simultaneously.

Instead of processing elements one-by-one (scalar mode), SIMD packs multiple values into a vector register and operates on them in parallel. This is not merely a theoretical abstraction; it is an essential part of every mainstream CPU architecture today:

- x86_64 supports SSE, AVX, AVX2, AVX-512, and the upcoming AVX10.
- ARM supports NEON for mobile and embedded devices, and SVE/SVE2 for servers and high-performance Apple Silicon chips.
- RISC-V introduces the scalable vector extension (RVV), a flexible approach for variable-width SIMD.

A SIMD register is a wide container capable of holding multiple numbers at once. For example, a 256-bit AVX register can hold:

- Eight float (32-bit) elements, or
- Four double (64-bit) elements.

This leads to the intuitive performance difference:

scalar: 1 multiply per instruction

SIMD: 8 multiplies per instruction

But this is only the surface. Modern SIMD instructions support:

- Vector addition, subtraction, multiplication, division
- Fused multiply-add (FMA) for enhanced precision and performance
- Vector comparisons and masks
- Horizontal reductions (sum, min, max)
- Bitwise logic and packing/unpacking
- Data shuffles, blends, and permutations
- Efficient gather/scatter for non-contiguous memory

SIMD is not merely about doing “eight operations at once.” It is about enabling the CPU to reach its true arithmetic throughput by minimizing instruction overhead, maximizing register utilization, and leveraging dedicated vector execution pipelines.

1.2 Why SIMD Provides Real Performance Gains

SIMD is not an optional feature for modern software. It is a primary source of performance across nearly all computational workloads.

When processing large collections of numerical data, the bottleneck is rarely the control logic. The bottleneck is the number of arithmetic operations needed and the number of memory accesses performed. SIMD addresses both issues:

1. It processes multiple data items with one instruction.
2. It loads/stores multiple values in a single memory operation.

In well-structured code, the performance lift is substantial — often 2x to 16x, depending on vector width and algorithm structure.

1.2.1 Key Domains Where SIMD Dominates

- Linear algebra and scientific computing. Dot products, matrix multiplications, and vector transforms rely heavily on wide parallel arithmetic.
- Machine learning preprocessing and inference. Normalization, activation functions, quantization, and tensor operations map naturally to SIMD.
- Image and signal processing. Blurs, filters, convolutions, pixel-level transforms, and FFT operations are SIMD-rich workloads.
- Cryptography and compression. Many symmetric ciphers, hash functions, and encoding algorithms operate on blocks of data that vectorize extremely well.
- Audio and DSP processing. Mixers, equalizers, filters, and sample transformations benefit directly from vectorized operations.

- Game development, physics, and simulation engines. SIMD accelerates vector math, collision detection, and particle systems.
- Database and analytics engines. Modern query engines use SIMD to filter, compare, scan, and aggregate large datasets.

In fact, many of the fastest libraries today (BLAS, FFTW, OpenCV, Eigen, Halide, TensorFlow Lite) incorporate SIMD at their core.

1.3 Scalar vs Auto-Vectorized Code

One of the strongest advantages of Modern C++ (C++17 and newer) is that well-structured loops can automatically turn into SIMD instructions through compiler auto-vectorization. This means that even if you never explicitly write SIMD code, the compiler may still generate vectorized machine instructions for you.

Consider the following scalar function:

```
// Scalar: sum of squares
float sumsq(const float* data, size_t n) {
    float s = 0.0f;
    for (size_t i = 0; i < n; ++i)
        s += data[i] * data[i];
    return s;
}
```

With an optimized build:

```
g++ -O3 -march=native
```

most modern compilers transform the inner loop into a vectorized sequence:

- On x86_64: AVX2 or AVX-512 instructions

- On ARM: NEON or SVE instructions
- On RISC-V: RVV instructions if available

The compiler may collapse eight or sixteen multiplications and additions into a single fused vector loop, drastically improving performance.

1.3.1 Why Auto-Vectorization Works Well Today

Compilers (GCC, Clang/LLVM, MSVC) have become significantly more sophisticated since 2020:

- Improved loop analysis and alias detection
- Better detection of pure or side-effect-free operations
- Stronger heuristics for vectorizing reduction loops (sum, min, max)
- Support for masked operations and predicated vector loops
- More robust support for vectorizing `std::transform` and `std::reduce`

Modern C++ techniques like using `std::span`, avoiding unnecessary pointer aliasing, and structuring clean loops help compilers generate highly efficient SIMD code automatically.

Exercises (Ch. 1)

1. Why is SIMD beneficial for array-based workloads compared to scalar processing?
2. How many 32-bit floating-point values fit inside a 512-bit SIMD register?
3. Identify three operations from your own projects that would likely benefit from SIMD acceleration.

Solutions (Ch. 1)

- SIMD allows parallel arithmetic on multiple elements using a single instruction, greatly increasing arithmetic throughput and reducing instruction overhead.
- A 512-bit register can store $512/32 = 16$ floats.
- Answers will vary, but common examples include pixel operations, signal filters, vector transforms, statistical reductions, or buffer processing.

Chapter 2

Writing SIMD-Friendly Modern C++

Writing SIMD-friendly code is one of the most essential steps toward achieving high performance in Modern C++. Even with powerful tools such as `std::experimental::simd`, intrinsics, or high-level SIMD libraries, your first and biggest performance wins often come from enabling the compiler to auto-vectorize your loops. This requires writing code that the compiler can easily analyze and safely convert into SIMD instructions. This chapter provides a detailed, practical guide to structuring loops, organizing data, and applying modern C++ idioms that maximize SIMD opportunities. These rules apply across all major architectures: x86 (SSE/AVX/AVX-512), ARM (NEON/SVE), and RISC-V (RVV).

2.1 Rules for Auto-Vectorization

Auto-vectorization occurs when the compiler transforms a scalar loop into a SIMD loop without requiring explicit programmer intervention. However, compilers must be mathematically certain that vectorization will not change program semantics. If the compiler detects potential aliasing, dependencies, side effects, or complex control flow,

it will conservatively fall back to scalar code.

To maximize auto-vectorization opportunities, follow these core principles:

- Loops must use simple linear indexing. Patterns like `for (i = 0; i < n; ++i)` with direct access `a[i]` allow the compiler to reason about memory access.
- Memory must be contiguous. Containers that guarantee contiguous storage include:
 - `std::vector<T>`
 - `std::array<T, N>`
 - Raw C arrays (`T*`)
 - `std::span<T>`
- Aliasing must be eliminated or minimized. If the compiler cannot guarantee that two pointers refer to non-overlapping memory, vectorization is disabled.
- Branches must be avoided inside the loop. Data-dependent if statements inhibit SIMD because SIMD lanes must execute the same operations.
- Avoid complex function calls inside hot loops. Virtual calls, function pointers, and non-inlineable functions block vectorization.
- Ensure the loop has no hidden side effects. Accessing global variables, modifying shared state, or using volatile memory breaks SIMD safety.
- Use trivially copyable arithmetic types. SIMD excels with `float`, `double`, `int32_t`, etc.

2.2 Understanding Aliasing

Aliasing occurs when two pointers may refer to overlapping memory. If the compiler cannot prove that:

$$a[i] \neq b[j]$$

it cannot safely vectorize the loop, because writing to `a[i]` might overwrite memory that will later be read from `b[j]`. As a result, the compiler must generate scalar code.

Modern C++ helps eliminate aliasing problems using:

- `const` qualifiers
- distinct function parameters
- standard containers
- `std::span` which communicates both bounds and contiguity

Refactoring pointer-based code into spans or vectors immediately improves SIMD performance.

2.3 Bad vs Good Example

2.3.1 Bad Example: Aliasing Prevents Vectorization

```
void scale(float* a, float* b, size_t n) {  
    for (size_t i = 0; i < n; ++i)  
        a[i] = a[i] + 1.5f * b[i];  
}
```

Here the compiler cannot guarantee that `a` and `b` do not overlap, so the loop cannot be vectorized safely.

2.3.2 Good Example: Using `std::span` to Remove Aliasing Doubts

```
#include <span>

void scale(std::span<float> a, std::span<const float> b) {
    for (size_t i = 0; i < a.size(); ++i)
        a[i] += 1.5f * b[i];
}
```

Benefits include:

- `b` is `const`: writes to `a` cannot affect `b`.
- `span` conveys contiguity and size explicitly.
- Simple, linear indexing.

Most compilers will auto-vectorize this into AVX2/AVX-512/NEON/SVE instructions.

2.4 Using `std::transform` and C++ Algorithms

The C++ standard algorithms are highly SIMD-friendly because they express intent clearly and reduce manual indexing. Modern execution policies such as `par_unseq` explicitly allow vectorization.

```
float sumsq(const std::vector<float>& v) {
    return std::transform_reduce(
        std::execution::par_unseq,
        v.begin(), v.end(),
        0.0f,
        std::plus<>{},
        [](float x){ return x * x; }
    );
}
```

Key properties:

- `par_unseq` enables both vectorization and parallelization.
- The lambda is side-effect free.
- Data is guaranteed contiguous in `std::vector`.

Compilers commonly generate very efficient SIMD for this pattern.

2.5 Additional Guidelines for SIMD-Friendly C++

2.5.1 Favor Data-Oriented Design

Keep data in arrays, spans, and vectors of plain arithmetic types. Avoid complex class wrappers when processing numeric data.

2.5.2 Minimize the Work Inside the Loop

Smaller loop bodies yield:

- better inlining,
- simpler dependency graphs,
- more predictable vectorization.

2.5.3 Avoid Data-Dependent Branching

Branches such as `if (a[i] > 0)` must be lowered into masked operations, which is far less efficient. Branchless code is preferable.

2.5.4 Use Compiler Vectorization Reports

Examples:

```
g++ -O3 -fopt-info-vec
```

```
clang++ -O3 -Rpass=loop-vectorize
```

```
msvc: /Qvec-report:2
```

These tools reveal why loops did or did not vectorize.

Exercises (Ch. 2)

1. Explain why `std::span` improves auto-vectorization.
2. Rewrite a loop computing $c[i] = a[i] * b[i] + d[i]$ using `std::transform`.
3. Identify a function in your own codebase that could be rewritten to become SIMD-friendly.

Solutions (Ch. 2)

- Span improves vectorization because it guarantees contiguity, provides explicit bounds, and improves const-correctness, reducing aliasing uncertainty.
- Transform solution:

```
std::transform(  
    a.begin(), a.end(),  
    b.begin(),  
    c.begin(),  
    [&](float aa, float bb){
```

```
    return aa * bb + d[i]; // or zip d properly  
}  
);
```

- Answers vary. Common examples: normalization loops, filters, reductions, numerical kernels, image/audio buffers.

Chapter 3

Portable Explicit SIMD Using `std::experimental::simd`

While auto-vectorization provides substantial performance improvements, some workloads require predictable and explicit control over SIMD execution. The C++ Parallelism TS 2 introduces a powerful abstraction: `std::experimental::simd`. This library provides a portable, expressive, and type-safe interface for explicit data-parallel programming without writing ISA-specific assembly or intrinsics.

Unlike intrinsics, which depend on architecture-specific APIs such as `immintrin.h` (x86) or `arm_neon.h` (ARM), the `simd` abstraction generates the appropriate instructions for the target architecture. This allows developers to write SIMD code that is clean, modern, and portable across x86, ARM, and RISC-V.

3.1 Basic Example

The following example demonstrates a simple, portable vectorized scaling operation. It adds $1.5f * b[i]$ to each element of `a[i]` using native SIMD width for the target

architecture:

```
#include <experimental/simd>
namespace stdx = std::experimental;

void scale_simd(float* a, float* b, size_t n) {
    using simd_t = stdx::native_simd<float>;
    using size_type = std::size_t;

    size_type i = 0;
    for (; i + simd_t::size() <= n; i += simd_t::size()) {
        simd_t va(&a[i], stdx::element_aligned);
        simd_t vb(&b[i], stdx::element_aligned);
        va = va + 1.5f * vb;
        va.copy_to(&a[i], stdx::element_aligned);
    }

    for (; i < n; ++i)
        a[i] += 1.5f * b[i];
}
```

Important points about this example:

- `native_simd` uses the widest available SIMD registers (SSE/AVX/AVX-512, NEON, SVE, or RVV).
- The loop processes elements in chunks of `simd_t::size()`.
- Remaining elements (tail) are processed using scalar operations.
- The constructor and `copy_to` use `element_aligned`, ensuring optimal loads/stores.

This code compiles into efficient SIMD instructions across all major architectures while remaining portable and easy to maintain.

3.2 Masked Operations

Masks allow selective computation on elements of a SIMD vector. Modern SIMD architectures support predicate registers that operate on multiple lanes simultaneously. The TS exposes this via `simd_mask`:

```
stdx::simd_mask<float> mask = (va > vb);  
auto r = stdx::where(mask, va + vb);  
r.copy_to(out, stdx::element_aligned);
```

This sequence:

1. Compares `va` and `vb` element-wise, producing a mask.
2. Executes the addition only where the mask is true.
3. Leaves other elements unchanged or unspecified based on where semantics.

Masked operations allow:

- conditional arithmetic,
- selective updates,
- branchless logic within SIMD loops,
- replacement of if-statements that traditionally block vectorization.

This approach maps naturally to predication on ARM SVE/SVE2, AVX-512 mask registers, and RVV vector masks.

3.3 More Useful SIMD Operations

The `std::experimental::simd` interface includes many utilities that mirror intrinsics but with better portability:

3.3.1 Elementwise Math

```
va = std::sin(va);  
vb = std::sqrt(vb);  
vc = std::abs(vc);
```

These operations use vectorized math routines when available.

3.3.2 Blending and Conditional Assignment

```
auto mask = (va < 0.0f);  
va = stdx::where(mask, -va, va);
```

3.3.3 Reduction

```
float sum = stdx::reduce(va);
```

3.3.4 Gather/Scatter

Supported on architectures that implement it (AVX2+, SVE, RVV):

```
simd_t r = stdx::simd<float>(a, idx);  
r.copy_to(a, idx);
```

3.4 Performance Notes

- SIMD size depends on architecture: 128, 256, 512 bits, or scalable width.
- Code is portable but still compiled into native vector instructions.
- Masks allow branchless logic, crucial for high-performance kernels.
- Memory alignment improves throughput but is not mandatory thanks to `element_aligned`.

Portable SIMD using `std::experimental::simd` frequently delivers performance close to handwritten intrinsics while being far more readable and maintainable.

Exercises (Ch. 3)

1. Modify the basic example so that each element is replaced by its square root.
2. Write a masked operation that updates `a[i]` only when `a[i] > 0`.
3. Create a reduction that computes the sum of all elements after scaling.

Solutions (Ch. 3)

- Square root update:

```
va = std::sqrt(va);
```

- Masked update for positive elements:

```
auto mask = (va > 0.0f);  
va = stdx::where(mask, va * 2.0f, va);
```

- Reduction example:

```
float total = stdx::reduce(va);
```

Chapter 4

SIMD Intrinsics — When You Need Maximum Control

While auto-vectorization and portable SIMD via `std::experimental::simd` offer high productivity and portability, there are situations where a developer needs absolute and explicit control over SIMD execution. These scenarios arise when:

- You need maximum performance in a critical hot loop.
- You must use features not yet exposed by portable abstractions (e.g., AVX-512 masks, special blends, gathers/scatters).
- You need deterministic register usage or instruction sequences.
- You are developing platform-specific algorithms (game engines, cryptography, DSP, physics engines).

In such cases, SIMD intrinsics — architecture-specific low-level APIs — allow you to directly access the CPU's vector instructions. This chapter focuses on the two dominant instruction sets used in modern computing:

- x86_64 SIMD intrinsics (SSE/AVX/AVX2/AVX-512)
- ARM NEON intrinsics (ARMv8-A, Apple Silicon, mobile devices)

These intrinsics map nearly 1:1 to CPU instructions and give total control at the cost of portability and readability. They are powerful tools when used with care and understanding.

4.1 x86 SIMD Intrinsics Overview

The `<immintrin.h>` header provides the full set of intrinsics for Intel and AMD CPUs, covering:

- SSE and SSE2/SSE3/SSE4.x instructions (128-bit)
- AVX and AVX2 (256-bit)
- AVX-512 (512-bit)
- FMA instructions (fused multiply-add)

Key intrinsic types include:

- `__m128` and `__m128i` for 128-bit vectors
- `__m256` and `__m256i` for 256-bit vectors
- `__m512` and `__m512i` for 512-bit vectors

Operations include arithmetic, logic, shuffling, masks, blends, gathers, scatters, reductions, and type conversions.

4.2 x86 AVX2 Example

Below is a simple AVX2 example that adds two arrays of eight floats using 256-bit registers:

```
#include <immintrin.h>

void add_avx(const float* a, const float* b, float* c) {
    __m256 va = _mm256_loadu_ps(a);
    __m256 vb = _mm256_loadu_ps(b);
    __m256 vc = _mm256_add_ps(va, vb);
    _mm256_storeu_ps(c, vc);
}
```

This sequence maps directly to the following CPU operations:

- Load 8 floats from memory into a 256-bit SIMD register.
- Load another 8 floats from memory.
- Perform 8 parallel floating-point additions.
- Store the result back to memory.

To scale this example across large arrays, you simply repeat this pattern inside a loop, advancing by 8 elements per step.

4.3 ARM NEON Intrinsics Overview

ARM NEON is the ubiquitous SIMD architecture found in:

- ARM Cortex-A processors (Android phones, tablets)

- ARM servers
- Apple Silicon (M1, M2, M3, etc.)
- Embedded systems and IoT devices

NEON uses 128-bit SIMD registers:

- `float32x4_t` for four 32-bit floats
- `int32x4_t`, `uint8x16_t`, etc. for integer vectors

NEON intrinsics are exposed through the `<arm_neon.h>` header.

4.4 ARM NEON Example (2023+)

The following example demonstrates how to add four floats using NEON intrinsics:

```
#include <arm_neon.h>

void add_neon(const float* a, const float* b, float* c) {
    float32x4_t va = vld1q_f32(a);
    float32x4_t vb = vld1q_f32(b);
    float32x4_t vc = vaddq_f32(va, vb);
    vst1q_f32(c, vc);
}
```

This compiles into efficient NEON instructions for any ARMv8-A CPU:

- `vld1q_f32`: Load a vector of 4 floats
- `vaddq_f32`: Add vectors in parallel
- `vst1q_f32`: Store the result

To operate over longer arrays, each iteration increments the pointer by four elements.

4.5 When to Use Intrinsics Instead of Portable SIMD

Intrinsics should be used when:

- Maximum performance is required.
- You need ISA-specific features (AVX-512 masks, NEON saturating arithmetic).
- Portable abstractions do not expose needed instructions.
- You require precise instruction scheduling.
- You are writing performance-critical libraries (DSP, cryptography, codecs).

However:

- Intrinsics are not portable.
- They increase code complexity.
- They require per-architecture branching or separate implementations.
- Maintainers must understand CPU microarchitecture.

When used appropriately, intrinsics unlock the highest possible performance of the underlying hardware.

Chapter 5

xsimd: A Modern Header-Only SIMD Library

While `std::experimental::simd` provides a standard and portable abstraction for SIMD programming, many production-grade libraries and scientific computing frameworks rely on high-performance, battle-tested SIMD layers. One of the most widely used libraries in this space is `xsimd`, a modern header-only SIMD wrapper that provides a clean C++ API for writing portable SIMD code without depending directly on low-level intrinsics. `xsimd` is used by:

- the `xtensor` numerical library (NumPy-like for C++),
- machine-learning preprocessing layers,
- high-performance data-analytics pipelines,
- custom physics simulations and signal-processing engines,
- C++ backends that require portable vectorization across x86, ARM, and RISC-V.

The key design goal of xsimd is to give developers the power of SIMD instructions with the elegance of high-level C++ code, while still generating highly optimized AVX2, AVX-512, NEON, or RVV instructions depending on the target system.

5.1 xsimd Overview

xsimd achieves portability by abstracting SIMD operations into batches. A batch represents a SIMD register of architecture-dependent width:

- On AVX2 CPUs: `batch<float>` uses 256-bit registers (8 floats).
- On AVX-512 CPUs: it uses 512-bit registers (16 floats).
- On ARM NEON: it uses 128-bit registers (4 floats).
- On RISC-V RVV: it adapts to scalable vectors.

The control logic, arithmetic operations, loads, stores, and mathematical functions are all exposed in natural C++ syntax:

- `batch a, b;` for SIMD registers
- `a + b`, `a * b`, `sin(a)`, `sqrt(a)` for vectorized math
- `load_unaligned`, `store_unaligned` for memory interaction

This means that xsimd allows developers to write vectorized code that is both portable and intuitive.

5.2 xsimd Scaling Example

The following example demonstrates how to scale one vector by another using a SIMD batch, with automatic handling of any tail elements:

```
#include <xsimd/xsimd.hpp>

void scale_xsimd(float* a, float* b, size_t n) {
    using batch = xsimd::batch<float>;
    size_t step = batch::size;

    size_t i = 0;
    for (; i + step <= n; i += step) {
        batch va = batch::load_unaligned(a + i);
        batch vb = batch::load_unaligned(b + i);
        batch vr = va + 1.5f * vb;
        vr.store_unaligned(a + i);
    }

    for (; i < n; ++i)
        a[i] += 1.5f * b[i];
}
```

This function behaves similarly to a hand-written AVX2 or NEON loop, but with several advantages:

- Portability: The same code compiles into AVX2/AVX-512/NEON/RVV instructions without modification.
- Automatic vector width selection: `batch::size` adjusts based on the CPU architecture.

- Readable arithmetic: Expressions like `va + 1.5f * vb` are compiled into efficient SIMD instructions.
- Graceful fallback: The tail loop ensures all elements are processed, even when the size is not a multiple of the SIMD width.

5.3 Vectorized Math with xsimd

xsimd supports a broad range of mathematical operations beyond simple arithmetic:

- Trigonometric: `sin`, `cos`, `tan`
- Exponential: `exp`, `log`
- Reduction: `hadd` for horizontal addition
- Comparison and masks: `va > vb`, `xsimd::select`
- FMA (fused multiply-add) where supported

Example:

```
batch x = batch::load_unaligned(input);  
batch y = xsimd::sin(x) + xsimd::sqrt(x);  
batch::store_unaligned(output, y);
```

All operations are vectorized across the SIMD lanes.

5.4 Blending and Masking

xsimd provides mask-based selection similar to AVX-512 and SVE predication:

```
batch a = ...;
batch b = ...;
auto mask = (a > b);
batch r = xsimd::select(mask, a, b);
```

This offers fast branchless conditional logic inside SIMD kernels.

5.5 Alignment Considerations

`xsimd` supports both:

- Aligned loads/stores for optimal performance
- Unaligned loads/stores for safety and flexibility

Modern CPUs handle unaligned loads well, but aligned data can improve peak bandwidth in some cases. `xsimd` makes both options explicit and simple.

5.6 When to Choose `xsimd`

Use `xsimd` when:

- You require portable SIMD across many architectures.
- You want readable SIMD code without touching intrinsics.
- You need high-performance batch operations in numerics or ML.
- You prefer a stable, well-tested library over raw intrinsics.

Compared to `std::experimental::simd`, `xsimd`:

- provides earlier adoption and broader compiler support,
- includes a larger set of math functions,
- integrates seamlessly with xtensor and HPC applications.

Exercises (Ch. 4)

1. Modify the example to compute $a[i] = \sqrt{a[i]} + b[i]$ using xsimd.
2. Write an xsimd loop that clamps each element of a to the range $[0, 1]$.
3. Implement a masked selection choosing the maximum between two arrays.

Solutions (Ch. 4)

- Square root update:

```
batch va = batch::load_unaligned(a + i);
batch vb = batch::load_unaligned(b + i);
batch vr = xsimd::sqrt(va) + vb;
vr.store_unaligned(a + i);
```

- Clamp example:

```
batch va = batch::load_unaligned(a + i);
batch vr = xsimd::max(batch(0.0f), xsimd::min(va, batch(1.0f)));
vr.store_unaligned(a + i);
```

- Masked max:

```
auto mask = (va > vb);  
batch vr = xsimd::select(mask, va, vb);  
vr.store_unaligned(out + i);
```


Final Conclusion

SIMD is no longer a niche optimization technique reserved for specialists working in assembly language, graphics engines, or scientific computing libraries. It has become a fundamental and indispensable part of Modern C++ performance engineering.

Today's hardware architectures — whether based on x86, ARM, or RISC-V — provide increasingly wide vector instruction sets, scalable vector architectures, and powerful fused operations that allow programmers to extract vast computational throughput from a single CPU core. Ignoring SIMD means leaving most of the processor's real capabilities unused.

Throughout this booklet, we explored the complete journey of becoming proficient with SIMD in Modern C++:

- Writing SIMD-friendly loops. You learned how to structure loops so that compilers can reliably auto-vectorize them. This includes avoiding unnecessary aliasing, maintaining simple indexing patterns, ensuring data contiguity, and eliminating hidden dependencies that block vectorization.
- Leveraging strong compiler auto-vectorization. Modern compilers (GCC, Clang/LLVM, MSVC) have become remarkably effective at turning clean C++ loops into SIMD instructions. Understanding how they analyze loops, and how to help them succeed, provides significant performance gains with minimal effort.

- Using portable explicit SIMD with `std::experimental::simd`. The Parallelism TS 2 introduces a powerful, expressive, and portable abstraction for explicit SIMD programming. You gained practical examples showing how `native_simd`, masks, reductions, and vector math operations map cleanly across architectures without writing intrinsics.
- Using intrinsics for full control. When maximum performance or precise control is required, intrinsics provide direct access to architecture-specific instructions such as AVX2, AVX-512, NEON, and RVV. These allow highly tuned kernels at the cost of portability, and are essential when developing performance-critical libraries and systems.
- Using high-level portable SIMD libraries like `xsimd`. Libraries such as `xsimd` deliver robust, readable, and efficient SIMD abstractions built on top of multiple architectures. They simplify writing numerical kernels, image processing filters, and signal-processing pipelines without sacrificing performance or clarity.

Together, these layers form the complete Modern C++ SIMD ecosystem:

- High-level auto-vectorization for effortless performance.
- Portable explicit SIMD for predictable and maintainable kernels.
- Architecture-specific intrinsics for expert-level optimization.
- Cross-platform SIMD libraries for rapid development and portability.

Mastering SIMD provides developers with a profound advantage. It allows you to write software that adapts automatically to new CPU generations, uses hardware more efficiently, and achieves performance close to that of hand-written vectorized assembly — all while staying within clean, modern, expressive C++.

As CPUs continue to evolve, vector widths will increase, predication models will mature, and new scalable vector architectures like ARM SVE and RISC-V RVV will become mainstream. The knowledge and techniques presented in this booklet ensure that your C++ code will continue to scale gracefully with these advancements.

SIMD is not just an optimization; it is an essential mindset for modern high-performance C++ programming. With the tools, techniques, and understanding presented in this work, you are well-equipped to build robust, efficient, and future-proof software for the computing landscapes of 2025 and beyond.

References

The following references represent a curated and authoritative collection of modern materials (2020–2025) essential for mastering SIMD programming in Modern C++, compiler vectorization, CPU architecture, and data-parallel abstractions. These sources span ISO proposals, architecture manuals, academic-quality optimization guides, and high-level frameworks. Together they provide both theoretical foundations and practical implementation details.

- ISO C++ Committee. P0214: Data-Parallel Types for C++. (2020–2024). Official proposal defining the standardization of `std::simd`. Provides complete design rationale, examples, ABI considerations, and portable semantics for SIMD operations.
- ISO C++ Committee Papers on Parallelism TS 2. (2020–2025). Working drafts and proposals evolving the data-parallel model adopted by compilers and high-level frameworks.
- GCC Team. GCC Auto-Vectorization and Loop Optimization Documentation. (2021–2025). Covers vectorizer heuristics, supported transforms, vectorization conditions, debug flags, and architecture-specific optimizations.
- LLVM/Clang Development Team. Clang Loop Vectorization and SLP Vectorizer Guide. (2021–2025). Explains the inner workings of Clang’s vectorizer, SLP

algorithms, cost models, masked vectorization, and reporting tools.

- Intel Corporation. Intel 64 and IA-32 Architectures Optimization Reference Manual. (2021–2025). Comprehensive documentation of AVX2, AVX-512, AVX10, FMA instructions, load/store pipelines, register layout, and scheduling.
- Intel Corporation. Intel AVX-512 and AVX10 ISA Extensions. (2021–2025). Detailed instruction descriptions, masking semantics, rounding control, and performance guidelines for 512-bit SIMD.
- ARM Ltd.. ARM NEON Programmer’s Guide. (2020–2025). Covers vector lanes, NEON load/store, arithmetic, shuffles, saturating ops, and architectural constraints.
- ARM Ltd.. ARM SVE and SVE2 Architecture Reference. (2020–2025). Defines scalable vector architecture, predication, vector-length-agnostic design, and advanced vector operations for HPC.
- RISC-V International. RISC-V Vector Extension (RVV) Specification, v1.0+. (2020–2025). Definitive reference for scalable vectors, masking, segment loads/stores, vector length register semantics, and vectorized control flow.
- xtensor / xsimd Project. xsimd: C++ SIMD Library Documentation. (2022–2025). High-level, portable SIMD abstraction layer integrating with NumPy-like tensor structures, enabling cross-platform vectorization.
- Agner Fog. Optimizing Software in C++. Updated 2023. Renowned performance guide covering microarchitecture, instruction throughput, memory hierarchy, vector pipelines, and optimization strategies for real-world code.

- Agner Fog. Instruction Tables, Microarchitecture Guides. (2020–2025). Detailed latency/throughput tables for x86 instructions, crucial for hand-tuned SIMD optimizations.
- Microsoft Visual C++ Team. MSVC Vectorization and Parallelization Documentation. (2022–2025). Including /Qvec-report, SIMD intrinsic coverage, and ISA-specific optimizations for Windows platforms.
- CppCon / Herb Sutter. Talks on Parallelism, Performance, and Modern C++. (2020–2024). Foundational presentations explaining the philosophy of modern C++ performance, data-parallel design, and architectural abstractions.
- NVIDIA. SIMT and SIMD Comparison Material. (2022–2025). Helps clarify architectural differences between GPU SIMT and CPU SIMD for HPC developers.
- Google Benchmark / Performance Engineering Notes. (2020–2025). Practical guidance for benchmarking vectorized code, understanding cache/memory effects, and avoiding measurement pitfalls.
- University Research Papers on Vectorization and Compilers. (2021–2025). Peer-reviewed studies on loop dependence analysis, superword-level parallelism (SLP), and masked vectorization.