

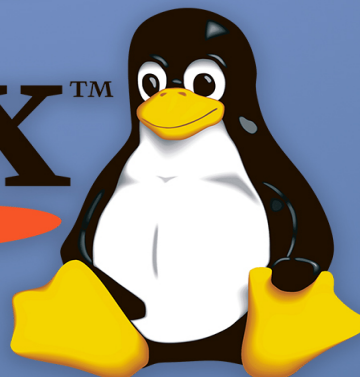
Low-Level Programming on Linux (x86-64)

Introduction to x86-64 Architecture and Memory



1

Linux™



Low-Level Programming on Linux (x86-64) :

Introduction to x86-64 Architecture and Memory

Prepared by Ayman Alheraki

simplifycpp.org

November 2025

Contents

Contents	2
Author's Introduction	13
1 Overview of Computer Architecture	15
1.1 Instruction Set Architecture vs. Microarchitecture	15
1.1.1 Instruction Set Architecture (ISA)	15
1.1.2 Microarchitecture	16
1.1.3 The Compatibility Contract	17
1.1.4 Importance for Low-Level Programmers	17
1.2 Hardware Abstraction in Modern Systems	19
1.2.1 Purpose of Hardware Abstraction	19
1.2.2 Privilege Levels and Protected Execution	19
1.2.3 Virtual Memory and Address Translation	20
1.2.4 Device Abstraction and Drivers	21
1.2.5 Architectural Feature Discovery	21
1.2.6 Significance for Low-Level Programming	22
1.3 The Role of the Operating System	23
1.3.1 Process Abstraction and Execution Management	23
1.3.2 Memory Management and Protection	24

1.3.3	System Calls and the Privileged Interface	24
1.3.4	Device Control and I/O Coordination	25
1.3.5	Security, Isolation, and Execution Integrity	25
1.3.6	Significance for Low-Level Programming	26
1.4	Where Assembly Fits in the Stack	27
1.4.1	Machine Code and Architectural Semantics	27
1.4.2	Assembly and High-Level Languages	27
1.4.3	Assembly in System-Level Software	28
1.4.4	Assembly and the ABI Boundary	29
1.4.5	Why Assembly Still Matters	29
2	Internal Structure of x86-64 CPUs	31
2.1	Execution Units	31
2.1.1	Role in the Out-of-Order Pipeline	31
2.1.2	Functional Classes of Execution Units	32
2.1.3	Port-Based Issue Architecture	33
2.1.4	Latency and Throughput Considerations	33
2.1.5	Implications for Low-Level Software	34
2.2	Decode Engine & Micro-Operations	35
2.2.1	Complexity of the x86-64 Instruction Format	35
2.2.2	Micro-Operations (μ ops)	36
2.2.3	The μ op Cache and Decode Bypass	36
2.2.4	Fusion and Decomposition	37
2.2.5	Interaction with Out-of-Order Scheduling	37
2.2.6	Significance for Low-Level Development	37
2.3	Reorder Buffer & Out-of-Order Execution	39
2.3.1	Motivation for Out-of-Order Execution	39
2.3.2	Structure and Role of the Reorder Buffer	39

2.3.3	Execution and Result Commitment	40
2.3.4	Handling Speculation and Mispredictions	40
2.3.5	Memory Ordering and the Load/Store Queue	41
2.3.6	Practical Implications for Low-Level Programming	41
2.4	Register Rename Logic	43
2.4.1	Architectural vs. Physical Registers	43
2.4.2	Rename Table and Allocation	43
2.4.3	Retirement and Version Release	44
2.4.4	Elimination of False Dependencies	45
2.4.5	Interaction with Branch Speculation	45
2.4.6	Impact on Low-Level Performance	45
2.4.7	Summary	46
3	The Memory Hierarchy	47
3.1	L1 / L2 / L3 Cache Layers	47
3.1.1	L1 Cache: Fast, Private, and Specialized	47
3.1.2	L2 Cache: Larger and Still Private	48
3.1.3	L3 Cache: Shared and Designed for Multi-Core Coordination	48
3.1.4	Relationship to DRAM and Access Latency	49
3.1.5	Implications for Low-Level Programming	50
3.2	Cache Lines and Latency	51
3.2.1	Cache Line Size	51
3.2.2	Cache Misses and Latency Sources	51
3.2.3	Stride Patterns and Cache Efficiency	52
3.2.4	Alignment and Alias Boundaries	53
3.2.5	Prefetching and Latency Hiding	53
3.2.6	Implications for Low-Level Programming on Linux	54
3.3	NUMA and Multi-Core Memory Effects	55

3.3.1	Memory Nodes and Locality	55
3.3.2	Multi-Core Workloads and Data Placement	56
3.3.3	Interconnect Contention and Bandwidth Sharing	56
3.3.4	Effects on Synchronization and Shared State	57
3.3.5	NUMA on Linux: Allocation and Scheduling Behavior	57
3.3.6	Implications for Low-Level Programming	58
3.4	Cache Miss Penalties	59
3.4.1	Cost Escalation Across the Hierarchy	59
3.4.2	Stall Behavior and Instruction-Level Parallelism Limits	59
3.4.3	Memory-Level Parallelism	60
3.4.4	Causes of Cache Miss Penalties	60
3.4.5	Measuring Cache Misses on Linux	61
3.4.6	Strategies to Reduce Cache Miss Impact	61
4	Registers and Flags	63
4.1	General Purpose Registers	63
4.1.1	Register Width and Naming Hierarchy	63
4.1.2	Roles and Conventional Usage	64
4.1.3	The Stack Pointer and Call Frame Structure	65
4.1.4	Pointer and Addressing Roles	65
4.1.5	Performance Considerations	66
4.1.6	Significance for Low-Level Programming	66
4.2	RFLAGS Bit Meanings	67
4.2.1	Structure of RFLAGS	67
4.2.2	Status Flags (Affect Conditional Instructions)	67
4.2.3	Control and Direction Flags	68
4.2.4	System Flags (Privilege and Execution Context)	69
4.2.5	Flag Dependencies and Performance Considerations	69

4.2.6	Practical Guidelines for Low-Level Programming	69
4.3	Special Registers (RIP, RSP, RBP)	71
4.3.1	Instruction Pointer (RIP)	71
4.3.2	Stack Pointer (RSP)	72
4.3.3	Base Pointer (RBP) and Stack Frames	72
4.3.4	Stack Discipline and Calling Semantics	73
4.3.5	Interaction with the Kernel and Context Switching	73
4.3.6	Importance for Low-Level Development	74
4.4	SIMD Registers (XMM / YMM)	75
4.4.1	XMM Registers (128-bit SIMD)	75
4.4.2	YMM Registers (256-bit SIMD)	76
4.4.3	ABI and Function Calling Considerations (Linux / System V AMD64) .	76
4.4.4	Alignment and Performance Behavior	77
4.4.5	Significance for Low-Level Programming	78
5	Instruction Flow	79
5.1	Fetch	79
5.1.1	Instruction Pointer and Sequential Flow	79
5.1.2	Instruction Cache and Front-End Bandwidth	80
5.1.3	Instruction Flow Optimization via μ op Cache	80
5.1.4	Fetch and Branch Prediction	81
5.1.5	Effects of Instruction Density and Encoding	81
5.1.6	Significance for Low-Level Programming	82
5.2	Decode	83
5.2.1	Variable-Length Instruction Parsing	83
5.2.2	Translation Into Micro-Operations (μ ops)	84
5.2.3	Decode Width and Front-End Throughput	84
5.2.4	The Microcode ROM	85

5.2.5	μop Cache Short-Circuiting Decode	85
5.2.6	Practical Implications for Low-Level Development	86
5.3	Execute	87
5.3.1	Execution Units and μop Dispatch	87
5.3.2	Latency vs. Throughput	88
5.3.3	Operand Readiness and Dependency Chains	88
5.3.4	Speculative Execution	89
5.3.5	Memory Access in Execution	89
5.3.6	Implications for Low-Level Development	90
5.4	Writeback & Commit	91
5.4.1	Writeback to Physical Registers	91
5.4.2	The Reorder Buffer (ROB) and In-Order Retirement	91
5.4.3	Register Rename Resolution at Commit	92
5.4.4	Stores and Memory Visibility	93
5.4.5	Exception and Interrupt Precision	93
5.4.6	Performance Considerations	94
5.4.7	Summary	94
6	Virtual Memory & Paging	95
6.1	Address Translation	95
6.1.1	Virtual vs. Physical Addresses	95
6.1.2	Page-Based Translation	96
6.1.3	Multi-Level Page Tables (x86-64)	96
6.1.4	The TLB and Translation Caching	97
6.1.5	Privilege Enforcement During Translation	97
6.1.6	Virtual Memory Allows Memory Overcommitment	98
6.1.7	Practical Implications for Low-Level Programming	98
6.2	Page Tables (4-Level Paging)	99

6.2.1	Hierarchical Layout	99
6.2.2	Page Table Levels and Roles	100
6.2.3	Page Table Entry Format	100
6.2.4	Example Translation Walk (Conceptual)	101
6.2.5	Cost of a Page Walk	101
6.2.6	Practical Implications for Low-Level Programming	102
6.3	TLB and Misses	103
6.3.1	TLB Structure and Levels	103
6.3.2	TLB Miss Cost	103
6.3.3	Impact of Working Set Size	104
6.3.4	Page Size and TLB Efficiency	104
6.3.5	Access Patterns and TLB Behavior	105
6.3.6	Measuring and Diagnosing TLB Misses on Linux	105
6.3.7	Practical Strategies for Low-Level Optimization	106
6.4	Kernel vs User Address Space	107
6.4.1	Virtual Address Space Layout (Typical Linux x86-64)	107
6.4.2	Privilege Levels and Access Control	108
6.4.3	System Calls: Controlled Transition to Kernel Mode	108
6.4.4	Shared Kernel Mapping and Security Constraints	109
6.4.5	User Space Memory Management vs Kernel Memory Management	109
6.4.6	Practical Implications for Low-Level Programming	110
7	Interrupts and Exceptions	111
7.1	IDT Structure	111
7.1.1	IDT Location and Register	111
7.1.2	IDT Entry Format (Gate Descriptors)	112
7.1.3	Vector Number Assignment	112
7.1.4	Privilege Level Control	113

7.1.5	Interrupt Stack Table (IST) Integration	113
7.1.6	Practical Low-Level Considerations	114
7.2	Hardware vs Software Interrupts	115
7.2.1	Hardware Interrupts	115
7.2.2	Masking and Prioritization	116
7.2.3	Software Interrupts	116
7.2.4	Behavioral Differences	117
7.2.5	Kernel Handling Differences	117
7.2.6	Practical Low-Level Considerations	118
7.3	Exception Types	119
7.3.1	Classification of Exceptions	119
7.3.2	Faults (Recoverable Errors)	120
7.3.3	Traps (Observation Points)	120
7.3.4	Aborts (Non-Recoverable Errors)	121
7.3.5	Delivering Exceptions to User Space	121
7.3.6	Practical Low-Level Guidelines	121
7.4	Common Debug Traps	123
7.4.1	INT3 (Breakpoint Trap)	123
7.4.2	Single-Step Trap (Trap Flag in RFLAGS)	123
7.4.3	Hardware Breakpoints (DR0–DR7 Debug Registers)	124
7.4.4	Watchpoints and Data Access Traps	124
7.4.5	Page Protection Traps for Debugging	125
7.4.6	Interaction With ptrace and the Kernel	126
7.4.7	Practical Low-Level Guidelines	126
8	Processor Manuals	127
8.1	Intel Manuals Layout	127
8.1.1	Volume Structure	127

8.1.2	Architectural vs. Microarchitectural Information	128
8.1.3	Instruction Format in Volume 2	129
8.1.4	System Programming Details in Volume 3	130
8.1.5	Reading Model for Practical Development	130
8.1.6	Key Principle	131
8.2	AMD APM Structure	132
8.2.1	Organization of the APM Set	132
8.2.2	AMD-Originated Architectural Elements	133
8.2.3	Memory System and Paging Clarifications	133
8.2.4	Model-Specific Registers (Volume 4)	134
8.2.5	Areas Where AMD APM Adds Beyond Intel	134
8.2.6	Practical Interpretation for Low-Level Development	135
8.3	How to Navigate Instruction Tables	136
8.3.1	Opcode and Encoding Fields	136
8.3.2	Operand Types and Naming	136
8.3.3	RFLAGS Effects	137
8.3.4	Exceptions and Fault Conditions	138
8.3.5	Mode-Specific Behavior	138
8.3.6	Finding Latency and Throughput Information	139
8.3.7	Practical Workflow for Low-Level Programming	139
9	Practical Register Tracing with GDB	140
9.1	Inspecting Registers	140
9.1.1	Entering a Debugging Session	140
9.1.2	Displaying General-Purpose Registers	141
9.1.3	Inspecting Flags (RFLAGS)	142
9.1.4	Inspecting SIMD Registers (XMM/YMM)	142
9.1.5	Modifying Registers During Debugging	143

9.1.6	Recording State Across Breakpoints	143
9.1.7	Practical Guidelines	143
9.2	Single-Step Execution	145
9.2.1	Entering Single-Step Mode	145
9.2.2	Viewing the Current Instruction	145
9.2.3	Observing Register Changes Per Instruction	146
9.2.4	Branch and Loop Analysis	147
9.2.5	Stepping Through Function Calls	147
9.2.6	Single-Stepping and System Instructions	148
9.2.7	Practical Considerations	148
9.2.8	Summary	149
9.3	Understanding Disassembly Flow	150
9.3.1	Entry Points and Instruction Context	150
9.3.2	Sequential vs. Control-Flow Transitions	151
9.3.3	Function Calls and Stack Frame Interpretation	151
9.3.4	Understanding Compiler Optimization Effects	152
9.3.5	Branch Targets and Symbol Resolution	153
9.3.6	Memory Address Interpretation	153
9.3.7	Practical Strategy for Reading Disassembly	154
10	Project	155
10.1	Choose an Instruction	155
10.1.1	Criteria for Instruction Selection	155
10.1.2	Example Instruction Selection: <code>add rax, rbx</code>	157
10.1.3	Verify No Compiler Interference	157
10.1.4	Define the Measurement and Observation Plan	158
10.1.5	Expected Learning Outcome	158
10.2	Trace Its Execution Step-by-Step	160

10.2.1	Prepare a Minimal Test Harness	160
10.2.2	Launch Under GDB	161
10.2.3	Record the Pre-Execution State	161
10.2.4	Execute the Instruction One Step	162
10.2.5	Vary Operand Types for Comparative Tracing	162
10.2.6	Trace Memory Operand Variants	163
10.2.7	Document Observations Objectively	163
10.2.8	Key Takeaway	164
10.3	Generate a Microarchitectural Report	165
10.3.1	Identify μ op Count and Execution Units	165
10.3.2	Latency vs. Throughput Interpretation	166
10.3.3	Characterize Register Dependencies	166
10.3.4	Memory Hierarchy Impact (If Memory Operand Was Used)	167
10.3.5	Report Format	167
10.3.6	Example Summary (Conceptual)	168
10.3.7	Key Insight	168
Conclusion		170
References		172

Author's Introduction

The x86-64 architecture remains one of the most widely deployed and deeply influential computing platforms in modern systems. From personal workstations to large-scale servers and cloud environments, Linux on x86-64 continues to define the execution environment for software that demands reliability, performance, and direct access to hardware resources. Yet, for many developers, the internal structure of this architecture remains at a distance—often abstracted behind high-level languages, compilers, runtime environments, and operating system interfaces.

This booklet was written to close that gap.

Its purpose is not to teach assembly programming in isolation, nor to romanticize low-level techniques for their own sake. Instead, the goal is to provide a clear and structured understanding of how instructions, registers, memory hierarchies, and privilege boundaries actually function inside a modern Linux system on x86-64 hardware. The material is organized to build intuition: from pipeline stages and register semantics, to caching behavior, paging structures, interrupt dispatch, and the observable effects visible through GDB and performance inspection.

The modern processor is neither simple nor opaque—its behavior can be understood precisely when approached from the correct perspective: the intersection of architecture (what the CPU guarantees) and microarchitecture (how it achieves it). Throughout this booklet, we intentionally maintain that distinction. We examine real execution, trace register state transitions, analyze the behavior of instructions at the machine level, and connect each

observation to the architectural rules defined by Intel and AMD.

The project component in the final chapter reinforces this approach. By studying a single instruction end-to-end—from encoding, to execution, to microarchitectural behavior—we make the CPU’s operation tangible and verifiable. Such a method fosters technical confidence: understanding becomes demonstrable, measurable, and exact.

This booklet is intended for programmers, systems engineers, security researchers, and students who seek to strengthen their grasp of how software truly executes beneath abstraction. It assumes seriousness of purpose and a willingness to learn details that most software never exposes. The reward for that effort is the ability to navigate complex systems with clarity rather than approximation.

Low-level knowledge is not nostalgia. It is capability.

Stay Connected

For more discussions and valuable content about **Introduction to x86-64 Architecture and Memory**, I invite you to follow me on **LinkedIn**:

<https://linkedin.com/in/aymanalheraki>

You can also visit my personal website:

<https://simplifycpp.org>

Ayman Alheraki

Chapter 1

Overview of Computer Architecture

1.1 Instruction Set Architecture vs. Microarchitecture

A fundamental distinction in computer architecture is the separation between the **Instruction Set Architecture (ISA)** and the **microarchitecture** that implements it. Although both terms describe aspects of processor design, they operate at different abstraction levels and evolve under different constraints.

1.1.1 Instruction Set Architecture (ISA)

The **Instruction Set Architecture** defines the contract between software and hardware. It specifies the set of machine instructions, visible registers, addressing modes, memory model, and calling conventions that programs rely on. From the compiler's perspective, the ISA is the target model; generated machine code must conform to the ISA specification to run correctly.

The **x86-64 ISA** (also referred to as AMD64) defines:

- A 64-bit general-purpose register file (RAX, RBX, ..., R15).

- A hierarchical memory addressing model and segmentation rules (largely deprecated in 64-bit mode, though still architecturally present).
- The format, encoding, and semantics of instructions (e.g., MOV, ADD, LEA, CALL, RET).
- Support for SIMD/vector instructions such as SSE, AVX, and AVX-512 depending on CPU feature flags.
- The architectural state required for context switching (control registers, flags, descriptor tables).

The ISA is **stable and backward-compatible**. A binary compiled in 2005 must still execute on a processor released in 2025, unless it depends on newly introduced optional features. This stability is central to the longevity of the x86 line on desktop, server, and cloud environments.

1.1.2 Microarchitecture

While the ISA defines *what* the processor must do, the **microarchitecture** defines *how* the processor achieves it. Microarchitecture concerns internal implementation strategies such as:

- Pipeline depth and structure.
- Out-of-order execution and instruction scheduling logic.
- Branch prediction algorithms and speculative execution behavior.
- Cache hierarchy design (L1/L2/L3, replacement policies, latency trade-offs).
- Execution units and port distribution for ALU, FPU, and vector operations.

Different processors may implement the same ISA while offering widely different performance characteristics. For example:

Vendor / Generation	ISA	Microarchitectural Focus
AMD Zen 4 (2022+)	x86-64	Wider execution front-end, improved branch prediction, optimized AVX-512 support
Intel Alder Lake (2021+)	x86-64	Hybrid P-core + E-core design for performance efficiency scaling
VIA Centaur CHA (2020)	x86-64	Compact server-oriented design with integrated AI acceleration units

All execute the same x86-64 instructions, but their microarchitectures differ substantially in pipeline width, cache topology, power management, and speculative execution barriers. This separation allows hardware designers to innovate without breaking software compatibility.

1.1.3 The Compatibility Contract

Because the ISA is the exposed interface, any microarchitectural optimization must preserve architectural correctness. If an instruction is architecturally defined to update flags or create memory side effects in a particular sequence, the processor must present those outcomes identically—even if internally the instruction was reordered, fused, or speculatively executed and later rolled back.

This is especially relevant on Linux systems, where binaries compiled once may run across generations of hardware from various vendors. The ELF binary format, the dynamic loader, and the kernel’s feature detection mechanisms depend on ISA stability and explicit reporting of optional instruction set extensions through CPU feature flags.

1.1.4 Importance for Low-Level Programmers

Understanding the distinction between ISA and microarchitecture is essential when:

- Writing performance-critical loops and memory-intensive algorithms.

- Investigating latency vs. throughput behavior of instructions.
- Designing multi-threaded systems sensitive to cache coherency and NUMA placement.
- Analyzing performance counters (e.g., via `perf` on Linux) to locate pipeline stalls, TLB misses, and branch mispredictions.

Low-level optimization is effective only when the programmer recognizes **which behaviors are guaranteed by the ISA** and **which are performance characteristics of a specific microarchitecture**. The former is stable; the latter varies across CPU models and future generations.

1.2 Hardware Abstraction in Modern Systems

Modern computer systems rely on structured layers of abstraction that allow software to execute without requiring direct knowledge of the underlying physical hardware. This abstraction is essential for portability, stability, security isolation, and performance scalability across different generations of processors and platforms. For low-level programming on Linux, understanding which components are abstracted and which remain directly accessible is central to writing efficient and correct system-level software.

1.2.1 Purpose of Hardware Abstraction

Hardware abstraction separates **architectural behavior**—the operations defined by the Instruction Set Architecture (ISA)—from **system-level integration**, such as memory management, device communication, and process isolation. This enables:

- The same binary to run on different microarchitectures implementing the same ISA.
- The operating system to schedule and protect multiple programs concurrently.
- User applications to operate without requiring privileged access to hardware.

On Linux, this abstraction is formalized through the kernel, which mediates interaction between software and physical resources according to defined system call and driver interfaces.

1.2.2 Privilege Levels and Protected Execution

Modern x86-64 processors support hierarchical privilege modes:

- **Kernel Mode (Ring 0):** Full access to processor instructions, control registers, physical memory, and hardware devices.

- **User Mode (Ring 3):** Restricted environment in which applications execute. Direct hardware access is prohibited.

This privilege separation ensures that user applications cannot compromise system integrity. The transition between these modes occurs through controlled mechanisms such as system calls, interrupts, and exceptions.

On Linux, system calls (e.g., `read`, `write`, `mmap`, `clone`) form the stable boundary through which user programs request sensitive operations. The kernel validates, mediates, and executes the requests.

1.2.3 Virtual Memory and Address Translation

Hardware abstraction is tightly linked to the **virtual memory system**, which enables each process to operate as though it has continuous access to a unified address space. In reality, addresses are translated through a multi-level paging structure under kernel control and enforced by hardware Memory Management Units (MMUs).

Key components include:

- **Virtual Address (VA):** Address used by program instructions.
- **Physical Address (PA):** Actual memory cell location.
- **Page Tables:** Hierarchical structures mapping VA to PA.
- **TLB (Translation Lookaside Buffer):** Hardware cache for recent address translations.

Virtual memory provides:

- Isolation of processes from one another.
- Ability to run programs larger than available physical memory.

- Efficient memory sharing between processes (e.g., shared libraries).
- Controlled mapping of memory-mapped devices (e.g., GPUs, NICs).

1.2.4 Device Abstraction and Drivers

Hardware devices expose their functionality through registers, buffers, and control flags. Direct access to these components would require privileged instructions and would bind software to specific hardware models. To avoid this, Linux implements **device drivers**, typically within the kernel, which translate high-level operations into hardware-specific control sequences.

Examples:

- Storage devices abstracted through block device interfaces.
- Network interfaces abstracted through standardized socket APIs.
- GPUs exposed through memory-mapped command buffers and kernel-mediated execution queues.

From the perspective of a low-level programmer, device access often occurs through:

- `ioctl` calls for issuing device-specific control operations.
- Memory-mapped I/O regions created with `mmap`.
- Kernel modules when running in privileged context.

1.2.5 Architectural Feature Discovery

Modern processors include optional ISA extensions (e.g., AVX-512, SHA, AES-NI). Linux exposes available hardware capabilities via:

- `/proc/cpuinfo`
- `cpuid` instruction (for privileged feature probing in user-space where allowed)
- ELF HWCAP flags used by dynamic linkers for selecting optimized libraries.

Correctly detecting feature availability is critical for writing portable, optimized software, particularly in high-performance computing and cryptographic workloads.

1.2.6 Significance for Low-Level Programming

Hardware abstraction does **not** eliminate the need to understand the underlying hardware. Instead, it requires the programmer to distinguish between:

- **Architecturally guaranteed behavior** (e.g., paging semantics, system call ABI).
- **Microarchitectural performance characteristics** (e.g., cache locality, pipeline stalls, TLB pressure).
- **Kernel-mediated resource policies** (e.g., NUMA scheduling, huge page allocation).

Effective low-level development demands awareness of how abstractions influence latency, memory consistency, concurrency, and instruction throughput.

1.3 The Role of the Operating System

The **operating system (OS)** functions as the authoritative control layer between application software and underlying hardware. In modern x86-64 platforms, the OS is responsible for managing processor privilege boundaries, orchestrating resource allocation, maintaining isolation between programs, and providing a stable execution environment across diverse hardware implementations. Without the OS, applications would require direct access to hardware interfaces, leading to instability, lack of portability, and unsafe memory interaction. Linux, as the dominant system in server, cloud, and embedded environments, provides a well-defined separation between **user space** and **kernel space**, enforced through CPU privilege rings and controlled transitions.

1.3.1 Process Abstraction and Execution Management

The OS presents the illusion that each program executes independently with dedicated processor time and memory. In reality, the kernel manages **process scheduling**, dispatching threads across available logical cores and balancing workloads according to scheduling classes, CPU affinity, and energy constraints.

Modern Linux kernels use a **completely fair scheduling** (CFS) model, based on proportional time distribution, allowing predictable latency while maintaining high system throughput. On systems with heterogeneous cores or NUMA topologies, scheduling extends to locality-aware placement to reduce memory access cost.

Each process receives a **virtual address space**, isolating it from other processes and the kernel itself. The OS maintains this mapping through page tables and coherently coordinates updates with the MMU.

1.3.2 Memory Management and Protection

The OS supervises allocation, mapping, and revocation of memory regions. It ensures that:

- Instructions in user space cannot modify kernel memory.
- Each process has exclusive ownership of its writable regions.
- Shared memory is explicitly requested and controlled.

Linux provides multiple allocation mechanisms:

- **brk** and **sbrk** for expanding the heap.
- **mmap** for flexible mapping of anonymous memory or files.
- Transparent support for **huge pages** to reduce TLB pressure.
- **Copy-on-write** semantics to optimize process creation and shared library mappings.

The kernel enforces **access control** through page permissions (read, write, execute), which directly affects executable code layout and memory safety.

1.3.3 System Calls and the Privileged Interface

User-mode applications cannot perform sensitive operations directly. Instead, they request services through **system calls**, which transition execution to kernel mode via controlled traps (`syscall` instruction on x86-64).

System calls handle operations such as:

- File and device I/O.
- Thread creation, signaling, and synchronization.

- Network communication.
- Memory mapping and page protection changes.
- Process lifecycle management.

This interface forms the **stable contract** between user applications and the kernel. Its correctness and binary compatibility enable long-term software portability.

1.3.4 Device Control and I/O Coordination

Hardware devices present complexity and variability in register layout, signaling models, and DMA mechanisms. The OS abstracts these differences via **device drivers**, which execute in kernel space and expose uniform access interfaces to user applications.

Linux provides:

- Block device abstraction for disks and SSDs.
- Network stack abstraction for NICs.
- Character devices for serial and peripheral interfaces.
- Unified device naming under `/dev` and enumeration under `/sys`.

Drivers cooperate with the kernel's **interrupt handling** and **I/O scheduling** frameworks, ensuring fairness, performance, and predictable latency.

1.3.5 Security, Isolation, and Execution Integrity

The OS enforces **execution integrity** through mechanisms including:

- Access control via UNIX permissions, ACLs, and capabilities.

- Mandatory security frameworks such as SELinux or AppArmor.
- ASLR (Address Space Layout Randomization) to increase exploit complexity.
- Seccomp filters to restrict system calls.

These controls are continuously tightened in modern Linux distributions due to expanded multi-tenant and containerized deployment environments.

1.3.6 Significance for Low-Level Programming

For low-level work, the OS is not a mere runtime supervisor — it is a defining participant in software behavior. Performance, safety, and predictability depend on understanding how the OS:

- Schedules execution on cores.
- Manages caches and memory mappings.
- Arbitrates access to hardware through drivers.
- Uses system calls and synchronization primitives to coordinate concurrency.

Developing performant and correct low-level software requires reasoning not only about the ISA and microarchitecture, but also about the policies and constraints imposed by the Linux kernel.

1.4 Where Assembly Fits in the Stack

Assembly language occupies a distinct position between high-level languages and raw machine instructions. It provides a symbolic, human-readable representation of the operations defined by the Instruction Set Architecture (ISA) while exposing direct control over registers, memory addressing modes, calling conventions, and instruction sequencing. In the software stack on Linux x86-64 systems, assembly forms the lowest-level programming interface that remains practical for direct human engineering.

1.4.1 Machine Code and Architectural Semantics

At runtime, processors execute **machine code**, a sequence of binary opcodes encoded according to the ISA. These encodings specify the exact operations performed by the CPU's execution units. Assembly language exists as a **one-to-one or near one-to-one textual representation** of these binary encodings. Each assembly instruction maps to one or more machine instructions after assembler translation.

For example:

```
add rax, rbx
```

is a symbolic representation of an operation that adds the value in RBX to RAX, affecting general-purpose registers and condition flags according to the architectural definition.

The role of assembly is to express machine code precisely, without the abstraction layers introduced by compilers.

1.4.2 Assembly and High-Level Languages

High-level languages (e.g., C, C++, Rust) describe computation in terms of expressions, control flow constructs, and data abstractions. Compilers map these abstractions into assembly

before further translation to machine code. This mapping is not purely mechanical; compilers perform:

- Instruction selection
- Register allocation
- Optimization for pipeline throughput and cache locality
- Inlining and elimination of redundant operations
- Vectorization where hardware supports AVX or SVE-equivalent extensions

Thus, assembly appears as an **intermediate interface** that exposes hardware capabilities while still remaining writable and analyzable by humans. Understanding assembly is essential to interpreting compiler output, diagnosing performance bottlenecks, and verifying correctness in low-level systems work.

1.4.3 Assembly in System-Level Software

On Linux x86-64, assembly commonly appears in:

- **System call paths**, where user-space transitions into kernel mode through instructions such as `syscall`.
- **Context switch and interrupt entry code**, which must manipulate processor state with precise control.
- **Boot loaders and firmware**, where no higher-level runtime is yet available.
- **Optimized math, cryptography, hashing, and compression routines**, where careful instruction scheduling and register usage yield significant performance benefits.

- **Runtime library internals**, such as those found in `libc` and threading implementations.

These areas demand predictable execution behavior at the instruction level, which high-level languages cannot always guarantee without explicit constraints.

1.4.4 Assembly and the ABI Boundary

The **Application Binary Interface (ABI)** defines how functions receive arguments, return values, reference stack frames, and preserve registers across calls. On x86-64 Linux, the System V ABI specifies:

- Argument passing in registers (`RDI`, `RSI`, `RDX`, `RCX`, `R8`, `R9`).
- Stack alignment and calling sequence.
- Callee-saved and caller-saved register conventions.

Assembly is the language where the ABI is most visible and most strictly enforced. Incorrect handling of the ABI results in silent corruption of registers, unpredictable stack state, and undefined program behavior. Therefore, writing correct assembly requires precise adherence to ABI rules.

1.4.5 Why Assembly Still Matters

Despite advances in compiler technology, assembly remains relevant because:

- **Performance-critical inner loops** may demand manual instruction scheduling beyond compiler heuristics.
- **Hardware features** (e.g., new vector instructions, model-specific registers, serialization instructions) may not yet be exposed by compilers.

- **Low-level correctness** in operating systems, hypervisors, or embedded runtimes requires exact manipulation of architectural state.
- **Security hardening** sometimes necessitates hand-written control-flow and data sanitization sequences.

In modern systems, understanding assembly is less about writing entire programs in it and more about **reading, analyzing, optimizing, and interfacing with hardware behavior where precision is required.**

Chapter 2

Internal Structure of x86-64 CPUs

2.1 Execution Units

Modern x86-64 CPUs employ deeply parallel internal structures to exploit instruction-level parallelism beyond the sequential semantics defined by the ISA. At the core of this design are **execution units**—specialized hardware blocks responsible for performing arithmetic, logical, memory, and control operations. While the ISA describes *what* operations exist, execution units define *how* those operations are physically carried out and at what throughput and latency.

2.1.1 Role in the Out-of-Order Pipeline

Contemporary x86-64 processors operate as **out-of-order superscalar** machines. Instructions are decoded, renamed, and scheduled independently of their original program order. Once operands are ready, instructions are dispatched to execution units that work concurrently. After computation, results are committed in program order to preserve architectural correctness. The execution units are fed by a scheduling system that continuously identifies ready

instructions, enabling fine-grained exploitation of parallelism. This allows multiple instructions to complete during the same cycle, even though software expresses them sequentially.

2.1.2 Functional Classes of Execution Units

Execution units are not uniform; each is optimized for a specific class of operations. Typical categories include:

Execution Unit Type	Role	Example Operations
Integer ALUs	Arithmetic and logical operations on general-purpose registers	ADD, SUB, XOR, shifts
Address Generation Units (AGUs)	Compute effective memory addresses	Base+Index*Scale+Offset addressing
Floating-Point Units (FPUs)	IEEE 754-compliant scalar arithmetic	ADDSD, MULSS, SQRTSD
Vector / SIMD Units	Parallel arithmetic on packed data	SSE, AVX, AVX2, AVX-512 instructions
Branch Units	Predict, evaluate, and resolve control flow	CMP, TEST, JMP, conditional branches
Load/Store Units	Memory reads and writes between registers and cache hierarchy	MOV rax, [rbx], MOV [rcx], rdx

A given core contains multiple units of each type, allowing simultaneous execution of mixed instruction classes.

2.1.3 Port-Based Issue Architecture

In current x86-64 implementations (e.g., AMD Zen and Intel Core microarchitectures), execution units are organized across **ports**. Each port provides access to a subset of execution units. Scheduling logic assigns micro-operations to ports dynamically.

For example:

- Multiple ALUs across separate ports allow several integer instructions to retire per cycle.
- Separate load and store ports allow sustained memory bandwidth when combined with adequate cache locality.
- Vector units occupy ports that differ in width and throughput depending on AVX generation.

Understanding port usage is essential for low-level optimization, as bottlenecks arise when too many instructions compete for the same port or unit class.

2.1.4 Latency and Throughput Considerations

Execution units have two performance characteristics:

- **Latency:** The number of cycles required for a result to become available.
- **Throughput:** The number of operations that can be completed per cycle.

For example, a floating-point multiply may have a latency of several cycles but a throughput of one per cycle, allowing sustained pipelined computation when operand dependencies are minimized. Conversely, dependency chains or misaligned loads can significantly limit parallelism.

Low-level programmers must distinguish between:

- **Data-dependent chains**, which are limited by latency.
- **Independent parallel operations**, which scale with throughput.

2.1.5 Implications for Low-Level Software

For performance-critical programming on Linux x86-64 systems:

- Instruction sequences must minimize unnecessary dependencies.
- Loops should be structured to utilize vector and integer units concurrently.
- Memory access patterns must align with AGU and load/store unit throughput capacities.
- Branch behavior must be predictable to avoid misprediction penalties that stall execution pipelines.

Software that disregards execution unit behavior often becomes bound not by clock frequency but by port contention, cache stalls, or dependency serialization.

2.2 Decode Engine & Micro-Operations

In modern x86-64 processors, the instructions executed by software are not processed directly in their original architectural form. Instead, complex variable-length x86-64 instructions are transformed internally into simpler, fixed-format operations known as **micro-operations** (**μops**). This translation enables advanced parallel scheduling, out-of-order execution, and pipeline optimizations that are not possible when operating directly on the architectural instruction stream.

2.2.1 Complexity of the x86-64 Instruction Format

The x86-64 ISA retains backward compatibility with decades of architectural evolution. Instructions vary in length, operand specification, addressing modes, and implicit state effects. While this flexibility is beneficial for code density and compatibility, it complicates direct hardware execution.

To manage this complexity, the **decode engine** parses instruction bytes and interprets prefixes, opcodes, ModR/M encodings, displacement constants, and immediate operands. The decoder identifies the operation class and produces one or more corresponding μops that express the instruction's semantics in a hardware-friendly representation.

For example, an instruction that performs a memory-to-register arithmetic operation may decode into:

- One μop for effective address computation.
- One μop for the load.
- One μop for the arithmetic operation.

2.2.2 Micro-Operations (μ ops)

A **micro-operation** is an atomic hardware-level operation representing the smallest executable step in the pipeline. μ ops specify explicit source operands, destination registers, and required functional units. They remove the implicit state and overlapping functionality common in high-level x86 instructions.

Key properties:

- Fixed internal format.
- Designed to map directly to execution units.
- Enable independent scheduling and parallel dispatch.
- Facilitate register renaming by separating architectural and physical state.

Whereas x86 instructions express semantics, μ ops express **execution structure**.

2.2.3 The μ op Cache and Decode Bypass

Decoding x86 instructions is expensive relative to pipeline cycle time. To prevent decoding from becoming a bottleneck, modern x86-64 CPUs include a **μ op cache** (sometimes referred to as a decoded instruction cache). When the CPU encounters previously executed instruction sequences, it retrieves pre-decoded μ ops directly, bypassing the decode stage entirely.

This improves:

- Power efficiency by reducing front-end work.
- Throughput by allowing the pipeline to supply more instructions per cycle.
- Performance predictability in tight loops and frequently executed code paths.

Code alignment, loop structure, and branch patterns influence the effectiveness of μ op caching.

2.2.4 Fusion and Decomposition

Some instructions can be **fused** into a single μop if their semantics allow combined execution (e.g., compare-and-branch sequences). Conversely, some complex operations may be **decomposed** into multiple μops , particularly under SIMD or memory-intensive patterns.

Fusion improves throughput by reducing scheduler and execution-unit pressure.

Decomposition exposes parallelism where dependency chains permit.

Understanding which forms fuse and which do not is essential for performance tuning and predictable latency behavior.

2.2.5 Interaction with Out-of-Order Scheduling

Once instructions are decomposed into μops , they enter the out-of-order execution pipeline:

- The register renaming stage maps architectural registers to physical ones.
- The scheduler tracks operand readiness.
- μops are dispatched to execution units when dependencies resolve.

The decode stage therefore functions as the **interface between the complex ISA and the high-performance dynamic execution core**.

2.2.6 Significance for Low-Level Development

For system-level and performance-critical programming:

- Seemingly simple instructions may translate into multiple μops , affecting throughput.
- Memory-based arithmetic instructions often cost more μops than register-only equivalents.

- Branch layout influences whether μ ops are served from cache or require decoding.
- Instruction sequences should be structured for predictable decode and scheduling behavior.

Thus, performance-aware assembly and compiler-guided optimization must consider **decoded μ op count and front-end throughput**, not merely instruction count.

2.3 Reorder Buffer & Out-of-Order Execution

Modern x86-64 CPUs achieve performance not through higher clock speeds alone, but through the ability to execute instructions **out of program order** while still producing results consistent with the logical, sequential semantics defined by the ISA. This capability is enabled by two core mechanisms: **register renaming** (covered earlier) and the **Reorder Buffer (ROB)**, which together allow instructions to execute as soon as their operands are ready, regardless of their original position in the instruction stream.

2.3.1 Motivation for Out-of-Order Execution

Instruction streams generated by compilers often contain unavoidable **latency-producing operations**, such as cache misses, memory dependencies, and multi-cycle arithmetic instructions. If the processor were forced to wait for each instruction to complete before issuing the next, utilization of execution units would be extremely low.

Out-of-order (OoO) execution allows the CPU to:

- Continue executing independent instructions while others are stalled.
- Maximize occupancy of execution units.
- Reduce the visible latency of slow operations.

This yields increased throughput without altering program correctness.

2.3.2 Structure and Role of the Reorder Buffer

The **Reorder Buffer (ROB)** is a hardware structure that tracks in-flight instructions— μ ops that have been issued but have not yet committed their results to the architectural state. Each entry in the ROB stores:

- The original program order index of the instruction.
- Destination register identifiers (architectural $\rightarrow$ physical mapping reference).
- Status flags indicating whether execution has completed.
- Exception or fault information if generated during execution.

The ROB maintains the illusion that the program executes strictly in sequence, even though the internal execution order may differ.

2.3.3 Execution and Result Commitment

The execution flow proceeds in three conceptual stages:

1. **Dispatch:** μ ops are issued into the ROB and made available to the scheduling window.
2. **Execute:** When operands are ready, μ ops are sent to execution units and produce results stored in physical registers.
3. **Retirement (Commitment):** μ ops retire in *program order* when their execution has completed and no faults are pending. Only at this point do results become architecturally visible.

This structured separation between execution and commitment prevents speculation and parallelization from altering logical outcomes.

2.3.4 Handling Speculation and Mispredictions

The CPU aggressively speculates on legal future states—such as branch outcomes—executing instructions before certainty is established. If speculation is later proven incorrect:

- The ROB **invalidates** all μ ops newer than the mispredicted instruction.

- The architectural state is restored using the committed state in the ROB.
- Execution restarts at the correct instruction boundary.

This rollback capability is a defining requirement for speculative and out-of-order pipelines.

2.3.5 Memory Ordering and the Load/Store Queue

While register dependencies can be resolved using register renaming and the ROB, **memory operations must preserve architectural ordering constraints**. To enforce this, CPUs maintain **load and store queues** parallel to the ROB:

- The **store queue** holds pending writes until commit, preserving program order.
- The **load queue** resolves possible dependencies between speculative loads and earlier unresolved stores to the same address.

This ensures memory semantics remain consistent with the architectural model while still enabling aggressive parallel memory access when safe.

2.3.6 Practical Implications for Low-Level Programming

For performance-sensitive development on Linux x86-64 systems:

- Long dependency chains (e.g., repeated arithmetic on the same register) constrain OoO execution and reduce pipeline utilization.
- Independent computations should be structured to allow the ROB and scheduler to overlap work.
- Memory access patterns must reduce aliasing and false dependencies to avoid load/store queue stalls.

- Misaligned or unpredictable branches cause ROB flushes, resulting in lost parallelism.

Effective low-level optimization therefore requires **designing instruction sequences that maximize instruction independence and minimize speculative rollback penalties.**

2.4 Register Rename Logic

The **register rename logic** is central to the out-of-order execution model in modern x86-64 CPUs. Although the x86-64 ISA exposes a fixed set of **architectural registers** (such as RAX, RBX, RSP, vector registers, and status flags), the processor internally maintains a significantly larger pool of **physical registers**. Register renaming allows the CPU to map architectural registers to these physical registers dynamically, eliminating false dependencies and enabling instructions to execute in parallel even when they appear sequential in software.

2.4.1 Architectural vs. Physical Registers

Architectural registers represent the stable, programmer-visible state defined by the ISA. However, instructions may temporarily overwrite registers before previous results are needed. To support parallel execution without overwriting still-needed values, the processor assigns a **unique physical register** to each instruction's destination.

This mapping ensures:

- Multiple in-flight instructions can use logical registers without interfering with one another.
- The pipeline does not stall due to write-after-read (WAR) or write-after-write (WAW) dependencies, which are artificial and not required by program semantics.

True dependencies—read-after-write (RAW)—are the only constraints that must be enforced.

2.4.2 Rename Table and Allocation

The rename logic uses two key structures:

Structure	Function
Register Alias Table (RAT)	Maps architectural register identifiers to current physical registers.
Free List	Tracks available physical registers for new allocations.

When an instruction that writes a register enters the pipeline:

1. The RAT assigns a **new physical register** for the destination.
2. The previous physical register remains valid until the instruction commits.
3. The ROB tracks which physical register version corresponds to each instruction for recovery and commit ordering.

This allows multiple versions of a register to coexist simultaneously.

2.4.3 Retirement and Version Release

When a μop **commits** in program order, its physical register becomes the authoritative architectural value. At that moment:

- The prior physical register version that it replaced is added back to the free list.
- The rename state becomes stable and recoverable.

This ensures that precise exceptions and branch misprediction recovery can restore architectural register state by rolling back to the correct RAT and physical register mappings.

2.4.4 Elimination of False Dependencies

Without renaming, the apparent reuse of registers in software would serialize execution.

Renaming removes false dependencies:

- **WAW hazard:** Two instructions writing the same architectural register get distinct physical registers.
- **WAR hazard:** A read of a logical register can safely proceed even if a future instruction will overwrite the register, because renaming ensures they refer to different physical registers.

Only **RAW hazards** remain, and these are tracked precisely via dependency graphs in the execution scheduler.

2.4.5 Interaction with Branch Speculation

Because speculative instructions also receive renamed registers:

- Branch mispredictions do **not** corrupt architectural state.
- The RAT is checkpointed at speculation boundaries.
- When speculation fails, the RAT reverts to a prior checkpoint, discarding later renamings.

This allows deep speculation without risking irreversible register updates.

2.4.6 Impact on Low-Level Performance

For performance-sensitive assembly and compiler-guided optimization:

- Maintaining **independent instruction streams** helps the rename logic expose parallelism to the scheduler.
- Long dependency chains force the scheduler to wait on single register lineages, reducing throughput.
- Excessive register pressure may cause spilling, introducing additional memory accesses with higher latency.

Well-structured code distributes work across registers to maximize the effective use of physical register resources.

2.4.7 Summary

Register renaming transforms a small logical register file into a far larger dynamic execution state, enabling:

- Parallel execution of independent operations.
- Precise recovery from speculation failures.
- Enforcement of architectural semantics while allowing microarchitectural freedom.

It is a foundational mechanism enabling the high-performance behavior of modern x86-64 processors.

Chapter 3

The Memory Hierarchy

3.1 L1 / L2 / L3 Cache Layers

The performance of x86-64 processors is fundamentally shaped by the speed gap between the CPU's execution units and main memory. While arithmetic operations complete in a few cycles, accessing DRAM may require hundreds. To bridge this disparity, CPUs employ a **multi-level cache hierarchy**, storing recently accessed data in progressively larger and slower layers of memory located physically closer to the processor core.

Modern Linux-capable x86-64 processors universally implement a **three-level cache hierarchy** optimized for spatial locality, temporal locality, and simultaneous multi-core workloads.

3.1.1 L1 Cache: Fast, Private, and Specialized

The **L1 cache** is the smallest and fastest cache layer, located closest to the core pipeline. It is typically split into:

- **L1D (Data Cache)** for load and store operations

- **L1I (Instruction Cache)** for decoded instruction fetch

Typical characteristics:

Property	L1 Data / Instruction Cache
Size	~32 KB per core (separate D/I caches)
Latency	~4 cycles
Scope	Private to each core

The L1 cache is optimized for **low latency**, ensuring the highest probability that tightly looped or frequently reused data is available immediately to the execution pipeline.

3.1.2 L2 Cache: Larger and Still Private

The **L2 cache** serves as an intermediate layer, backing both L1D and L1I caches.

Property	L2 Cache
Size	256 KB to 1 MB per core (microarchitecture dependent)
Latency	~10–14 cycles
Scope	Private to each core

The L2 cache balances speed with capacity. It acts as a **buffer zone**, reducing the frequency with which cores must access the shared L3 cache or main memory, thereby lowering contention between concurrent threads on the same processor.

3.1.3 L3 Cache: Shared and Designed for Multi-Core Coordination

The **L3 cache** (also known as the Last Level Cache, LLC) is typically shared across all cores in a CPU package.

Property	L3 Cache
Size	8 MB to >64 MB, depending on core count and generation
Latency	~30–50 cycles
Scope	Shared among cores on the same die

The shared nature of L3 allows:

- Efficient data exchange among threads running on different cores.
- Reduced DRAM traffic by servicing accesses at the socket level.
- Cache coherence protocols to maintain unified memory visibility.

The L3 cache is generally **inclusive or non-inclusive** depending on vendor design. Inclusive caches replicate data from lower levels for faster coherence checks, while non-inclusive caches avoid redundancy to maximize capacity.

3.1.4 Relationship to DRAM and Access Latency

Cache performance directly influences execution timing:

Level	Approx. Latency	Notes
L1	~4 cycles	Near-instant for dependent instructions
L2	~10–14 cycles	Slightly visible but generally tolerable
L3	~30–50 cycles	Can stall pipelines if heavily contested
DRAM	~200+ cycles	Requires prefetching and locality-aware layout

When a request misses at all cache layers, the processor must stall until data is fetched from main memory, significantly impacting performance.

3.1.5 Implications for Low-Level Programming

Performance-sensitive software on Linux x86-64 must consider cache hierarchy explicitly:

- Arrange data to maximize **spatial locality**, placing frequently accessed values near each other in memory.
- Ensure **temporal locality**, reusing recently accessed data wherever possible.
- Align data structures to cache line boundaries to avoid false sharing in multi-threaded workloads.
- Design parallel algorithms with awareness of shared vs. private cache layers.

Cache-conscious programming often determines whether performance matches theoretical execution throughput or becomes bottlenecked by memory fetch delays.

3.2 Cache Lines and Latency

The behavior of the cache hierarchy is governed not only by its layered structure, but also by its **fundamental unit of storage**: the **cache line**. A cache line is the smallest block of memory that can be transferred between memory hierarchy levels. Understanding cache line size, alignment, and access patterns is essential for predicting performance characteristics on x86-64 systems.

3.2.1 Cache Line Size

On modern x86-64 processors, a **cache line is typically 64 bytes**. Whenever data is read from memory, the CPU does not fetch a single byte or word—it fetches an entire 64-byte aligned block containing that data.

This design exploits **spatial locality**: the assumption that nearby memory addresses are likely to be accessed together, particularly in loops and contiguous data structures.

For example, when accessing `array[i]`, elements `array[i+1]`, `array[i+2]`, etc., often reside in the same cache line and can be accessed with no additional latency. Conversely, accessing elements that fall on different cache lines may incur additional fetch overhead.

3.2.2 Cache Misses and Latency Sources

Caches reduce access latency, but not all accesses hit the nearest cache. Misses fall into several types:

Miss Type	Cause	Typical Example
Cold Miss	First access to data not yet cached	Initial traversal of a large dataset
Capacity Miss	Working set exceeds cache size	Operations on data larger than available cache space

Miss Type	Cause	Typical Example
Conflict Miss	Multiple addresses map to the same cache set	Unfavorable data alignment or stride patterns

Latency increases sharply as requests move further from the core:

Level	Approximate Access Latency	Effect on Execution
L1	~4 cycles	Typically hidden by instruction overlap
L2	~10–14 cycles	Noticeable in dependency chains
L3	~30–50 cycles	Can stall pipelines if frequent
DRAM	~200+ cycles	Usually becomes the dominant performance bottleneck

When accessing memory with insufficient locality, execution units may stall while waiting for data, regardless of available computational throughput.

3.2.3 Stride Patterns and Cache Efficiency

Memory access patterns strongly influence cache behavior:

- **Unit stride (contiguous access)** optimizes caching because sequential elements occupy the same or adjacent cache lines.
- **Large stride access** (e.g., accessing every Nth element) can defeat spatial locality, causing repeated cache misses.
- **Irregular or pointer-heavy traversal** (e.g., linked structures) prevents the CPU from effectively prefetching data and often results in increased memory stalls.

For high-performance workloads, designing data structures around contiguous layouts often yields significant gains.

3.2.4 Alignment and Alias Boundaries

Because cache lines are aligned to 64-byte boundaries, performance depends on how data structures align with these boundaries:

- If two frequently accessed fields share a cache line, access may be efficient due to locality.
- If independent data accessed by different threads shares a cache line, **false sharing** can occur, leading to unnecessary cache invalidations—addressed in the next section.
- Aligning critical data to cache line boundaries ensures predictable and isolated memory access behavior.

Compilers and allocators typically provide alignment guarantees, but low-level programmers often enforce explicit alignment in performance-sensitive code.

3.2.5 Prefetching and Latency Hiding

Modern processors automatically **prefetch** data ahead of demand when sequential access patterns are detected. However, prefetching mechanisms fail when:

- Access patterns are irregular or data-dependent.
- Working sets exceed cache capacity.
- Cache conflicts cause eviction before reuse.

In such cases, explicit software prefetch instructions or loop restructuring may be required to maintain throughput.

3.2.6 Implications for Low-Level Programming on Linux

Effective low-level memory performance requires:

- Structuring data for contiguous access.
- Minimizing cross-cache-line accesses inside tight loops.
- Ensuring memory alignment to cache line boundaries.
- Avoiding unpredictable or pointer-chasing data dependencies.
- Measuring performance using hardware counters (`perf`) to identify cache miss sources.

At scale, performance is rarely limited by arithmetic throughput; it is most often limited by how efficiently data moves between levels of the memory hierarchy.

3.3 NUMA and Multi-Core Memory Effects

As core counts increase, the memory hierarchy must scale to support parallel execution without creating bandwidth bottlenecks or latency collapse. In modern x86-64 systems—especially multi-socket servers and high-core-count workstations—memory is organized under a **Non-Uniform Memory Access (NUMA)** model. Under NUMA, memory is physically divided into **nodes**, each associated with a subset of CPU cores. Accessing memory local to a core's own node is faster than accessing memory attached to a different node.

This architecture preserves scalability but requires programmers to understand the performance implications of memory placement across nodes.

3.3.1 Memory Nodes and Locality

Each NUMA node contains:

Component	Description
A set of CPU cores	Execution units associated with the node
One or more memory controllers	Interfaces to DRAM modules
A portion of system DRAM	Local memory bank(s)

A core accessing its node's local memory benefits from lower latency and higher sustained bandwidth. Access across nodes requires traversal through interconnect links (e.g., AMD Infinity Fabric or Intel UPI), increasing latency and reducing bandwidth.

Typical latency relationship:

Memory Access Type	Approximate Relative Latency
L1 / L2 / L3 Cache	Lowest
Local DRAM	Baseline (1×)
Remote DRAM	~1.3× to 2.5× depending on architecture and load

For memory-intensive workloads, this difference is often performance-critical.

3.3.2 Multi-Core Workloads and Data Placement

On Linux, threads may migrate between cores for load balancing. If a thread migrates to a different NUMA node while retaining pointers to memory allocated on the original node, it incurs **remote memory access penalties**.

Therefore, high-performance multi-threaded code must ensure:

- Memory is allocated **on the same node** where the accessing threads execute.
- Threads remain **affined** to specific cores or NUMA nodes when appropriate.

Linux provides **NUMA policy controls** (e.g., `numactl`, `mbind`, and process affinity masks) to explicitly bind computation and memory.

3.3.3 Interconnect Contention and Bandwidth Sharing

When many cores access remote memory or when several memory-intensive threads share bandwidth on the same node, contention arises. This manifests as:

- Increased access latency
- Decreased memory throughput
- Unpredictable performance scaling when additional threads are added

NUMA-aware scheduling and memory allocation can mitigate this by distributing memory and work evenly across nodes.

3.3.4 Effects on Synchronization and Shared State

When two threads running on different NUMA nodes share the same cache line:

- Memory coherence protocols must coordinate updates across nodes.
- This introduces **cross-node coherence traffic**, which is significantly slower than intra-node communication.
- False sharing and shared lock variables become especially expensive.

Contention on shared data structures frequently becomes the dominant scaling limiter on multicore NUMA systems.

NUMA-efficient synchronization strategies include:

- Per-node data replication
- Sharded queues or hash tables
- Avoiding cache-line-sharing in lock variables and frequently-updated counters

3.3.5 NUMA on Linux: Allocation and Scheduling Behavior

Linux employs **first-touch allocation**: memory is physically allocated to the NUMA node of the core that first writes to it. Consequently:

- Initialization loops should run on the same cores that will later use the data.
- Shared data should be initialized in parallel by the threads that will operate on their respective regions.

Failure to respect first-touch causes silent performance degradation that is often misdiagnosed as algorithmic inefficiency.

3.3.6 Implications for Low-Level Programming

To write NUMA-efficient code on Linux:

- Ensure data locality matches execution locality.
- Avoid unnecessary cross-node synchronization and shared state.
- Measure node-local vs. remote memory access using profiling tools (e.g., `perf`, `numastat`).
- Design data structures and workloads for **partitioned parallelism**, not shared access.

In sustained multi-core workloads, performance is frequently determined not by instruction throughput but by how effectively memory access patterns align with the NUMA topology.

3.4 Cache Miss Penalties

The performance characteristics of modern x86-64 processors depend heavily on the assumption that frequently accessed data resides in the upper levels of the cache hierarchy. When this assumption fails, the processor must retrieve data from progressively slower memory layers. Each **cache miss** introduces a delay measured in CPU cycles, during which execution units may stall or reorder surrounding instructions to mitigate latency. The scale of these penalties makes cache behavior a dominant factor in real-world performance.

3.4.1 Cost Escalation Across the Hierarchy

Memory access time increases sharply as data travels further from the core:

Location	Approx. Latency (Cycles)	Access Scope
L1 Cache	~4 cycles	Private per core
L2 Cache	~10–14 cycles	Private per core
L3 Cache (LLC)	~30–50 cycles	Shared across cores
DRAM (Main Memory)	~200+ cycles	System-wide
Remote DRAM (NUMA)	~250–450+ cycles	Cross-node

Even a small increase in miss frequency can dominate total execution time. For example, a tight loop with a dependency chain cannot progress if each load must retrieve data from DRAM.

3.4.2 Stall Behavior and Instruction-Level Parallelism Limits

Out-of-order execution can hide some memory latency by scheduling independent instructions around a stall. However, this is only effective when:

- Independent instructions are available in the pipeline window.

- The reorder buffer (ROB) does not fill while waiting on a memory response.
- The miss does not block core-wide structures (e.g., load/store queue exhaustion).

Once critical dependency chains require data that has not yet arrived, the pipeline must **wait**, halting forward progress regardless of available ALUs or vector units.

3.4.3 Memory-Level Parallelism

Modern cores issue multiple outstanding memory requests simultaneously, leveraging **memory-level parallelism (MLP)**. However, MLP is limited by:

- Load/store queue depth
- Hardware prefetch capability
- Cache miss concurrency constraints

Sequential pointer chasing, for instance, defeats MLP because each load depends on the result of the previous one, forcing strictly serialized memory access.

Algorithms structured to operate on arrays or blocks instead of linked structures can sustain much higher throughput.

3.4.4 Causes of Cache Miss Penalties

Cache misses occur for several primary reasons:

Miss Type	Description	Typical Cause
Cold Miss	First reference to data	Large streaming or startup behavior

Miss Type	Description	Typical Cause
Capacity Miss	Working set > cache capacity	High-dimensional or poorly tiled data
Conflict Miss	Addresses map to same cache sets	Poor alignment or stride access
Coherency Miss	Cache line invalidated by another core	Multi-threaded shared data update

Resolving penalties depends on identifying the type responsible and restructuring data access accordingly.

3.4.5 Measuring Cache Misses on Linux

Linux provides several tools that expose cache miss statistics at runtime:

- **perf stat** for overall L1/L2/L3/DRAM miss counts
- **perf record/report** to locate miss-heavy code regions
- **numastat** to identify local vs. remote DRAM access
- **lscpu** and `/sys/devices/system/node/` for NUMA topology inspection

Optimization begins by quantifying where misses occur and how severe they are relative to execution time.

3.4.6 Strategies to Reduce Cache Miss Impact

Low-level performance tuning requires designing data layouts and algorithms around the memory hierarchy. Key strategies include:

- **Maximizing spatial locality:** Arrange data contiguously and iterate in memory order.

- **Maximizing temporal locality:** Reuse values before eviction from cache.
- **Blocking / tiling algorithms:** Operate on sub-regions sized to fit within cache.
- **Structure flattening:** Replace pointer-heavy decomposed structures with compact arrays.
- **Prefetching:** Allow hardware (or software hints) to fetch data before it is needed.
- **Thread-core affinity:** Ensure access to memory local to the executing core or NUMA node.

These techniques shift execution costs back toward the L1/L2 regimes where performance is dominated by compute throughput instead of memory stalls.

Chapter 4

Registers and Flags

4.1 General Purpose Registers

In the x86-64 architecture, **general purpose registers (GPRs)** form the core of the programmer-visible execution state. These registers hold integer values, pointers, offsets, loop counters, condition outcomes, and function arguments. While the ISA exposes a familiar set of named registers, modern processors internally map these architectural registers to larger pools of **physical registers** through register renaming, as discussed previously. From the programmer's perspective, however, the architectural register set defines the interface for computation at the assembly level.

4.1.1 Register Width and Naming Hierarchy

x86-64 expands the legacy 32-bit register set to **64 bits per register**, introducing new register names and preserving backward-compatible partial-register access. Each register can be addressed at multiple granularities:

64-bit	32-bit	16-bit	8-bit (low)	8-bit (high, legacy only)
RAX	EAX	AX	AL	AH
RBX	EBX	BX	BL	BH
RCX	ECX	CX	CL	CH
RDX	EDX	DX	DL	DH
RDI	EDI	DI	DIL	—
RSI	ESI	SI	SIL	—
RBP	EBP	BP	BPL	—
RSP	ESP	SP	SPL	—
R8–R15	R8D–R15D	R8W–R15W	R8B–R15B	—

The upper-byte (AH, BH, CH, DH) registers persist only for backward compatibility; modern code avoids them because they introduce encoding penalties and reduce renaming flexibility.

4.1.2 Roles and Conventional Usage

While all GPRs are functionally similar, conventional usage patterns exist, especially under the **System V AMD64 ABI** used by Linux:

Register	Common Role
RAX	Return value register; often used for arithmetic accumulations
RBX	Callee-saved register, preserved across function calls
RCX	General counter or argument register; used in shift operations
RDX	Argument register; used for division remainder
RSI / RDI	Primary argument registers for string, memory copy, and compare operations
RBP	Frame pointer (optional; compilers may omit)
RSP	Stack pointer; strictly reserved for stack operations
R8–R11	Additional argument or temporary registers

Register	Common Role
R12–R15	Callee-saved general-purpose registers

These conventions ensure the interoperability of compiled code with libraries, system calls, and the kernel.

4.1.3 The Stack Pointer and Call Frame Structure

RSP points to the top of the stack and is implicitly referenced by `CALL`, `RET`, `PUSH`, and `POP` instructions. Maintaining correct stack alignment is essential:

- The System V ABI requires **16-byte alignment before `call`**.
- Stack misalignment may cause performance penalties and break SIMD call requirements.

RBP may serve as a **frame pointer**, but modern compilers frequently omit it to free an additional register for computation.

4.1.4 Pointer and Addressing Roles

Because x86-64 defines a flat, 64-bit virtual address space, any general-purpose register can legally hold pointers. However:

- Pointer arithmetic is typically performed using scaled-index addressing ($\text{Base} + \text{Index} \times \text{Scale} + \text{Displacement}$).
- RIP (instruction pointer) is not directly writable but supports **RIP-relative addressing**, improving position-independent code efficiency.

RIP-relative addressing is central to Linux shared libraries and dynamic linking.

4.1.5 Performance Considerations

Although all GPRs share identical architectural semantics, their **microarchitectural treatment may differ**:

- Instructions using AH/BH/CH/DH can block register renaming pathways.
- Partial-register writes (writing EAX while reading RAX) may introduce false dependencies unless handled carefully by the microarchitecture.
- Frequent stack pointer updates may limit pipeline scheduling when functions push and pop excessively.

Performance-critical code tends to:

- Favor full 64-bit register operations.
- Avoid partial-register dependencies.
- Minimize unnecessary stack traffic.

4.1.6 Significance for Low-Level Programming

A precise understanding of the GPR set enables:

- Efficient function interface design.
- Predictable stack and frame construction.
- Manual loop and data movement optimization.
- Correct interaction with system calls and kernel interfaces.

The general-purpose register file is not merely a storage resource; it is the **primary lever of performance, correctness, and expressiveness** in x86-64 assembly programming.

4.2 RFLAGS Bit Meanings

The **RFLAGS register** (often simply referred to as RFLAGS) contains a set of condition and control bits that reflect the outcome of arithmetic, control CPU behavior, and influence branching decisions. Unlike general-purpose registers, RFLAGS is not freely used for computation; instead, it expresses **processor state** relevant to conditional execution and exception handling.

Understanding RFLAGS is essential for interpreting the effects of arithmetic instructions, designing efficient branch logic, and ensuring correct behavior during low-level debugging and performance tuning.

4.2.1 Structure of RFLAGS

RFLAGS is a 64-bit register; however, only a subset of bits are architecturally meaningful. The lower bits are inherited from the legacy 32-bit EFLAGS and 16-bit FLAGS registers.

Key flag categories:

Category	Purpose
Status Flags	Reflect arithmetic and logical operation outcomes
Control Flags	Modify processor operational modes
System Flags	Interface with privilege and execution environment

4.2.2 Status Flags (Affect Conditional Instructions)

Status flags are set or cleared by most arithmetic and logical instructions:

Bit	Name	Meaning	Typical Use
CF	Carry Flag	Indicates unsigned overflow (carry/borrow)	Unsigned comparisons, multi-word arithmetic

Bit	Name	Meaning	Typical Use
ZF	Zero Flag	Set when result equals zero	Loop termination, equality checks
SF	Sign Flag	Reflects the sign bit of the result	Signed comparisons
OF	Overflow Flag	Indicates signed arithmetic overflow	Detecting range errors in signed math
PF	Parity Flag	Indicates even parity of the low byte	Legacy; rarely performance-relevant today
AF	Adjust Flag	Assists BCD operations	Retained for compatibility, seldom used

In practice, **CF**, **ZF**, **SF**, and **OF** are the primary flags used in modern conditional logic and branching.

4.2.3 Control and Direction Flags

Control flags modify how instructions operate:

Bit	Name	Effect
DF	Direction Flag	Determines iteration direction in string/memory instructions (forward = 0, backward = 1)
IF	Interrupt Flag	Enables or disables maskable hardware interrupts
TF	Trap Flag	Enables single-step debugging via exceptions

Ordinary application code rarely manipulates these directly, as the kernel controls interrupt-threading semantics.

4.2.4 System Flags (Privilege and Execution Context)

These flags are relevant to kernel-mode and virtualization:

Bit	Name	Purpose
IOPL	I/O Privilege Level	Controls access to hardware I/O ports
VM	Virtual-8086 Mode	Enables execution of legacy 16-bit code in protected mode
RF	Resume Flag	Controls debug exception behavior during instruction restart

These fields are **normally immutable in user space** under Linux. Attempts to change them are either ignored or cause a fault based on privilege level.

4.2.5 Flag Dependencies and Performance Considerations

Because status flags encode results of arithmetic, they introduce **data and control dependencies**:

- Instructions that read flags (e.g., JZ, JNZ, ADC, SBB) cannot execute until the producing instruction completes.
- Chains of flag-dependent instructions may limit out-of-order scheduling freedom.
- Replacing flag-dependent instructions with explicit register comparisons (when possible) may increase throughput in performance-critical loops.

Additionally, **partial flag updates** (e.g., INC does not modify CF) can lead to subtle correctness issues if flag state is assumed rather than explicitly re-established.

4.2.6 Practical Guidelines for Low-Level Programming

For robust and efficient use of RFLAGS in Linux system-level programming:

- **Do not assume flag state persists** across calls or operations; recompute using `CMP` or `TEST` where clarity is required.
- Favor **branchless forms** (e.g., `CMOV`, `setcc`) to reduce misprediction penalties when branch patterns are irregular.
- Be aware of **instruction selection effects**, as ALU operations and comparisons differ in which flags they modify.

The RFLAGS register is not merely a passive result indicator; it is a **control element** that shapes program flow, scheduling behavior, and execution predictability.

4.3 Special Registers (RIP, RSP, RBP)

While the general-purpose registers provide flexible storage for computation and data handling, several **special-purpose registers** serve defined structural roles in program execution. In x86-64 systems, **RIP**, **RSP**, and **RBP** govern instruction sequencing and stack-based function organization. These registers are central to function calls, stack frame layout, control flow, and position-independent execution in Linux environments.

4.3.1 Instruction Pointer (RIP)

The **RIP** register contains the address of the instruction currently being executed. Unlike general-purpose registers, RIP is **not directly writable** via standard arithmetic instructions. Instead, it changes implicitly:

- Sequential execution increments RIP after each instruction.
- Control flow instructions (`CALL`, `JMP`, conditional branches) update RIP explicitly.
- Exceptions, interrupts, and system calls push the old RIP and load a new one.

A defining feature of x86-64 is **RIP-relative addressing**, where instructions can reference memory relative to the current instruction location:

```
mov rax, [rip + offset]
```

This enables **position-independent code (PIC)**, which is required for:

- Shared libraries on Linux
- ASLR (Address Space Layout Randomization)
- Dynamically loaded runtime components

RIP-relative addressing reduces the need for absolute relocations, allowing ELF loaders and linkers to manage shared objects efficiently.

4.3.2 Stack Pointer (RSP)

The **RSP** register points to the top of the current thread's runtime stack. It is updated implicitly or explicitly by:

- `CALL` and `RET` instructions (push return address; pop return location)
- `PUSH` and `POP` instructions
- Manual stack frame adjustments via `ADD` / `SUB`

The System V AMD64 ABI used on Linux mandates:

- **16-byte alignment** of RSP before executing a `CALL` instruction.
- Misalignment may degrade performance or break ABI compliance for SIMD-using functions.

Because RSP is central to control flow recovery, any corruption of RSP typically results in **immediate program failure**, often through segmentation faults.

4.3.3 Base Pointer (RBP) and Stack Frames

The **RBP** register traditionally serves as the **frame pointer**, marking a stable reference within a function's stack frame. A typical prologue:

```
push rbp
mov  rbp, rsp
sub  rsp, <frame_size>
```

This establishes a structured frame where local variables, spilled registers, and call arguments are accessed at fixed offsets from RBP. However, modern compilers frequently employ **frame**

pointer omission (FPO) to free RBP for use as a general-purpose register, relying instead on RSP-relative addressing.

FPO increases available registers and improves instruction scheduling freedom, but complicates debugging or manual stack inspection unless Call Frame Information (CFI) metadata is present.

4.3.4 Stack Discipline and Calling Semantics

The correctness of function calls depends on maintaining consistent stack structure:

- Callers push return addresses via `CALL`.
- Callees optionally preserve RBP and adjust RSP for local storage.
- On return, the stack must be restored exactly to its pre-call state.

Failure to maintain stack integrity leads to subtle control-flow corruption, often detected only when returning from deep call chains.

4.3.5 Interaction with the Kernel and Context Switching

During a context switch, Linux saves and restores:

- RIP (execution resume point)
- RSP (user stack state)
- RBP and other callee-saved registers when necessary

This allows threads and processes to suspend and resume execution transparently, preserving program state as if continuously running.

4.3.6 Importance for Low-Level Development

Understanding these special registers is essential when:

- Writing assembly routines or hand-crafted calling conventions
- Implementing system call stubs or signal trampolines
- Designing coroutine, fiber, or green-thread runtime switch mechanisms
- Debugging segmentation faults and stack corruption
- Interpreting disassembly and stack traces

RIP defines *where* execution is; RSP defines *the structure supporting execution*; RBP defines *a reference anchor when stack-based addressing is used*. Together, they form the backbone of procedural program control flow on Linux x86-64.

4.4 SIMD Registers (XMM / YMM)

Modern x86-64 processors include **SIMD (Single Instruction, Multiple Data)** register sets that enable parallel operations on packed data. These registers support vectorized arithmetic, logical operations, data rearrangement, and memory movement. On Linux, SIMD is critical for high-performance workloads in graphics, numerical computing, cryptography, compression, signal processing, and machine learning inference.

The primary SIMD registers addressed in general-purpose code today are **XMM** and **YMM**, corresponding to **SSE** and **AVX/AVX2** instructions.

4.4.1 XMM Registers (128-bit SIMD)

The **XMM registers** are 128-bit registers exposed through SSE and SSE2–SSE4 instruction sets. There are 16 architectural XMM registers in 64-bit mode:

```
XMM0, XMM1, XMM2, ..., XMM15
```

Each XMM register can hold packed:

Data Type	Elements (per XMM register)
8-bit integers	16 elements
16-bit integers	8 elements
32-bit integers or floats	4 elements
64-bit integers or doubles	2 elements

Example operations (Intel syntax):

```
movaps  xmm0, xmm1      ; Move aligned 128-bit vector
addps   xmm0, xmm2      ; Parallel addition of 4 floats
mulpd   xmm1, xmm3      ; Parallel multiplication of 2 doubles
```

XMM registers form the basis of **vectorized computation**, allowing multiple data elements to be processed per instruction.

4.4.2 YMM Registers (256-bit SIMD)

The **YMM registers** extend the XMM registers to 256 bits through the **AVX and AVX2** instruction sets. Each YMM register overlays an XMM register:

```
YMM0 overlays XMM0
YMM1 overlays XMM1
...
YMM15 overlays XMM15
```

Thus, the upper 128 bits of a YMM register are simply an extension of the corresponding XMM register.

Data capacity per YMM register:

Data Type	Elements (per YMM register)
8-bit integers	32 elements
16-bit integers	16 elements
32-bit integers or floats	8 elements
64-bit integers or doubles	4 elements

Example operations:

```
vmovaps ymm0, ymm1 ; Move aligned 256-bit vector
vaddps ymm2, ymm2, ymm3 ; Add 8 single-precision floats in parallel
vpmullw ymm4, ymm4, ymm5 ; Multiply 16 packed 16-bit integers
```

The v-prefixed AVX instructions introduce **three-operand form**, improving register reuse efficiency by retaining source operands.

4.4.3 ABI and Function Calling Considerations (Linux / System V AMD64)

Under the System V AMD64 ABI (Linux, BSD, UNIX-like systems):

- SIMD registers **XMM0–XMM7** are used for passing **floating-point and vector function arguments**.
- **XMM0** is used for returning floating-point and vector results.
- All SIMD registers are **caller-saved**, meaning:
 - A function may freely modify them.
 - Callers must save/restore if values must persist across calls.

For example, caller saving:

```
sub    rsp, 32
movaps [rsp], xmm0
; call function
movaps xmm0, [rsp]
add    rsp, 32
```

4.4.4 Alignment and Performance Behavior

To achieve maximum SIMD throughput:

- Data intended for XMM/YMM operations should be **16-byte aligned** (for XMM) and **32-byte aligned** (for YMM).
- Misaligned loads are tolerated on modern CPUs but may incur **additional latency** and reduce effective memory bandwidth.

Large vector loops benefit from:

- **Sequential (unit-stride) memory access**

- Avoidance of mixing XMM and YMM instructions without `vzeroupper` to prevent **AVX-SSE transition penalties**

Example of preventing transition stalls:

```
vzeroupper      ; Clear upper halves before executing SSE code
```

4.4.5 Significance for Low-Level Programming

SIMD registers are not merely wider storage units; they define a **data parallel execution model** that directly influences:

- Algorithm design (favoring block and batch computation)
- Data structure layout (favoring contiguous arrays over linked structures)
- Performance scalability across cores and NUMA regions

Mastery of XMM/YMM registers allows low-level programmers to shift workloads from being **memory latency-bound** to **compute-throughput optimized**.

Chapter 5

Instruction Flow

5.1 Fetch

The **fetch stage** is the entry point of the instruction execution pipeline. Its role is to retrieve the next instruction bytes from memory based on the current value of the **instruction pointer (RIP)**. The fetch unit operates at the boundary between the memory hierarchy and the CPU's front-end decoding hardware. Its efficiency directly influences overall throughput, as all later pipeline stages depend on a continuous supply of instruction data.

5.1.1 Instruction Pointer and Sequential Flow

At any given moment, **RIP** identifies the address of the instruction being executed. After an instruction completes, RIP typically increments to point to the next instruction:

```
; Conceptual flow (not explicit assembly)
RIP ← RIP + instruction_length
```

However, control flow instructions alter RIP explicitly:

```
call    target_label    ; Save return address to stack, jump to target
jmp     target_label    ; Unconditional branch
jz      target_label    ; Conditional branch based on zero flag
```

Fetching must therefore anticipate both **sequential** and **non-sequential** instruction flow.

5.1.2 Instruction Cache and Front-End Bandwidth

The fetch unit reads instruction bytes through the **L1 Instruction Cache (L1I)**. To achieve sustained throughput, the pipeline must fetch enough bytes each cycle to keep the decode stage supplied.

Typical constraints:

- L1I latency: ~3–4 cycles
- Fetch width: Several instruction bytes per cycle, architecture-dependent

Instruction alignment matters: sequences that cross cache-line boundaries increase fetch overhead.

5.1.3 Instruction Flow Optimization via μ op Cache

Modern x86-64 CPUs include a **micro-operation (μ op) cache** that stores previously decoded instructions in execution-ready form. If an instruction sequence already resides in the μ op cache, the fetch stage **bypasses the L1I and decoding entirely**, increasing throughput and reducing power consumption.

This optimization is most effective for:

- Tight loops
- Frequently executed code paths

- Shared library stubs and runtime kernels

Therefore, code layout and loop positioning have direct performance consequences.

5.1.4 Fetch and Branch Prediction

Because the fetch unit must determine *which* instruction to retrieve next, it cooperates closely with the **branch predictor**. When a branch is encountered and the outcome is not yet known, the fetch unit proceeds based on predicted direction. If the prediction is correct, execution proceeds without interruption. If incorrect:

- The **reorder buffer (ROB)** rolls back speculative work.
- The fetch unit restarts from the correct path.
- Pipeline bubbles (stalls) occur, reducing throughput.

Thus, **fetch efficiency depends on branch predictability**.

5.1.5 Effects of Instruction Density and Encoding

x86-64 instructions vary in length (1 to 15 bytes), which impacts fetch alignment:

- Dense encoding improves instruction cache utilization.
- Long or prefix-heavy instructions can reduce the effective fetch rate.
- Mixed code containing frequent rip-relative loads and vector instructions may increase L1I pressure.

Optimized code typically favors shorter encodings where equivalent functionality is preserved:

```
xor eax, eax      ; Efficient idiom for zeroing
```

Instead of:

```
mov eax, 0        ; Larger encoding with identical effect
```

5.1.6 Significance for Low-Level Programming

At the system level, the fetch stage determines whether the CPU consistently supplies instructions fast enough to keep the execution pipeline busy. Performance-sensitive software must therefore consider:

- Code layout to avoid crossing cache line boundaries.
- Loop structure to maximize μop cache residency.
- Branch predictability to minimize misfetch penalties.
- Choosing instruction encodings that reduce front-end bandwidth demand.

The fetch stage is not simply the beginning of execution—it sets the **maximum sustainable throughput ceiling** for the entire core.

5.2 Decode

The **decode stage** translates variable-length x86-64 machine instructions into an internal representation suitable for execution. Because x86-64 instructions have flexible encodings—with prefixes, ModR/M fields, displacement values, and immediates—the decode stage is structurally more complex than the fetch stage. Its main output is a sequence of **micro-operations (μops)**, each representing a primitive hardware action.

The decode stage forms the critical link between the architectural instruction set visible to programmers and the microarchitectural machinery responsible for actual execution.

5.2.1 Variable-Length Instruction Parsing

Unlike RISC architectures with fixed-width instructions, x86-64 instructions range from 1 to 15 bytes. The decoder must:

1. Identify and interpret prefixes (REX, segment overrides, SIMD width flags).
2. Decode the primary opcode byte.
3. Parse addressing mode bytes (ModR/M, optional SIB).
4. Extract immediate constants and displacement offsets.

This parsing determines:

- Operand types (register, memory, immediate)
- Instruction length (affects RIP advancement)
- Computation type (integer, vector, control flow)

Example (Intel syntax):

```
mov rax, [rbx + rcx*4 + 16]    ; Requires decode of ModR/M, SIB,  
↪ displacement
```

The decode engine extracts **operation semantics** from a compact and expressive binary format.

5.2.2 Translation Into Micro-Operations (μops)

Most x86-64 instructions are not executed directly. Instead, each instruction becomes one or more μops. A μop corresponds to a single scheduling and execution-ready hardware action.

Examples:

Instruction	Possible μop Breakdown
add rax, rbx	Single ALU μop
mov rax, [rdx]	Address calculation μop + load μop
mul rax	Multi-step pipeline sequence via dedicated multiply unit

Complex instructions (e.g., string operations) may produce multiple μops, while certain fused instruction pairs may form a single μop under microarchitectural optimization.

5.2.3 Decode Width and Front-End Throughput

Modern x86-64 CPUs decode **multiple instructions per cycle**, but the exact decode bandwidth varies by microarchitecture. If decode cannot supply μops fast enough, the **front-end becomes the bottleneck**, not the execution units.

Performance-sensitive code therefore benefits from:

- Shorter instruction encodings
- Avoidance of unnecessary prefixes
- Compact loop bodies that fit in μop and instruction caches

5.2.4 The Microcode ROM

Some instructions—especially those that perform complex or rarely used operations—cannot be expressed efficiently as a small set of μ ops. For these, the decode stage triggers a **microcode assist**, dispatching a **microcode sequence** from an internal ROM-like structure. Examples include:

- Certain system instructions (CPUID, LGDT, WRMSR)
- Legacy instructions with complex side effects
- Floating-point exception paths

Microcode execution is **slower**, and thus avoided in performance-critical paths.

5.2.5 μ op Cache Short-Circuiting Decode

If a sequence of instructions is cached in the CPU's **μ op cache**, the decode stage can be bypassed. In such cases, the front-end supplies μ ops directly from the μ op cache to the scheduler.

Conditions improving μ op cache usage:

- Loop bodies placed contiguously
- Predictable branch flow
- Avoiding interleaving code paths inside hot loops

Effective μ op caching **raises front-end throughput** and reduces power cost.

5.2.6 Practical Implications for Low-Level Development

To avoid decode bottlenecks and ensure efficient front-end feed:

- Prefer instructions that map to few μ ops.
- Use addressing forms that avoid redundant complexity.
- Maintain compact, predictable loop structures to maximize μ op cache residency.
- Favor three-operand v -prefixed AVX instructions over SSE forms to reduce register traffic.

The decode stage does not simply interpret instruction encodings; it defines **how efficiently the processor can begin execution**. If decode throughput is insufficient, no amount of execution-unit parallelism can compensate.

5.3 Execute

The **execute stage** is where micro-operations (μ ops) are issued to the appropriate **execution units** and produce results. While the architectural view of execution suggests that instructions run sequentially, the microarchitectural reality is that modern x86-64 CPUs perform **out-of-order, parallel, and speculative execution**, driven by operand readiness and hardware scheduling policies. The execute stage does not directly follow program order; instead, it proceeds according to data dependencies and resource availability.

5.3.1 Execution Units and μ op Dispatch

After decode and register renaming, μ ops enter the **scheduler** (or reservation stations). When all input operands for a given μ op are available, the scheduler selects an appropriate **execution port**, which in turn routes the μ op to a matching execution unit.

Execution units are specialized:

Unit Type	Operations (Intel Syntax Examples)
Integer ALU	<code>add rax, rbx, xor rdx, rcx, shl rax, 1</code>
Address Generation Unit (AGU)	<code>mov rax, [rbx + rcx*4 + 8]</code>
Floating-Point Unit (FPU)	<code>addsd xmm0, xmm1, mulsd xmm2, xmm3</code>
SIMD / Vector Unit	<code>vaddps ymm0, ymm1, ymm2, vpmullw ymm3, ymm4, ymm5</code>
Load Unit	Reads from memory into registers
Store Unit	Writes register data into memory

Multiple μ ops may issue per cycle when independent and when execution units are available.

5.3.2 Latency vs. Throughput

Execution performance depends on two key metrics:

Metric	Meaning
Latency	Cycles before a result becomes available for dependent μ ops
Throughput	How many μ ops of that kind can complete per cycle

Example (typical values, microarchitecture-dependent):

Operation	Latency	Throughput
Integer add	~ 1 cycle	1–2 per cycle
Multiply (int)	~ 3 cycles	1 per cycle
Load from L1	~ 4 cycles	Multiple in parallel if not conflicting
Load from DRAM	$\sim 200+$ cycles	Stalls pipeline if dependency exists

To maintain pipeline utilization, **avoid long dependency chains**:

```
; Poor: serialized dependence
add rax, rbx
add rax, rcx
add rax, rdx

; Better: parallelizable accumulation
mov r8, rbx
add r8, rcx
add rax, r8
add rax, rdx
```

5.3.3 Operand Readiness and Dependency Chains

The scheduler issues a μ op only if:

- Its source registers are ready.
- Memory dependencies are resolved.
- A suitable execution port is free.

True data dependencies (RAW: Read-After-Write) limit parallelism. False dependencies (WAR/WAW) are eliminated via register renaming, allowing multiple live versions of registers.

This is a key enabler of **out-of-order execution**.

5.3.4 Speculative Execution

When control flow is uncertain, the CPU **speculates** by executing μ ops along a predicted path. If the prediction is correct, results commit normally. If incorrect:

- The **Reorder Buffer (ROB)** discards speculative results.
- The pipeline refills along the correct path.

Speculation increases throughput but makes execution **non-deterministic** in timing, especially for branch-heavy workloads.

5.3.5 Memory Access in Execution

Memory interactions are significantly slower than register operations, so load/store performance strongly influences execution behavior:

- **Load-use latency** stalls occur when a μ op depends on a memory load that has not yet completed.
- High-performance loops minimize memory pressure and maximize register reuse.

- Strided or irregular memory patterns reduce effective parallelism.

Memory hints (e.g., software prefetching) can reduce latency where predictable access exists.

5.3.6 Implications for Low-Level Development

Optimizing execution requires:

- Reducing dependency depth to expose instruction-level parallelism.
- Structuring loops to maintain steady feed into execution units.
- Preferring operations with high throughput and predictable latency.
- Ensuring memory access patterns align with cache behavior.

In modern x86-64 CPUs, raw compute resources are rarely the limiting factor; **data availability and dependency scheduling define practical performance.**

5.4 Writeback & Commit

After execution, instructions must update the **architectural state** of the processor in a way that preserves program correctness. While execution may occur *out of order* and speculatively, the processor must ensure that the final, visible effects of instructions appear as if they executed strictly *in program order*. The **writeback** and **commit** (or **retirement**) stages achieve this synchronization.

5.4.1 Writeback to Physical Registers

During execution, each μop writes its result to a **physical register**, not directly to an architectural register. This preserves multiple in-flight versions of the same logical register and prevents overwriting values needed by earlier instructions.

Example (Intel syntax conceptual flow):

```
add rax, rbx      ; RAX_new stored in a new physical register entry
```

At this point:

- The **result exists internally**, but
- It is **not yet visible** to the architectural register file.

This is critical for supporting speculation and rollback.

5.4.2 The Reorder Buffer (ROB) and In-Order Retirement

The **Reorder Buffer (ROB)** tracks every in-flight μop in the original program order. When a μop completes execution, the ROB marks it as **ready to retire**, but it does not commit immediately. Retirement occurs only when:

- All previous instructions have also completed.
- No exceptions, interrupts, or branch mispredictions are pending.

This ensures the program's architectural state remains **consistent and precise**.

If a misprediction or exception occurs, the ROB discards all uncommitted μ ops:

```
Speculative instructions → Invalidated  
Committed instructions → Preserved
```

This rollback mechanism is central to modern speculative execution.

5.4.3 Register Rename Resolution at Commit

At commit, the mapping of physical registers to architectural registers is **updated**:

- The newly computed physical register is designated as the current architectural value.
- The prior physical register version is returned to the **free list**, making it available for future renaming.

Thus:

- Writeback updates **physical state**, while
- Commit updates **architectural state**.

This separation allows the CPU to run ahead of the visible program state without losing correctness.

5.4.4 Stores and Memory Visibility

Memory updates do not commit at execution time. They are held in the **store buffer** until the instruction retires. Only then can the write propagate to the cache hierarchy, ensuring correct memory ordering.

Example flow:

```
mov [rdi], rax      ; Store result only visible after commit
```

This guarantees:

- Speculative stores do **not** overwrite real program data.
- Failed speculation incurs **no observable memory change**.

5.4.5 Exception and Interrupt Precision

The commit stage preserves **precise exceptions**: faults appear as if they occurred exactly at the faulting instruction in program order.

If a fault occurs:

- The ROB flushes younger instructions.
- State is restored to the last known committed point.
- Execution resumes at the fault handler.

This model allows debugging, signal handling, and system call transitions to remain deterministic.

5.4.6 Performance Considerations

Writeback and commit can form a pipeline bottleneck if:

- Too many μ ops complete at once relative to the retirement rate.
- Long dependency chains serialize execution and reduce ROB turnover.
- Frequent mispredictions trigger pipeline flushes.

Optimized code reduces pressure on commit by:

- Limiting speculative control flow hazards.
- Avoiding excessive register pressure that triggers spills.
- Maintaining memory-access patterns that avoid coherence invalidations.

5.4.7 Summary

The writeback and commit stages ensure that:

Stage	Purpose
Writeback	Records results in <i>physical</i> registers
Commit (Retirement)	Makes results <i>architecturally visible</i> in program order

This separation enables high-performance, speculative, out-of-order cores to preserve the **illusion of sequential execution**, while exploiting deep parallelism internally.

Chapter 6

Virtual Memory & Paging

6.1 Address Translation

Modern x86-64 systems do not use physical memory addresses directly when executing user-space programs. Instead, each process operates within a **virtual address space**, where the addresses referenced in instructions—load, store, call, jump—are **virtual addresses (VAs)**. These addresses must be translated into **physical addresses (PAs)** before data can be fetched from DRAM. The component responsible for this translation is the **Memory Management Unit (MMU)**, guided by page table structures maintained by the operating system.

Address translation provides memory isolation, controlled sharing, efficient allocation, and the ability to run programs independently of physical memory layout.

6.1.1 Virtual vs. Physical Addresses

A program may reference:

```
mov rax, [rbx]           ; Load from virtual address stored in RBX
```

The value in RBX is *not* a physical address. It is an index into a process-specific virtual address space. The MMU translates this virtual address to a physical location using multiple levels of lookup.

Key properties:

Address Type	Meaning	Managed By
Virtual Address (VA)	Seen by software	OS + CPU MMU
Physical Address (PA)	Seen by DRAM, caches, memory controllers	Hardware only

The CPU **never exposes** physical addresses to user space directly.

6.1.2 Page-Based Translation

Virtual memory is divided into fixed-size blocks called **pages**. The standard page size on x86-64 Linux is **4 KB**, though 2 MB (huge pages) and 1 GB (gigantic pages) are also supported for performance.

Translation occurs via **page tables**, a hierarchical structure mapping:

```
Virtual Address (48 bits typical)
  VPN      Offset
    translated via MMU  no change
```

The lower bits of the address provide the offset within the page and are not translated.

```
→ Page Fault
→ Kernel exception handler invoked
→ Possible segmentation fault if illegal
```

6.1.3 Multi-Level Page Tables (x86-64)

x86-64 uses a **4-level paging hierarchy** for standard virtual memory:

Level	Name	Purpose
Level 4	PML4	Selects the highest region of address space
Level 3	PDPT	Points to a page directory
Level 2	PD	Points to a page table
Level 1	PT	Maps a virtual page to a physical frame

The CPU performs this walk automatically when no cached translation exists.

6.1.4 The TLB and Translation Caching

A **Translation Lookaside Buffer (TLB)** caches recent VA→PA translations to avoid repeatedly walking page tables.

If a translation is found in the TLB (a **TLB hit**), translation is effectively free.

If not (**TLB miss**):

- The CPU performs a page-table walk.
- This walk requires multiple memory accesses.
- Latency increases sharply (dozens to hundreds of cycles, depending on cache state).

Thus, TLB behavior directly impacts performance.

6.1.5 Privilege Enforcement During Translation

During translation, the MMU also checks:

- **Access permissions** (read/write/execute)
- **User vs. kernel mode validity**
- **Shared vs. private mapping semantics**

If an access violates permissions:

- Page Fault
- Kernel exception handler invoked
- Possible segmentation fault if illegal

This is the foundation of **process isolation** on Linux.

6.1.6 Virtual Memory Allows Memory Overcommitment

Because virtual addresses do not map 1:1 to physical memory:

- A process may reference more memory than physically installed.
- Pages may be paged out to disk if inactive.
- The OS may lazily allocate memory pages on first access (demand paging).

This enables running large applications efficiently—at the cost of potential page-fault latency.

6.1.7 Practical Implications for Low-Level Programming

Performance depends on **locality and translation efficiency**:

Strategy	Benefit
Maintain spatial locality	Reduces page transitions and TLB pressure
Use large pages when appropriate	Reduces TLB entries required
Avoid random pointer chasing	Prevents high TLB miss rates
Keep hot data within the same page	Minimizes page walk stalls

Address translation is not a passive mechanism—it is a **performance-critical component** of execution on Linux x86-64 systems.

6.2 Page Tables (4-Level Paging)

The x86-64 architecture implements virtual memory through a hierarchical page table system. Each process has its own page table structure that maps **virtual addresses (VAs)** to **physical addresses (PAs)**. This structure enforces memory protection, isolation between processes, shared-memory regions, and controlled access to kernel memory. On Linux, page tables are constructed and maintained by the kernel, while the **Memory Management Unit (MMU)** performs translation automatically during execution.

6.2.1 Hierarchical Layout

x86-64 uses **4-level paging** for 48-bit virtual address spaces (common in modern Linux systems). The virtual address is divided into fields that index successive table levels:

```
[ PML4 | PDPT | PD | PT | Offset ]
  9 bits  9 bits  9 bits  9 bits  12 bits
```

This yields:

Component	Bits	Meaning
PML4 Index	9	Selects root table entry
PDPT Index	9	Selects page directory pointer
PD Index	9	Selects page directory
PT Index	9	Selects page table
Offset	12	Byte offset inside a 4 KB page

Each level contains **512 entries**, and each entry is **8 bytes**, forming **4 KB per page table page**.

6.2.2 Page Table Levels and Roles

Level	Name	Points To	Final Mapping Resolution
L4	PML4	A PDPT	Defines top-level region of VA space
L3	PDPT	A Page Directory	May also contain 1 GB page entries
L2	Page Directory (PD)	A Page Table	May also contain 2 MB page entries
L1	Page Table (PT)	Physical memory frame numbers	Always resolves to 4 KB pages

Linux may select **large pages** by using entries at L2 or L3 that directly map large physical regions, bypassing lower levels.

6.2.3 Page Table Entry Format

Each Page Table Entry (PTE) encodes the physical frame address plus **access control flags**:

Common flags:

Bit	Name	Effect
P	Present	Page is mapped and accessible
RW	Read/Write	Allows writes if set
US	User/Supervisor	Controls user vs. kernel access
NX	No-Execute	Prevents code execution from page
A	Accessed	Hardware-set when page is touched
D	Dirty	Set on writes to pages (used in paging decisions)

These bits enforce **memory correctness**, **security boundaries**, and **executable memory restrictions**.

6.2.4 Example Translation Walk (Conceptual)

For instruction:

```
mov rax, [rbx]      ; Virtual address in RBX
```

The MMU performs:

1. **Index PML4** using high bits of VA → fetches PDPT pointer
2. **Index PDPT** → fetches PD pointer
3. **Index PD** → fetches PT pointer
4. **Index PT** → yields physical frame number
5. Combine with offset → Physical Address

If any entry has $P = 0$, a **page fault** occurs and Linux handles it (either mapping a page or signaling segmentation fault).

6.2.5 Cost of a Page Walk

A page walk may require **4–5 memory accesses**, depending on cache residency. If the walk is not cached, latency can reach **hundreds of cycles**.

However, recent translations are cached in the **Translation Lookaside Buffer (TLB)**. Thus:

- **TLB hit** → translation is effectively free.
- **TLB miss** → full page walk → significant stall.

Performance depends heavily on **TLB locality**, not just cache locality.

6.2.6 Practical Implications for Low-Level Programming

To optimize memory translation overhead:

Guideline	Benefit
Keep working sets within a small page footprint	Reduces TLB pressure
Prefer contiguous data structures	Increases spatial locality
Use large pages (2 MB or 1 GB) for large arrays	Reduces number of PTEs and TLB entries
Avoid pointer chasing	Prevents unpredictable TLB miss patterns

Low-level performance tuning is as dependent on **address translation behavior** as on cache behavior.

6.3 TLB and Misses

The **Translation Lookaside Buffer (TLB)** is a hardware cache that stores recent virtual-to-physical address translations. Because page tables are multi-level and require multiple memory references per lookup, directly translating every memory access through page tables would be prohibitively slow. The TLB avoids these repeated walks by caching translation results. When a virtual address lookup hits in the TLB, the physical address is obtained in essentially **zero additional latency**, allowing the pipeline to continue uninterrupted.

However, when a translation is not present in the TLB (**TLB miss**), the processor must perform a full **page table walk**, significantly increasing memory access cost.

6.3.1 TLB Structure and Levels

Modern x86-64 CPUs have several TLBs, typically split by cache level and page size:

TLB Level	Coverage	Notes
L1 TLB	Small, very fast	Serves most core accesses
L2 / Unified TLB	Larger, slower	Shared among more operations
Page-Size-Specific TLBs	Separate entries for 4 KB, 2 MB, and 1 GB pages	Reduces conflict pressure

L1 TLB hits provide translation with minimal latency. Accesses missing in L1 fall back to L2 TLB. If both miss, the hardware performs a page walk.

6.3.2 TLB Miss Cost

When a TLB miss occurs, the CPU performs a page walk through the hierarchy:

PML4 → PDPT → PD → PT → Physical Frame

If these entries are not present in cache:

- Each lookup may require one or more memory accesses.
- Total cost may reach **dozens to hundreds of cycles**, depending on memory locality.

Therefore, frequent TLB misses can overshadow arithmetic and even L1/L2 cache performance.

6.3.3 Impact of Working Set Size

The TLB can only store a limited number of translations. If the working set spans more pages than the TLB can hold, translations are constantly evicted and reloaded, causing **TLB thrashing**.

This occurs often in workloads that:

- Operate on large arrays with sparse or strided access
- Walk pointer-based structures (trees, linked lists, graphs)
- Use process-local memory across a large virtual footprint

Underlying principle: **Performance degrades when memory access patterns exceed translation locality.**

6.3.4 Page Size and TLB Efficiency

Using **large pages** reduces the number of TLB entries required to map the same address range:

Page Size	Data per TLB Entry	Effect on TLB Pressure
4 KB	Small coverage	Highest pressure
2 MB (Huge Page)	512× more coverage	Significant reduction in misses
1 GB Page	Very large coverage	Useful only for specialized memory layouts

Linux supports large pages via:

- `mmap` with `MAP_HUGETLB`
- Transparent Huge Pages (THP), when enabled

Large pages reduce TLB miss frequency, but misuse can increase memory fragmentation.

6.3.5 Access Patterns and TLB Behavior

Memory access patterns determine TLB efficiency:

Pattern	Result
Contiguous traversal	Excellent TLB locality
Strided access with power-of-two spacing	Possible TLB/set conflicts
Pointer chasing (linked structures)	Poor predictability → frequent misses
NUMA remote access	Translation remains local, but data fetch costs increase

TLB misses are especially destructive in pointer-heavy algorithms, where each load depends on the previous pointer.

6.3.6 Measuring and Diagnosing TLB Misses on Linux

Tools such as `perf` allow monitoring of TLB hit/miss behavior:

```
perf stat -e dTLB-load-misses,dTLB-loads ./program
```

High miss ratios indicate that the working set's **page footprint**, not computation, is the limiting factor.

6.3.7 Practical Strategies for Low-Level Optimization

Strategy	Effect
Keep hot data on as few pages as possible	Decreases TLB churn
Use arrays or SoA (Structure of Arrays) over linked lists	Increases sequential access locality
Use huge pages for large contiguous buffers	Reduces number of TLB entries required
Align data to page boundaries to avoid unintended spread	Ensures predictable mapping
Prefetch data when access patterns are known	Helps pipeline hide translation latency

The key principle: **performance is maximized when virtual address usage is stable, compact, and predictable.**

6.4 Kernel vs User Address Space

In x86-64 Linux, the virtual address space is divided into two distinct regions: **user space** and **kernel space**. Although every process shares the same *kernel* virtual address mappings, each process has its own private *user* virtual memory. This separation is enforced through page permission mechanisms and hardware privilege levels, ensuring that user applications cannot directly modify or access kernel memory.

This isolation prevents accidental corruption of system state and forms the foundation of security and process containment.

6.4.1 Virtual Address Space Layout (Typical Linux x86-64)

A standard 48-bit canonical virtual address space is structured as follows:

```
+-----+ High virtual addresses
|      Kernel Space      | (shared among all processes)
|      (typically at     |
|  ffffffff8000000000 → ffff...) |
+-----+
|      User Space        | (unique to each process)
|  0000000000400000 → 00007f...) +
|      Heap, Stack, Code |
+-----+ Low virtual addresses
```

Key characteristics:

Region	Visibility	Access Mode	Purpose
User Space	Private per process	Ring 3 (unprivileged)	Code, heap, stack, mapped libs

Region	Visibility	Access Mode	Purpose
Kernel Space	Shared globally	Ring 0 (privileged)	Kernel code, drivers, device mappings, page tables

Although both are mapped simultaneously in the page tables, **user code cannot access kernel space** because of page protection flags and CPU privilege level enforcement.

6.4.2 Privilege Levels and Access Control

The CPU enforces separation using **Rings**:

Ring	Privilege	Executed Code
Ring 0	Highest	Kernel, drivers, interrupt handlers
Ring 3	Lowest	User applications, runtime libraries

When executing in user mode:

- Memory pages with `US=0` (kernel-only) cause a **page fault** if accessed.
- Instructions that modify control registers or I/O ports trigger **general protection faults**.

Thus, **correctness is enforced by hardware, not just software convention**.

6.4.3 System Calls: Controlled Transition to Kernel Mode

User applications interact with the kernel only through controlled entry points such as system calls:

```

mov rax, 1           ; sys_write
mov rdi, 1           ; stdout
mov rsi, msg
mov rdx, msg_len
syscall              ; transition to kernel space

```

During this transition:

- The CPU switches to **Ring 0**.
- The kernel stack (not the user stack) becomes active.
- Execution continues at a kernel-defined entry point.

When the call completes, control returns to user space by restoring `RIP`, `RSP`, and flags.

6.4.4 Shared Kernel Mapping and Security Constraints

For performance reasons, kernel memory remains mapped even while running user processes—this enables:

- Fast syscall entry/exit
- Efficient interrupt handling
- Context switching without reloading page tables

However, to mitigate side-channel attacks (post-2018 designs), **Kernel Page-Table Isolation (KPTI)** may separate user and kernel page tables selectively, preventing speculative access across privilege boundaries.

6.4.5 User Space Memory Management vs Kernel Memory Management

Property	User Space	Kernel Space
Allocation Mechanism	<code>malloc()</code> , <code>mmap()</code> , <code>brk()</code>	Slab allocators, buddy system

Property	User Space	Kernel Space
Fault Handling	Page faults may grow heap or map pages	Faults often signify kernel bugs
Memory Lifetime	Controlled by process	Long-lived and shared across system
Execution Ability	Depends on NX flag	Kernel regions generally executable unless hardened

The user space is **dynamic and process-scoped**, while kernel memory is **global and persistent**.

6.4.6 Practical Implications for Low-Level Programming

For performance and correctness:

Guideline	Reason
Never assume kernel addresses are valid in user mode	Hardware forbids access
Ensure user–kernel transitions (system calls) are minimized in tight loops	Syscall and context-switch overhead
Prefer user-space synchronization mechanisms when possible	Kernel locks are slower and globally visible
Use <code>mmap()</code> to control layout, alignment, and permissions explicitly	Heap allocators abstract page behavior

In essence, the CPU executes all code in a single virtual address space, but **privilege enforcement, page permissions, and syscall transitions keep user processes isolated and safe**.

Chapter 7

Interrupts and Exceptions

7.1 IDT Structure

The **Interrupt Descriptor Table (IDT)** is a core structure in x86-64 systems that defines how the processor responds to **interrupts** (external events) and **exceptions** (internal faults or conditions). Each entry in the IDT maps an **interrupt vector number** to a specific handler routine located in privileged memory. When an interrupt or exception occurs, the CPU consults the IDT to determine what code to execute and with what privilege level.

The IDT is essential for **context switching**, **kernel fault handling**, **system call entry**, and **hardware device signaling**.

7.1.1 IDT Location and Register

The processor stores the base address and size of the IDT in the **IDTR** register. The kernel sets this register during initialization:

```
lidt [idtr]      ; Load IDT pointer (base + limit)
sidt [buffer]    ; Store current IDT pointer (for debugging or diagnostics)
```

The IDTR contains:

Field	Meaning
Base	Virtual address of the IDT table in memory
Limit	Size in bytes minus one

The IDT itself must reside in **kernel-accessible memory**.

7.1.2 IDT Entry Format (Gate Descriptors)

Each IDT entry is a **16-byte descriptor** containing:

Field	Purpose
Offset (Low/Mid/High)	64-bit pointer to handler function
Selector	Code segment selector (kernel code segment in 64-bit mode)
IST Index	Optional alternate stack selector (for critical exceptions)
Type and Attributes	Defines gate type, present bit, privilege level

The **gate type** determines how the CPU transitions to the handler:

Gate Type	Meaning
Interrupt Gate	Disables further interrupts on entry
Trap Gate	Leaves interrupts enabled
Task Gate	Used for legacy TSS-based switching (rare in x86-64 Linux)

Linux uses **interrupt gates** for hardware interrupts and **trap gates** for exceptions where continued interrupt reception is safe.

7.1.3 Vector Number Assignment

The IDT is indexed by a **vector number** (0–255):

Range	Usage
0–31	Exceptions (CPU-generated faults)
32–255	Hardware interrupts, system-specific handlers

Examples:

Vector	Condition
0x00	Divide-by-zero error
0x0D	General protection fault
0x0E	Page fault
0x20 and up	IRQ-mapped hardware interrupt lines

Linux remaps legacy PIC/IOAPIC hardware interrupts starting at vector 32 to avoid collisions with CPU exceptions.

7.1.4 Privilege Level Control

IDT entries contain a **Descriptor Privilege Level (DPL)**, which dictates from which CPU privilege levels the interrupt may be triggered:

DPL	Who Can Invoke
0	Kernel only
3	User mode permitted to trigger (e.g., system call entry)

User-mode software **cannot** invoke most interrupt vectors directly due to DPL restrictions. This enforces kernel protection boundaries.

7.1.5 Interrupt Stack Table (IST) Integration

x86-64 provides the **IST** field to allow certain critical exceptions to switch to a dedicated, known-good stack, regardless of the current stack state.

Handlers that commonly use IST entries:

Exception	Reason
Double Fault	Primary/secondary stack corruption handling
Non-Maskable Interrupt	CPU state salvage paths

This mechanism prevents kernel crashes when the current kernel stack is invalid or overflowed.

7.1.6 Practical Low-Level Considerations

When writing low-level routines:

- The handler must preserve registers per System V AMD64 ABI before calling C-level routines.
- Handlers must acknowledge hardware interrupts (e.g., APIC or PIC end-of-interrupt signaling).
- Invalid IDT entries or misconfigured privilege bits typically result in a **triple fault**, causing immediate system reset.

Correct IDT setup is foundational for **process scheduling**, **preemption**, **signal delivery**, and **debugging traps** in Linux.

7.2 Hardware vs Software Interrupts

Interrupts allow the processor to **pause** the currently executing instruction sequence and transfer control to a handler routine. Although all interrupts follow the same **IDT-based dispatch mechanism**, they originate from two fundamentally different sources: **hardware** and **software**. Understanding this distinction is essential for reasoning about system responsiveness, privilege transitions, and kernel control flow on Linux x86-64 systems.

7.2.1 Hardware Interrupts

Hardware interrupts are generated by external devices that require CPU attention. Common sources include:

- Timers (e.g., HPET, APIC timer)
- Network interfaces
- Storage controllers (NVMe, SATA, SCSI)
- PCIe expansion hardware
- Power and thermal management units

Hardware interrupts are delivered via the system's **interrupt controller**, historically the PIC, now replaced in modern systems by the **Advanced Programmable Interrupt Controller (APIC)** architecture (local APIC + I/O APIC).

When a hardware interrupt occurs:

1. The CPU stops executing the current instruction stream.
2. The interrupt vector is looked up in the **IDT**.

3. Control transfers to a kernel-mode **interrupt service routine (ISR)**.
4. The kernel acknowledges the interrupt and performs the required processing.
5. Execution resumes where it left off.

Hardware interrupts **always** transition execution to **Ring 0** (kernel mode).

Example flow (conceptual, not literal code):

```
; interrupt arrives
push interrupt_frame
jmp [IDT + vector*16] ; CPU performs this automatically
```

7.2.2 Masking and Prioritization

The CPU can **mask** (block) certain external interrupts by clearing the IF (Interrupt Flag) bit in RFLAGS:

```
cli ; Clear interrupt flag (disable maskable interrupts)
sti ; Set interrupt flag (enable interrupts)
```

Non-maskable interrupts (NMIs) cannot be blocked and are reserved for critical hardware events.

7.2.3 Software Interrupts

Software interrupts are generated explicitly by instructions. They are used to request kernel services or trigger diagnostic/exception paths.

Relevant instructions:

```
int 0x80 ; Legacy Linux system call gateway
syscall ; Modern, fast system call entry
int3 ; Breakpoint, used by debuggers
```

Unlike hardware interrupts, software interrupts:

- Are **synchronous** (execution continues from the triggering instruction's semantic boundary)
- Occur under program control
- Typically initiate **controlled privilege transitions** (user $\rightarrow$ kernel)

7.2.4 Behavioral Differences

Property	Hardware Interrupt	Software Interrupt
Origin	External device or hardware event	CPU instruction execution
Timing	Asynchronous	Synchronous
Privilege Transition	Always enters kernel mode	Often enters kernel mode; depends on interrupt vector DPL
Typical Use	Device signaling, scheduling ticks	System calls, debugging, exception testing

Hardware interrupts drive **asynchronous concurrency**, while software interrupts initiate **explicit service requests** or architectural exceptions.

7.2.5 Kernel Handling Differences

Hardware interrupts typically require:

- Acknowledgment signaling (e.g., APIC End-of-Interrupt register write)
- Deferred processing via **bottom halves** or **tasklets** to avoid long handler times

Software interrupts generally:

- Follow well-defined, ABI-stable entry paths (e.g., system call stubs)
- Execute in kernel context but may return immediately to user mode

Linux ensures that interrupt handlers run **on the kernel stack**, isolated from user stack corruption.

7.2.6 Practical Low-Level Considerations

For performance-critical system work:

Guideline	Reason
Avoid long computations inside hardware interrupt handlers	Prevent system-wide latency stalls
Minimize system call frequency in hot loops	Each software interrupt induces a privilege transition cost
Use <code>syscall</code> instead of <code>int 0x80</code>	Faster entry/exit path on x86-64
Ensure correct interrupt masking when modifying shared state	Prevent inconsistent kernel data structures

Interrupt interactions define the **timing, responsiveness, and concurrency structure** of the operating system.

7.3 Exception Types

Exceptions are events generated **by the CPU itself** when it detects abnormal or special execution conditions. Unlike hardware interrupts, exceptions are **synchronous**—they occur *as a direct result* of the instruction currently executing. The CPU halts instruction retirement at the faulting point, pushes an exception frame, and dispatches control through the **IDT** to a handler in the kernel.

Exceptions enforce memory safety, arithmetic correctness, control-transfer integrity, and program debugging semantics on Linux x86-64.

7.3.1 Classification of Exceptions

Exceptions are categorized into **faults**, **traps**, and **aborts**, based on when control is transferred and whether the faulting instruction can be re-executed.

Type	When Delivered	Can Instruction Restart?	Typical Use
Fault	Before the faulting instruction completes	Yes	Page faults, protection faults
Trap	After the instruction completes	Yes (returns to next instruction)	Breakpoints, single-step debugging
Abort	Unpredictable state; cannot restart	No	Severe corruption conditions (e.g., double fault)

This classification affects how the kernel resolves or propagates the exception.

7.3.2 Faults (Recoverable Errors)

Faults occur **before** the instruction commits architectural state. The OS may correct the cause and **resume execution**.

Common examples:

Vector	Name	Meaning
0x00	Divide Error	<code>div</code> or <code>idiv</code> with zero divisor
0x0D	General Protection Fault	Illegal memory or privilege operation
0x0E	Page Fault	Virtual address not mapped or access not permitted

Example of a fault-triggering instruction:

```
mov rax, [invalid_address] ; Causes page fault if unmapped
```

In Linux, page faults typically trigger:

- Demand paging (map new physical memory), or
- Segmentation fault signal (SIGSEGV) if unresolved.

7.3.3 Traps (Observation Points)

Traps occur **after** instruction execution and allow execution to resume at the next instruction. They serve primarily **debugging and monitoring** purposes.

Vector	Name	Trigger
0x03	Breakpoint Trap	<code>int3</code> instruction or debugger instrumentation
0x01	Debug Trap	Single-step mode via TF flag in RFLAGS

Example breakpoint:

```
int3 ; Causes trap → debugger gains control
```

Traps do *not* indicate error; they indicate **intentional control transfer**.

7.3.4 Aborts (Non-Recoverable Errors)

Aborts indicate **serious conditions** where execution state cannot be reliably continued. They are not restartable.

Vector	Name	Cause
0x08	Double Fault	Exception while handling another exception (stack corruption, invalid IDT)
Machine Check	Hardware-detected CPU or memory integrity error	

Aborts generally result in process termination or system halt, depending on severity and kernel policy.

7.3.5 Delivering Exceptions to User Space

If the kernel **cannot resolve** an exception, it converts the event to a **signal** and delivers it to the affected process:

Exception Cause	Linux Signal
Page Fault with illegal access	SIGSEGV
Division by zero	SIGFPE
Breakpoint (<code>int 3</code>)	SIGTRAP
Illegal instruction	SIGILL

This provides a consistent software-visible fault model.

7.3.6 Practical Low-Level Guidelines

Guidance	Rationale
Avoid implicit faults by validating pointers before use	Prevents unnecessary kernel transitions
Use <code>int3</code> for debugging instead of modifying memory	Produces precise trap semantics
Understand page faults when designing memory allocators	Demand paging affects latency behavior
Never assume exceptions are rare in system-level code	Faults are integral to memory and process management

Exceptions are not merely error conditions—they are **mechanisms by which Linux enforces memory safety, privilege separation, and debugging visibility**.

7.4 Common Debug Traps

Debug traps are exceptions intentionally generated to support **inspection, control, and instrumentation** of program execution. Unlike hardware interrupts, debug traps are synchronous and occur as part of the instruction flow. They allow debuggers, profilers, and tracing tools to observe processor state without altering program semantics beyond control transfer.

Linux relies on these traps to implement **breakpoints, single-stepping, watchpoints, and trap-based system call tracing** (e.g., `ptrace`).

7.4.1 INT3 (Breakpoint Trap)

The most widely used debug trap is triggered via the `int3` instruction:

```
int3          ; One-byte breakpoint instruction
```

Characteristics:

- Vector: **0x03**
- Type: **Trap** (executes *after* the breakpoint instruction completes logically)
- Used by: GDB, perf, dynamic instrumentation frameworks (e.g., ftrace)

Since `int3` is a **single-byte opcode**, it can replace any instruction byte boundary without shifting code layout. Debuggers temporarily insert `int3` instructions into target code, then restore original bytes upon trap handling.

7.4.2 Single-Step Trap (Trap Flag in RFLAGS)

Setting the **Trap Flag (TF)** in `RFLAGS` causes the CPU to trigger a debug trap **after every instruction executed**:

```
pushfq
or     qword [rsp], 0x100    ; Set TF (bit 8)
popfq
```

When TF is set:

- Vector: **0x01**
- Use case: Step-through debugging, instruction-level profiling, controlled dynamic binary analysis.

Single-stepping is precise but expensive due to trap overhead on every instruction.

7.4.3 Hardware Breakpoints (DR0–DR7 Debug Registers)

The CPU provides dedicated **debug registers** for setting execution, read, or write watchpoints on specific addresses:

Register	Purpose
DR0–DR3	Hold watched linear addresses
DR6	Status (indicates which breakpoint triggered)
DR7	Control (enable flags, length/type specification)

Example: break when writing to a variable in memory.

These traps are handled **without modifying program code** and are essential for debugging data corruption issues.

7.4.4 Watchpoints and Data Access Traps

A watchpoint triggers when memory at a specific address is:

- Read

- Written
- Executed

The trap behavior:

Event	Debug Trap Vector	Typical Usage
Data write	0x01	Detecting unintended memory modification
Instruction fetch	0x01	Executing code in unexpected regions

Since watchpoints rely on debug registers, the number of hardware-enforced watchpoints is limited.

7.4.5 Page Protection Traps for Debugging

Temporarily marking memory **read-only** or **non-executable** can induce controlled page faults:

```
; marking memory execute-disabled allows trap-based tracing
```

This method is used in:

- Runtime code integrity enforcement
- JIT execution validation
- Self-modifying code protections

The resulting exception is a **fault**, not a trap, but is often interpreted as a *debug event* in Linux via SIGSEGV or SIGBUS.

7.4.6 Interaction With ptrace and the Kernel

Linux uses debug traps to implement **process tracing**:

- When a debug trap occurs, the kernel stops the process.
- If the process is being traced (`ptrace`), the kernel notifies the tracer.
- The tracer may inspect registers, memory, or resume execution.

This mechanism underlies debuggers, sandboxing frameworks, and system call tracing.

7.4.7 Practical Low-Level Guidelines

Guideline	Rationale
Use <code>int3</code> for temporary breakpoints	Safe and code-alignment-preserving
Use hardware watchpoints when debugging data corruption	Avoids modifying code regions
Avoid enabling single-step mode in performance-sensitive regions	Trap overhead is extremely high
Be aware that some traps change timing behavior	Debugging may mask concurrency issues

Debug traps are not merely debugging conveniences—they are **architectural mechanisms used to control, introspect, and enforce program correctness** in complex Linux systems.

Chapter 8

Processor Manuals

8.1 Intel Manuals Layout

Intel publishes the **Architectural and Developer Manuals** as the authoritative specification of the x86-64 ISA, privilege model, paging structures, and micro-architectural behavior. These manuals define the interface the CPU presents to software, independent of specific processor generations. Linux kernel internals, compilers, debuggers, and low-level system libraries are all ultimately constrained by the behaviors and guarantees formalized in these documents. The Intel manual set is organized into **conceptual layers**, progressing from architectural overview to detailed bit-level encodings.

8.1.1 Volume Structure

Modern x86-64 documentation is split into major volumes:

Volume	Focus	Content Scope
Volume 1	<i>Basic Architecture</i>	Execution environment, register models, memory addressing, privilege levels, system environment state
Volume 2	<i>Instruction Set Reference</i>	Complete instruction encoding, operand semantics, side effects, flag impact, latency notes
Volume 3	<i>System Programming</i>	Paging, segmentation, interrupt architecture, task switching, virtualization controls, system-level registers

Supplementary documents provide microarchitecture-specific details, errata, performance tuning guidance, and virtualization extensions.

8.1.2 Architectural vs. Microarchitectural Information

Intel distinguishes between:

Category	Definition	Relevance
Architectural Behavior	Behavior guaranteed across all compliant processors	Forms the basis for OS and compiler correctness
Microarchitectural Behavior	Implementation details that may vary between CPU generations	Impacts performance tuning and pipeline optimization

For example:

```
add rax, rbx ; Architectural effect: rax = rax + rbx
```

The execution width, scheduling ports, and latency characteristics of this instruction are microarchitecture-dependent.

8.1.3 Instruction Format in Volume 2

Every instruction entry defines:

- **Opcode encoding** (including legacy and REX prefixes)
- Allowed **operand types** (register, memory, immediate)
- ModR/M and SIB byte format rules
- **RFLAGS** bits modified
- Exceptions that may be raised
- Behavior in protected, long, and compatibility modes

Example excerpt (conceptual):

```
ADD r/m64, r64
Opcode: 01 /r
RFLAGS: CF, ZF, SF, OF updated
```

This format allows the developer to predict how the instruction interacts with:

- Register renaming
- Dependency chains
- Branch and flag-dependent execution

8.1.4 System Programming Details in Volume 3

Volume 3 specifies:

- **Page table formats** and access checks
- **Interrupt and exception dispatch**
- **Privilege transitions (Ring 3 → Ring 0)**
- **MSRs (Model-Specific Registers)**
- **Debug register behavior**
- **Virtualization controls (VMX / Intel VT-x)**

For Linux kernel work, Volume 3 provides the **exact hardware contract** that the kernel must satisfy to safely manage memory, scheduling, and trap handling.

8.1.5 Reading Model for Practical Development

To use the manuals effectively in low-level Linux programming:

Goal	Primary Reference
Understand execution and registers	Volume 1
Implement or analyze assembler output	Volume 2
Modify kernel or page tables, write exception handlers	Volume 3

The manuals are *specifications*, not tutorials. They define what *must* be true, not how to design efficient code. Performance insights require combining architectural rules with microarchitectural timing knowledge.

8.1.6 Key Principle

The Intel manuals define the **architectural contract**. Every compiler, debugger, kernel path, and calling convention is grounded in this specification. To write correct system-level code in Linux:

- **Volume 1 dictates what registers mean.**
- **Volume 2 dictates what instructions do.**
- **Volume 3 dictates how the CPU interacts with the operating system.**

Anything beyond these volumes is performance tuning—not correctness.

8.2 AMD APM Structure

While Intel’s manuals define the architectural baseline of the x86-64 ISA, AMD provides the **AMD64 Architecture Programmer’s Manual (APM)** series, which formalizes AMD’s extensions, implementation guidance, and architectural clarifications. The AMD APM is particularly relevant because AMD originated the **64-bit extension of x86**, meaning many architectural behaviors in modern Linux systems derive directly from AMD’s specifications. The APM does not simply rephrase Intel’s documentation; it highlights **differences, optional features, and extensions** that are essential for correctness and performance on AMD CPUs.

8.2.1 Organization of the APM Set

The AMD APM is divided into volumes, with each serving a distinct role:

Volume	Scope	Focus
Volume 1	Application Programming	Execution environment, registers, addressing, instruction behavior overview
Volume 2	System Programming	Paging, segmentation, privilege levels, memory protection, system registers
Volume 3	Instruction Reference	Complete instruction set reference (semantic-level detail)
Volume 4	Model-Specific Registers and Extensions	MSRs, virtualization controls, power states, debugging and monitoring extensions

Where Intel volumes emphasize architectural stability, AMD’s APM emphasizes **extension semantics and specific implementation constraints**.

8.2.2 AMD-Originated Architectural Elements

Key elements standardized from AMD’s design:

Feature	Description
64-bit mode	Unified register extension and flat virtual address model
REX Prefix	Enables access to extended register set (R8–R15) and 64-bit operand size
RIP-Relative Addressing	Used extensively in modern Linux position-independent code
SYSCALL / SYSRET	Low-overhead system call mechanism chosen by Linux

Example system call entry (Intel syntax):

```
mov rax, 1          ; sys_write
syscall             ; Uses AMD-specified calling semantics
```

Linux on x86-64 defaults to the AMD-originated syscall ABI even when running on Intel CPUs.

8.2.3 Memory System and Paging Clarifications

AMD’s documentation provides explicit definitions of:

- Canonical address rules (upper address bits must be sign-extended)
- Large page behavior for 2 MB and 1 GB pages
- TLB invalidation semantics under multi-core and multi-threading
- Memory attribute precedence (PAT, MTRRs, page flags interplay)

Such details are essential for:

- Kernel page table manipulation
- NUMA-aware memory allocators
- High-performance shared memory IPC implementations

8.2.4 Model-Specific Registers (Volume 4)

AMD's MSRs expose low-level controls for:

Subsystem	Controls include
Performance counters	Core event selection, sampling, hardware profiling
Power and frequency scaling	P-state and C-state transition configuration
Virtualization	Nested paging (NPT) and VMCB controls
Debugging	Hardware breakpoints, watchpoints, trap routing behavior

Linux performance tooling (e.g., `perf`, `oprofile`) relies heavily on the accuracy of this documentation.

8.2.5 Areas Where AMD APM Adds Beyond Intel

Topic	AMD Contribution
Virtualization	SVM (Secure Virtual Machine) model and VMCB privilege control structure
Memory Management	NPT (Nested Page Tables) for hardware-assisted virtualization
Power Efficiency	Explicit architectural state transitions for energy-aware scheduling

Topic	AMD Contribution
Performance Monitoring	Expanded event models tailored to Zen microarchitectures

These influence Linux behavior in scheduling, NUMA balancing, and virtualization performance tuning.

8.2.6 Practical Interpretation for Low-Level Development

For Linux low-level work, the manuals should be used as follows:

Task	Primary Source
Understanding instruction semantics	AMD APM Volume 3
Configuring system registers for paging or privilege transitions	AMD APM Volume 2
Writing syscall stubs or ABI-bound assembly	AMD APM Volume 1
Using performance counters, debugging microarchitecture timing	AMD APM Volume 4

In performance-critical contexts, Intel and AMD CPUs behave differently even under the same ISA. The AMD APM is therefore **not optional** reading for low-level developers targeting Linux on heterogeneous x86-64 systems.

8.3 How to Navigate Instruction Tables

Instruction reference tables in the Intel and AMD manuals are structured to provide **encoding**, **operand rules**, **flag behavior**, and **exception conditions** for each instruction. Understanding how to efficiently read these tables is essential when writing hand-crafted assembly, validating compiler output, debugging low-level behavior, or predicting microarchitectural side effects.

The manuals are not tutorials—they are **specifications**. Effective use requires knowing where to look and how to interpret the structure consistently.

8.3.1 Opcode and Encoding Fields

Instruction descriptions begin with the opcode encoding. This encoding indicates the byte-level form of the instruction, including prefixes and ModR/M fields.

Example entry (conceptualized from the manual):

```
ADD r/m64, r64
Opcode: 01 /r
```

Meaning:

- 01 is the primary opcode byte.
- /r indicates that one operand is encoded in the **ModR/M** byte.

When decoding or encoding manually, the **ModR/M** byte selects registers or memory addressing modes.

8.3.2 Operand Types and Naming

Operands in the tables follow standardized notation:

Symbol	Meaning
r64	64-bit general-purpose register
r/m64	64-bit register or memory operand
imm32	32-bit immediate value (sign-extended in 64-bit mode)
xmm, ymm	SIMD registers

Example interpretation:

```
mov rax, [rbx + rcx*4 + 16]
```

Matches operand pattern:

```
mov r64, r/m64
```

Thus:

- First operand is a register.
- Second operand is memory referenced through ModR/M + SIB + displacement.

8.3.3 RFLAGS Effects

Every instruction entry specifies which **condition flags** are modified. Understanding these effects is critical for correct branching and arithmetic logic.

Example excerpt:

```
RFLAGS: CF, ZF, SF, OF updated; PF undefined; AF undefined
```

Implications:

- Branches dependent on ZF, SF, or OF are valid immediately.
- Avoid relying on PF or AF unless explicitly required.

Efficient low-level code avoids dependency chains based on condition flags unless necessary.

8.3.4 Exceptions and Fault Conditions

Instruction descriptions list **conditions that trigger CPU exceptions**, typically referencing:

Condition Type	Result
Invalid memory address	Page fault (#PF)
Write to read-only memory	General protection fault (#GP)
Divide-by-zero	Divide error (#DE)
Misaligned access (when required)	Alignment check fault (#AC)

This allows determining **whether the instruction can restart** after fault handling.

8.3.5 Mode-Specific Behavior

The tables specify differences between:

- **64-bit mode**
- **Compatibility mode**
- **Legacy 32-bit mode**

Examples:

- Some encodings require **REX prefixes** to access R8 {R15}.
- Operand size defaults differ between 32-bit and 64-bit modes.
- System instructions may be privileged (Ring 0 only).

When writing assembler for Linux:

Always assume **64-bit long mode** and verify whether REX prefixes are needed for register access.

8.3.6 Finding Latency and Throughput Information

Instruction tables **do not** include performance characteristics. These are microarchitectural and vary by CPU generation. Performance tuning requires consulting:

- Vendor optimization manuals
- Empirical profiling (e.g., `perf`, `pmu-tools`)
- Microbenchmarking

Correctness is defined in the instruction tables; performance is discovered separately.

8.3.7 Practical Workflow for Low-Level Programming

Task	Refer To	Notes
Determine correct opcode and operands	Instruction tables (Volume 2 or APM Volume 3)	Verify addressing modes
Understand privilege or system effects	System programming volumes (Intel Vol. 3 / AMD Vol. 2)	Required for OS-level and syscalls
Determine performance	Microarchitecture-specific guides and profiling	Varies across CPU families

The instruction tables define the **legal contract** between assembly code and the hardware. Mastery of these tables allows precise, predictable behavior across Linux system layers, compilers, linkers, and kernel boundaries.

Chapter 9

Practical Register Tracing with GDB

9.1 Inspecting Registers

When analyzing program behavior at the machine level, the first step is to observe the **state of the processor registers** during execution. On Linux, the standard debugger for low-level tracing is **GDB**, which provides direct access to **general-purpose registers**, **RFLAGS**, **instruction pointer (RIP)**, **stack pointer (RSP)**, **frame pointer (RBP)**, and SIMD register sets.

Inspecting registers is essential for verifying calling convention correctness, diagnosing memory corruption, tracing control flow, and validating hand-written assembly.

9.1.1 Entering a Debugging Session

Typical workflow:

```
gdb ./program
```

or attach to a running process:

```
gdb -p <pid>
```

Set a breakpoint to inspect registers at a specific location:

```
(gdb) break *0x400800
```

```
(gdb) run
```

After stopping, register state can be examined before the instruction executes.

9.1.2 Displaying General-Purpose Registers

To view all standard integer registers:

```
(gdb) info registers
```

Example output (conceptual):

```
rax 0x7fffffffde20  rbx 0x0
rcx 0x1             rdx 0x7ffff7dd18c0
rdi 0x1             rsi 0x7fffffffddf0
rbp 0x7fffffffddc0  rsp 0x7fffffffddb0
rip 0x4006a0 <main+16>
```

Key interpretation points:

Register	Typical Role
RIP	Address of next instruction
RSP	Top of stack (stack grows downward)
RBP	Frame pointer (if used)
RAX	Return value or accumulator
RDI, RSI, RDX, RCX, R8, R9	Function argument registers (System V AMD64 ABI)

9.1.3 Inspecting Flags (RFLAGS)

```
(gdb) info registers rflags
```

Example conceptual output:

```
rflags 0x246 [ ZF PF ]
```

Common flags of interest:

Flag	Meaning
ZF	Zero result
SF	Sign bit
CF	Unsigned overflow / carry
OF	Signed overflow

Flags are critical for analyzing conditional branches and arithmetic behavior.

9.1.4 Inspecting SIMD Registers (XMM/YMM)

To display XMM registers:

```
(gdb) info registers xmm0 xmm1 xmm2
```

For YMM (256-bit):

```
(gdb) print $ymm0
```

Example conceptual:

```
$1 = {v4_float = {1.0, 2.0, 3.0, 4.0}}
```

SIMD examination is vital when debugging vectorized code, multimedia operators, or cryptographic routines.

9.1.5 Modifying Registers During Debugging

Registers may be modified to test alternate control paths:

```
(gdb) set $rax = 0
(gdb) set $rip = 0x400750
```

This is particularly useful for:

- Forcing a specific branch outcome
- Reconstructing corrupted state
- Skipping faulty instructions during analysis

9.1.6 Recording State Across Breakpoints

For repeated inspection:

```
(gdb) display/16gx $rsp
(gdb) display $rax
(gdb) display/i $rip
```

These update each time the debugger stops.

9.1.7 Practical Guidelines

Guidance	Explanation
Break after function prologue to study stack layout	Validates calling convention behavior
Always inspect RSP before modifying registers	Ensures stack remains well-formed

Guidance	Explanation
Compare RIP and disassembly side by side	Clarifies execution flow
Use YMM/XMM views when debugging optimized builds	Compilers aggressively vectorize loops

Register tracing is not just observation—it provides insight into the **runtime execution model**, enabling reliable reasoning about correctness, data flow, and performance-critical behavior.

9.2 Single-Step Execution

Single-step execution allows the developer to observe program behavior one instruction at a time. On x86-64 Linux, GDB leverages the processor's built-in **Trap Flag (TF)** in `RFLAGS` to trigger a debug trap after each instruction retires. This provides precise insight into register changes, control flow decisions, stack behavior, and memory state transitions — making it indispensable for low-level debugging and reverse engineering.

9.2.1 Entering Single-Step Mode

While in GDB, after hitting a breakpoint or interrupt:

```
(gdb) stepi
```

or shorthand:

```
(gdb) si
```

Each `si` command executes **exactly one machine instruction** and stops again, allowing state inspection between steps.

To step multiple instructions while remaining in single-step mode:

```
(gdb) stepi 10
```

This executes 10 instructions before stopping again.

9.2.2 Viewing the Current Instruction

To disassemble the surrounding code at the instruction pointer (RIP):

```
(gdb) x/i $rip
```

or continuously track it:

```
(gdb) display/i $rip
```

Example conceptual output:

```
=> 0x4006a0 <main+16>:  mov rax, rdi
```

This ensures you always see the exact instruction about to execute.

9.2.3 Observing Register Changes Per Instruction

Before stepping:

```
(gdb) info registers
```

Execute one step:

```
(gdb) si
```

Inspect again:

```
(gdb) info registers
```

This reveals **data dependencies**, such as which instructions update RAX, RSP, or flags.

For automated continuous tracking:

```
(gdb) display $rax
```

```
(gdb) display $rsp
```

```
(gdb) display $rflags
```

These update with each step, providing a **per-instruction delta view**.

9.2.4 Branch and Loop Analysis

Single-stepping allows inspection of branch behavior:

```
cmp rax, rbx
jle .Lsmall_case
```

Step through:

```
(gdb) si
(gdb) si
```

Observe:

- Whether ZF, SF, or OF influenced the branch
- Whether RIP moved sequentially or jumped

This is especially important when debugging **off-by-one**, **sign mismatch**, or **overflow** logic errors.

9.2.5 Stepping Through Function Calls

To step *into* a function:

```
(gdb) stepi
```

To execute the call *without* entering the function body:

```
(gdb) nexti
```

This distinction is crucial when analyzing control flow vs. internal logic.

9.2.6 Single-Stepping and System Instructions

Some instructions cause architectural transitions:

Instruction	Effect
<code>syscall</code>	Switches to Ring 0, kernel stack
<code>ret</code>	Pops return address → changes RIP
<code>leave</code>	Restores frame pointer and stack pointer
<code>iretq</code>	Returns from interrupt or exception handler

Single-stepping across these verifies:

- **stack frame correctness**
- **ABI compliance**
- **kernel entry/exit transitions**

9.2.7 Practical Considerations

Risk	Explanation
Single-step overhead is very high	Each trap invokes kernel handling
Concurrency behavior may change	Debug traps serialize execution
Timing-sensitive bugs may disappear	Debugging alters scheduling and memory visibility

Therefore, single-step execution is a **precision diagnostic tool**, not a general-purpose runtime inspection method.

9.2.8 Summary

Single-stepping enables:

- Precise instruction-by-instruction understanding
- Immediate observation of register and flag changes
- Accurate analysis of branch conditions and call/return paths
- Verification of ABI and stack correctness

It is fundamental when working at the **assembly, ABI, and calling convention level** on Linux x86-64.

9.3 Understanding Disassembly Flow

When debugging at the machine level, the **disassembly** shown by GDB represents the actual instruction sequence the CPU executes, not the instruction order the programmer wrote in source form. Optimizations, inlining, register allocation, instruction scheduling, loop unrolling, and elimination of unused computation all reshape the control and data flow visible at the assembly level. Understanding how to interpret disassembly is essential for tracing program execution accurately and diagnosing low-level issues.

9.3.1 Entry Points and Instruction Context

To examine the disassembly at the current instruction pointer:

```
(gdb) x/10i $rip
```

This displays the next ten instructions to execute. Disassembly must be understood **in context**, which includes:

- Current value of RIP
- Visible register state (`info registers`)
- Stack frame layout (`info frame`)
- ABI-defined parameter passing rules

A standalone instruction rarely provides full meaning without these surrounding constraints.

9.3.2 Sequential vs. Control-Flow Transitions

Control flow in disassembly transitions through:

Instruction Category	Effect on Flow
Direct control flow (<code>jmp</code> , <code>call</code> , <code>ret</code>)	Explicit change in RIP
Conditional branches (<code>jz</code> , <code>jnz</code> , <code>jb</code> , <code>jbe</code> , ...)	Based on RFLAGS status
Computed jumps (<code>jmp rax</code> , <code>call [rbx]</code>)	Indirect control transfer

Example:

```
cmp rax, rbx
jl  .Lless
mov rdx, rax
```

If `SF` $\neq$ `OF`, the jump is taken and flow moves to `.Lless`; otherwise, execution continues sequentially.

Understanding the **dependency between arithmetic instructions and RFLAGS** is key to tracing execution.

9.3.3 Function Calls and Stack Frame Interpretation

Upon encountering a call:

```
call 0x400750 <func>
```

Control switches to the called function, and `RSP` is updated. After stepping into the call:

```
(gdb) stepi
```

Check stack state:

```
(gdb) info registers rsp rbp rip
```

If the function uses a frame pointer:

```
push rbp
mov rbp, rsp
```

This establishes a stable frame lattice:

- $RBP + \text{offset} \rightarrow$ parameters and saved state
- RSP moves for local storage

Recognizing frame structure allows reliable tracing through nested call stacks.

9.3.4 Understanding Compiler Optimization Effects

Optimized builds change disassembly flow substantially:

Optimization Effect	Observable Change in Disassembly
Inlining	Functions appear merged; no call boundary
Register allocation	Variables exist only in registers
Dead-code elimination	Some expected instructions do not exist
Loop unrolling	Multiple loop iterations appear inline
Instruction scheduling	Instructions reordered for execution pipeline efficiency

Thus, **expect disassembly to diverge from source order** when optimization is enabled.

To temporarily disable optimizations when learning:

```
gcc -O0 -g program.c -o program
```

9.3.5 Branch Targets and Symbol Resolution

To display symbolic labels in disassembly:

```
(gdb) set disassemble-next-line on
```

To see cross-reference of labels and addresses:

```
(gdb) info functions
```

Recognizing label structure (`.L0`, `.Ltmp1`, `.Lend`) is essential when inspecting compiler-generated control flow.

9.3.6 Memory Address Interpretation

Indirect addressing must be decoded carefully:

```
mov rax, [rbx + rcx*4 + 16]
```

Meaning:

- Base: RBX
- Index: RCX, scaled by 4
- Displacement: +16 bytes
- Loads from calculated memory address

Understanding addressing forms is necessary to trace pointer logic and data structure traversal.

9.3.7 Practical Strategy for Reading Disassembly

Step	Purpose
Identify <code>RIP</code> and show surrounding disassembly	Establish current execution point
Inspect register set	Understand operand sources
Observe flags before and after compares	Predict branch outcomes
Examine memory operands using <code>x/gx</code> <address>	Reveal structure layout
Track call/return patterns	Validate stack correctness

The key principle:

Disassembly describes **what the CPU will do, not what the source code intended.**

Effective debugging requires bridging this gap by interpreting register state, memory layout, and ABI-driven calling conventions in real time.

Chapter 10

Project

10.1 Choose an Instruction

The goal of this project phase is to select a **single x86-64 instruction** and analyze it thoroughly from the perspectives of **encoding, execution behavior, register and flag interactions, and practical use on Linux**. By reducing scope to one instruction, we isolate microarchitectural effects and observe how a single machine-level operation propagates through registers, the pipeline, and surrounding program logic.

This exercise emphasizes comprehension over performance. The objective is to understand *exactly how the instruction behaves*, not merely that it “works.”

10.1.1 Criteria for Instruction Selection

You may choose **any** x86-64 instruction, but the instruction must:

Requirement	Reason
Have clearly defined input and output operands	Enables explicit register tracing

Requirement	Reason
Affect at least one flag or architectural register	Allows observation of state transitions
Execute meaningfully in user mode on Linux	Avoids privileged/system-only instructions
Have non-trivial behavior under different operand patterns	Enables controlled experimentation

Suitable categories of instructions:

Category	Example Instructions	Notes
Integer arithmetic	add, sub, imul, idiv	Clear flag effects, easy to trace
Bitwise / logic	and, or, xor, shl, shr	Observably modifies specific bits in place
Data movement	mov, lea, push, pop	Demonstrates addressing and stack effects
Control flow	jmp, call, ret, jz, jnz	Shows branch and RFLAGS interaction
Memory access	mov [rdi], rax, mov rbx, [rsi]	Influenced by TLB/cache locality

Avoid at this stage:

- System instructions (syscall, wrmsr, rdtsc, invlpg)
- Privileged control registers (cr0{cr4, cr8})
- SIMD instructions (optional later stage)

These add complexity before the core pipeline model is established.

10.1.2 Example Instruction Selection: `add rax, rbx`

The instruction:

```
add rax, rbx
```

is an ideal teaching candidate because it:

- Has explicit input (RAX, RBX) and output (RAX) dependencies.
- Modifies key RFLAGS bits (SF, ZF, CF, OF).
- Interacts visibly with register renaming and out-of-order scheduling.
- Can be traced step-by-step in GDB.

This instruction also allows variation:

```
add eax, ebx      ; 32-bit operand, zero-extends into RAX
add al, bl         ; 8-bit operand, partial-register dependency behavior
```

Demonstrating how operand width affects pipeline behavior.

10.1.3 Verify No Compiler Interference

For pure instruction-level study, compile with:

```
gcc -nostdlib -static -o project project.s
```

Or restrict optimization:

```
gcc -O0 -g project.s -o project
```

This ensures the instruction appears **exactly** as written in the disassembly, without reordering or substitution.

10.1.4 Define the Measurement and Observation Plan

For the selected instruction, you will:

1. **Trace input register state** before execution.
2. **Execute** the instruction **one step** using GDB (`stepi`).
3. **Inspect changed registers and flags**.
4. Repeat with **different operand values** to observe flag behavior.
5. Repeat with **memory-based operands** to observe addressing effects:

```
add rax, [rdi]
```

6. Document execution differences under:
 - Different operand sizes
 - Different addressing modes
 - Different register dependency patterns

10.1.5 Expected Learning Outcome

By the end of this step, you should be able to:

- Predict the effect of the chosen instruction **before execution**.
- Confirm results by inspecting **registers and flags** in GDB.
- Understand the difference between architectural description and **actual execution**.
- Explain how operand selection influences data dependency and pipeline behavior.

This single instruction will serve as the foundation for the remaining project steps, where it will be traced across:

- The decode and μ op generation pipeline
- Register renaming and dependency scheduling
- Memory hierarchy interactions (if memory operands are used)
- GDB tracing and disassembly inspection

10.2 Trace Its Execution Step-by-Step

Having selected an instruction, the next objective is to **observe its behavior directly** by tracing execution one instruction at a time. This reveals the real effects on registers, flags, stack layout, control flow, and—when memory operands are used—the underlying addressing and memory access patterns. The goal is not to infer what the instruction *should* do, but to observe **what it actually does** as the CPU executes it.

For clarity, we assume the example instruction chosen is:

```
add rax, rbx
```

10.2.1 Prepare a Minimal Test Harness

Write the instruction in a standalone assembly file to prevent compiler transformations:

```
; project.s
global _start
section .text
_start:
    mov    rax, 5
    mov    rbx, 7
    add    rax, rbx
    mov    rdi, rax
    mov    rax, 60        ; exit syscall
    syscall
```

Assemble and link:

```
nasm -f elf64 project.s -o project.o
ld project.o -o project
```

10.2.2 Launch Under GDB

```
gdb ./project
```

Set a breakpoint just before the target instruction:

```
(gdb) break _start
```

```
(gdb) run
```

Step to the add instruction:

```
(gdb) stepi
```

```
(gdb) stepi
```

```
; Execution now positioned at `add rax, rbx`
```

10.2.3 Record the Pre-Execution State

Inspect registers:

```
(gdb) info registers rax rbx rflags
```

Conceptual example:

```
rax = 0x5  
rbx = 0x7  
rflags = 0x202
```

Note:

- RAX = 5 is the accumulator.
- RBX = 7 is the operand.
- RFLAGS holds the current arithmetic state (ZF, SF, CF, OF not yet affected).

10.2.4 Execute the Instruction One Step

```
(gdb) stepi
```

Now inspect:

```
(gdb) info registers rax rflags
```

Expected conceptual result:

```
rax = 0xc          ; 5 + 7 = 12 (decimal)
rflags: ZF=0, SF=0, CF=0, OF=0
```

Flags interpretation:

Flag	Meaning in This Case
ZF = 0	Result is non-zero
SF = 0	Result's sign bit is 0 (positive)
CF = 0	No unsigned overflow
OF = 0	No signed overflow

Thus, the *observed hardware state* matches the *architectural definition* of addition.

10.2.5 Vary Operand Types for Comparative Tracing

Change operand sizes to observe partial-register and zero-extension effects:

```
add eax, ebx      ; Writes 32 bits and zero-extends into RAX
add al, bl        ; 8-bit operation with partial-register hazards
```

Trace again and record:

- Whether upper bits change

- Whether flags differ
- Whether register renaming introduces unexpected stalls (seen via timing changes)

This step demonstrates that **operand width directly affects architectural and microarchitectural state**.

10.2.6 Trace Memory Operand Variants

Now test:

```
add rax, [rdi]
```

In GDB:

```
(gdb) x/gx $rdi      ; Inspect source memory
(gdb) stepi           ; Execute add
(gdb) info registers rax rflags
```

This reveals:

- How effective address computation is performed
- Whether a page fault would occur
- How memory locality affects data visibility

10.2.7 Document Observations Objectively

For each variant, capture:

Variant	Registers Before	Registers After	Flags Before/After	Notes

This **formal observation log** is a core deliverable for the project.

10.2.8 Key Takeaway

Tracing one instruction in isolation confirms:

- The **architectural state transitions** (registers + flags)
- The **execution semantics** match the ISA description
- operand width and addressing mode **materially influence instruction behavior**

This forms the conceptual bridge between **manual specification (manual volumes)** and **runtime machine execution (GDB trace)**.

10.3 Generate a Microarchitectural Report

After tracing the architectural effects of the chosen instruction (register and flag changes), the next step is to analyze how the instruction behaves inside the **microarchitecture** — the internal pipeline that executes instructions beyond the architectural abstraction. The microarchitectural report documents **how** the CPU executes the instruction: how many μ ops it generates, which execution units it uses, how it interacts with the reorder buffer, and whether data dependencies or memory hierarchy effects influence timing.

The goal is to translate **observed execution behavior** into a structured understanding of **pipeline utilization and performance constraints**.

10.3.1 Identify μ op Count and Execution Units

Every instruction decomposes into one or more **micro-operations (μ ops)**. The mapping is defined by the CPU implementation, not the ISA specification.

For the selected instruction (example):

```
add rax, rbx
```

Typical μ op behavior (for common x86-64 cores):

Instruction	μ op Count	Execution Port(s)	Resource Type
add rax, rbx	1 μ op	ALU pipeline (port 0 or port 1, core-dependent)	Integer arithmetic

If memory addressing is used:

```
add rax, [rdi]
```

The micro-op breakdown changes:

Component	μop Type
Address generation	AGU μop
Memory load	Load μop
Arithmetic	ALU μop

Total: **3 μops**, unless fused by the microarchitecture.

This difference is critical:

Register-to-register arithmetic is cheap; memory-to-register arithmetic is not.

10.3.2 Latency vs. Throughput Interpretation

Two timing metrics must be recorded:

Metric	Meaning
Latency	The number of cycles before the result is ready for dependent instructions
Throughput	The number of such instructions that can complete per cycle in parallel

For example (approximate, model-dependent):

```
Integer add latency  1 cycle
Integer add throughput 1{2 per cycle
Load-use latency (L1 hit) 4{6 cycles
```

Your report must distinguish **dependency latency** (true data dependency) from **instruction dispatch bandwidth** (parallelism).

10.3.3 Characterize Register Dependencies

Record whether execution stalled due to operand readiness. For example, a dependency chain:

```
add rax, rbx
add rax, rcx
add rax, rdx
```

forces serialized execution, because each `add` waits on the previous result.

In contrast:

```
mov r8, rbx
add r8, rcx
add rax, rdx
add rax, r8
```

allows partial parallel execution.

Document which registers were updated in parallel and which created bottlenecks.

10.3.4 Memory Hierarchy Impact (If Memory Operand Was Used)

For `add rax, [rdi]`, determine:

Level	Approximate Access Cost	Notes
L1 hit	4 cycles	Fast, typically hidden via OoO execution
L2 hit	10–12 cycles	May stall dependent μ ops
L3 hit	30–45 cycles	More severe stall
DRAM	150–300+ cycles	Full stall unless other work is parallelizable

This part of the report explains **why dependent code may pause**, even if the arithmetic is trivial.

10.3.5 Report Format

Your microarchitectural report should consist of:

1. Instruction Form Studied

e.g., `add rax, rbx` and/or `add rax, [rdi]`

2. μ op Breakdown

List μ ops produced, and which execution ports they use.

3. Latency and Throughput

Specify expected vs. observed behavior.

4. Dependency Effects

Demonstrate behavior with and without chaining.

5. Memory Interaction (If Applicable)

Identify L1/L2/L3/DRAM impact using controlled experiments.

6. Conclusions

Summarize the most significant performance determinants.

10.3.6 Example Summary (Conceptual)

- ``add rax, rbx`` translates to a single ALU μ op with ~1-cycle latency.
- ``add rax, [rdi]`` expands to AGU + load + ALU μ ops.
- Dependent arithmetic sequences serialize execution.
- Memory operands introduce variable latency based on cache state.
- Core execution throughput is limited by ALU port availability and
↪ dependency chains.

10.3.7 Key Insight

The ISA defines what the instruction **means**.

The microarchitecture determines how **fast** it executes.

This report closes the gap between architectural observation and performance-relevant internal execution behavior.

Conclusion

The material presented in this booklet has traced the fundamental components of execution on x86-64 Linux systems, moving from architectural definitions to practical observation. We examined how instructions are fetched, decoded, executed, and committed; how registers and flags shape program state; how memory is translated and cached; how interrupts and exceptions transition control to the kernel; and how these mechanisms become visible in debugging and low-level inspection.

The central lesson is that software does not run in isolation. Every instruction interacts with layers of hardware translation, privilege control, caching, prediction, and scheduling. Understanding these layers grants the programmer not only the ability to write or analyze low-level code, but the ability to reason about performance, correctness, and security with precision.

Modern processors are highly parallel, speculative, and adaptive. Their behavior cannot be fully understood from the instruction set alone. Instead, the programmer must distinguish between:

The architectural contract (what the instruction set guarantees),

And the microarchitectural implementation (how the processor achieves it).

This distinction is critical in systems programming, reverse engineering, security research, and performance optimization. It is also the foundation for understanding why different programs behave differently on different CPUs, why performance varies under different memory layouts, and why certain coding patterns cooperate better with the hardware pipeline than others.

The final project demonstrated that mastery begins not with memorizing all instructions, but with deeply understanding one instruction at a time, traced through registers, flags, memory dependencies, and microarchitectural effects. This approach scales: the same methodology applies to any instruction, any subsystem, and any execution context.

The work does not end here. Every concept introduced—paging, TLB behavior, speculative execution, vector registers, system calls—opens into deeper study. But the framework is now established. The reader should now be equipped not simply with knowledge, but with a method: observe, measure, correlate specification with execution, and draw conclusions grounded in real behavior.

The essence of low-level programming is clarity. Clarity of execution, clarity of data movement, and clarity of control flow. Once clarity is achieved, complexity becomes manageable.

— End of Booklet

References

1. intel-sdm Intel Corporation, *Intel 64 and IA-32 Architectures Software Developer's Manuals*, Volumes 1–3, latest revision. Authoritative specification of the x86-64 instruction set, privilege levels, paging, and system behavior.
2. amd-apm Advanced Micro Devices (AMD), *AMD64 Architecture Programmer's Manual*, Volumes 1–4, latest revision. Defines AMD 64-bit architectural extensions, memory model, and system register behavior.
3. sysv-abi System V Application Binary Interface, *AMD64 Architecture Processor Supplement*, latest maintained edition. Formal definition of calling conventions, register usage, stack layout, and ELF executable format.
4. linux-kernel-docs Linux Kernel Documentation, maintained in the official kernel source tree. Covers virtual memory implementation, scheduling, and system call interface behavior.
5. kerrisk Michael Kerrisk, *The Linux Programming Interface*. No Starch Press. Comprehensive reference on POSIX system calls and process environment fundamentals.
6. ldd Jonathan Corbet, Alessandro Rubini, Greg Kroah-Hartman, *Linux Device Drivers*, 3rd Edition. Conceptually relevant reference on privileges, kernel entry points, and

driver interaction.

7. intel-optimization Intel Corporation, *Intel 64 and IA-32 Architectures Optimization Reference Manual*. Defines microarchitectural execution characteristics, pipeline resource usage, and performance heuristics.
8. amd-zen-optimization Advanced Micro Devices (AMD), *Optimization Guidelines for AMD Zen Microarchitectures*. Describes port mappings, branch prediction behavior, TLB organization, and power considerations.
9. patterson-hennessy David Patterson and John Hennessy, *Computer Organization and Design: RISC-V Edition*. Morgan Kaufmann. Provides foundational insight into pipelining and memory hierarchies applicable to x86-64 understanding.
10. csapp Randal E. Bryant and David R. O'Hallaron, *Computer Systems: A Programmer's Perspective*, 3rd Edition. Explains linking, stack frames, memory layout, and machine-level program behavior.
11. gdb Free Software Foundation, *GDB User Documentation*. Defines debugger command semantics, register inspection, stepping models, and disassembly methods.
12. binutils GNU Project, *Binutils Manual Pages* (objdump, readelf, nm). Reference for ELF inspection, symbol tables, relocations, and executable structure.
13. xed Intel Corporation, *Intel XED Instruction Encoder/Decoder Reference*. Formal encoding structure for x86 instructions and operand decoding logic.
14. capstone Capstone Disassembly Framework, *Capstone Documentation*. Guidance for programmatic disassembly and instruction stream analysis.
15. tanenbaum-os Andrew S. Tanenbaum and Herbert Bos, *Modern Operating Systems*, 4th Edition. Conceptual background for process scheduling, virtual memory, and protection mechanisms.

16. tanenbaum-sco Andrew S. Tanenbaum, *Structured Computer Organization*, 6th Edition.
Clear conceptual explanation of hardware layering and control sequencing.