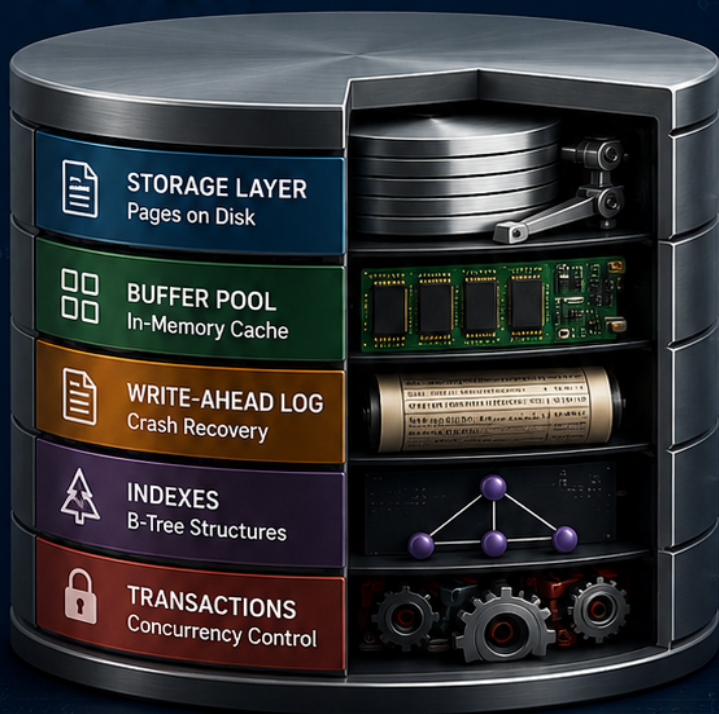


BUILDING DATABASE ENGINES FROM SCRATCH

BUILD A
MODERN
RELATIONAL
DATABASE
IN RUST

A HANDS-ON GUIDE TO STORAGE, CONCURRENCY,
QUERY PROCESSING, AND RELIABILITY IN
SYSTEMS PROGRAMMING



PERSISTENT
STORAGE



CONCURRENCY
& ISOLATION



CRASH
RECOVERY



B-TREE
INDEXES



COST-BASED
OPTIMIZATION



JOINS, SORTS &
AGGREGATIONS

JASON REED

Building Database Engines from Scratch

A Hands-On Guide to Storage, Concurrency, Query Processing, and Reliability in Systems Programming

Steve Publications

This book is available at

<https://leanpub.com/buildingdatabaseenginesfromscratch>

This version was published on 2026-08-30



This is a [Leanpub](#) book. Leanpub empowers authors and publishers with the Lean Publishing process. [Lean Publishing](#) is the act of publishing an in-progress ebook using lightweight tools and many iterations to get reader feedback, pivot until you have the right book and build traction once you do.

© 2026 Steve Publications

Contents

A Hands-On Guide to Storage, Concurrency, Query Processing, and Reliability in Systems Programming	1
Chapter 1: Why Build a Database?	2
The Hidden Complexity of Every Database Query	2
What You Will Build and Why It Matters	3
A Tour of the Architecture We Are Going to Construct	5
Implementation Roadmap: From Empty Project to Working Database	7
How to Read This Book and Use the Code	8
Chapter 2: Project Setup and Systems Foundations	9
Choosing Rust for Database Implementation	9
Project Structure and Build Configuration	10
Core Abstractions: Pages, Records, and Offsets	13
Error Handling Strategy Across the Engine	18
First Binary: Initializing an Empty Database File	24
Chapter 3: Storage Fundamentals – Disks, SSDs, and Pages	28
How Disks and SSDs Actually Work	28
Why Databases Think in Pages, Not Bytes	28
Building Our Page Abstraction Layer	28
Binary Data Layout: Endianness, Alignment, and Serialization	28
Implementing the Storage Manager Interface	28
Chapter 4: File Format and Slotted Page Records	29
Designing the Database File Format	29
Slotted Page Layout: Headers, Directories, and Records	29
Variable-Length Record Serialization	29
Free-Space Management Within Pages	29
Implementing Page Reads and Writes with Verification	29

Chapter 5: Buffer Pool Manager – Caching at Scale	30
The Purpose of Buffer Pools and Why They Are Hard	30
Pinning Semantics and Reference Counting	30
Dirty Pages, Write Policies, and Flush Behavior	30
Handling Memory Pressure and Eviction Under Load	30
Chapter 6: Logging, Durability, and Crash Recovery	31
The Durability Problem and Why fsync Matters	31
Write-Ahead Logging: Protocol and Invariants	31
Designing Our Log Record Format	31
Checkpointing Strategies and Implementation	31
Crash Recovery: Redo and Undo in Practice	31
Chapter 7: Indexing – B-Trees, Hash Indexes, and Beyond	32
Why Indexes Exist and What They Trade Away	32
Implementing a B+ Tree from Scratch	32
Page Splits, Merges, and Concurrency Considerations	32
Hash Indexes for Equality Lookups	32
LSM Trees and Alternative Index Architectures	32
Chapter 8: Transactions and Concurrency Control	33
ACID Properties: What They Mean in Implementation Terms	33
Locking Protocols and Two-Phase Locking	33
Deadlock Detection and Prevention	33
Latches vs Locks: Protecting Data Structures	33
Implementing MVCC with Snapshot Isolation	33
Chapter 9: SQL Parsing and Semantic Analysis	34
Designing Our SQL Subset and Grammar	34
Lexical Analysis: Tokenizing SQL Input	34
Recursive Descent Parsing and AST Construction	34
Symbol Tables and Schema Resolution	34
Chapter 10: Query Planning and Cost-Based Optimization	35
From AST to Logical Plan	35
Physical Operators and Execution Strategies	35
Building a Cost Model from Statistics	35
Cardinality Estimation: Histograms and Selectivity	35
Join Ordering and Query Rewrite Rules	35

Chapter 11: Query Execution Engine	36
Operator Architecture: Iterators and Pipelining	36
Table Scans and Index Scans	36
Filtering, Projection, and Expression Evaluation	36
Join Implementations: Nested-Loop, Hash, and Merge	36
Aggregation and Grouping with Hash Tables	36
Chapter 12: Performance Engineering and Optimization	37
Understanding Database Performance Characteristics	37
I/O Patterns: Sequential Versus Random Access	37
Cache Efficiency and Data Locality	37
Vectorized Execution vs Tuple-at-a-Time Processing	37
Profiling, Benchmarking, and Workload Generation	37
Systematic Optimization: Finding and Fixing Bottlenecks	37
Chapter 13: Reliability, Correctness, and Failure Handling	39
Crash Consistency: Guarantees and Invariants	39
Torn Pages and Partial Writes	39
Corruption Detection with Checksums and Validation	39
Disk-Full Conditions and Safe Failure Modes	39
Backup, Restore, and Integrity Verification	39
Chapter 14: Beyond Single Node – Distributed Database Concepts	40
Why Single-Node Databases Eventually Hit Limits	40
Replication Architectures: Leader-Follower and Multi-Leader	40
Log Replication and Consistency Models	40
Sharding and Distributed Query Execution	40
Consensus, Fault Tolerance, and Rebalancing	40
Chapter 15: Conclusion – From Educational Engine to Production System	41
Architecture Review: How All Subsystems Fit Together	41
What We Simplified and Why Production Systems Differ	41
Comparing Our Implementation to Real Databases	41
Paths Forward: Extending Your Database Engine	41
The Engineer’s Mindset for Building Reliable Systems	41
References	42

A Hands-On Guide to Storage, Concurrency, Query Processing, and Reliability in Systems Programming

This book teaches you how modern database systems work by having you build one. Over fifteen chapters we will design and implement a complete relational database engine in Rust: storage layers that read and write pages to disk, a buffer pool manager that caches hot data in memory, write-ahead logging for crash recovery, B-tree indexes for fast lookups, a transaction system with proper isolation and concurrency control, a SQL parser, a cost-based query optimizer, and an execution engine that runs joins, sorts, and aggregations. Every chapter adds working code to the same project so that by the end you have a coherent, functional database application rather than a collection of isolated examples. If you are a software engineer who wants to understand what happens between typing `SELECT * FROM users` and seeing results on screen, this is the book for you.

Chapter 1: Why Build a Database?

The Hidden Complexity of Every Database Query

A simple query hides enormous engineering. When you type `SELECT id, name FROM users WHERE age > 25 ORDER BY name LIMIT 10` into any modern database and press Enter, something remarkable happens in milliseconds. The system parses your text into a structured representation, validates that the table and columns exist, decides how to find the matching rows most efficiently, possibly consults one or more indexes, reads data from storage while managing concurrent access with other queries, applies filters and sorts, and returns exactly ten rows in alphabetical order. If another connection is modifying users at the same time, your query still sees a consistent view of the data. If the power fails mid-execution, committed transactions are never lost and incomplete ones never partially appear. None of this is magic. It is the product of decades of research and engineering distilled into layers of carefully designed subsystems that work together under strict invariants.

Most developers interact with databases as black boxes. You write SQL, you get results, and occasionally you add an index when something runs slow. This works fine for using a database but leaves you unable to answer deeper questions. Why does one query plan outperform another? What actually happens when you commit a transaction? How does the system guarantee durability across crashes? Why do LSM trees excel at write-heavy workloads while B-trees dominate OLTP systems? When should you choose PostgreSQL over SQLite or RocksDB?

The answers lie inside the engine. Every architectural decision in a database solves a concrete problem under real constraints: latency, throughput, memory pressure, disk I/O limits, correctness guarantees, and failure modes. You cannot fully internalize these trade-offs by reading about them abstractly. The most effective way to learn is to build one yourself.

This approach has proven its value across systems programming education. People who have built their own web servers understand HTTP, connection management, and request routing in ways that users never will. People who

have implemented a filesystem understand journaling, allocation strategies, and directory structures at an intuitive level. The same is true for databases. When you write the code that implements a buffer pool replacement algorithm, you understand cache behavior more deeply than any diagram could convey. When you implement crash recovery from a write-ahead log, durability stops being an abstract ACID property and becomes a series of concrete operations with real failure modes you must handle.

The database field is particularly rich for this kind of learning because the problems are well-defined and the solutions have been studied exhaustively. Unlike many areas of software engineering where best practices evolve rapidly and opinions diverge wildly, core database algorithms are mature and thoroughly understood. The B-tree index structure dates to 1970 [1]. Write-ahead logging was formalized in the 1980s through foundational work by Jim Gray, Andreas Reuter, and others [2]. Multiversion concurrency control traces back to David Reed's 1978 dissertation [3]. These algorithms are not relics. They power every major database system in production today: PostgreSQL, MySQL/InnoDB, SQLite, Oracle, SQL Server, CockroachDB, TiDB, Redis, RocksDB, and countless others all use variants of these same core ideas.

Building a database also teaches you systems programming at its most demanding. A database engine must manage memory carefully without garbage collection interference, perform I/O efficiently while respecting storage hardware characteristics, coordinate concurrent threads without deadlocks, handle partial failures gracefully, and maintain strict invariants under all conditions. These are the skills that separate competent engineers from exceptional ones.

What You Will Build and Why It Matters

Over the course of this book we will implement a complete relational database engine called `mydb`. It will be written in Rust for reasons we will discuss in the next chapter, and it will grow incrementally across all fifteen chapters. Each chapter adds substantial working functionality that integrates with everything built before it. By the end you will have a single coherent codebase that implements:

A storage layer that reads and writes fixed-size pages to disk files using proper binary serialization. A buffer pool manager that caches frequently

accessed pages in memory, uses a replacement algorithm to evict cold pages, tracks dirty pages, and flushes them back to disk with appropriate write policies. A write-ahead logging system that records all modifications before they are applied to data pages, enabling crash recovery through redo and undo operations. A B-tree index implementation that supports fast key lookups, range scans, and concurrent access patterns. A transaction manager providing ACID semantics through locking protocols or multiversion concurrency control with configurable isolation levels. A SQL parser that tokenizes input text, builds abstract syntax trees for a practical subset of the SQL language, and performs semantic validation against schemas. A query planner and optimizer that transforms parsed queries into logical plans, applies rewrite rules, estimates costs using statistics, and selects efficient physical execution strategies. An execution engine implementing core relational operators: scans, filters, projections, sorts, aggregations, and multiple join algorithms including nested-loop, hash, and merge joins.

This is not a toy. While we will simplify certain aspects for educational clarity and explicitly call out where production systems do more, the core algorithms you implement are the same ones used in real databases. The B-tree you build will use the same split and merge operations as PostgreSQL's indexes. The write-ahead log will follow the same protocol that guarantees durability in SQLite and MySQL. The buffer pool replacement policy will address the same fundamental problem that every database faces: how to keep hot data in memory while cold data is evicted to disk.

The project structure is designed so that you can actually use mydb as a working command-line database application. You will be able to create tables, insert data, build indexes, run queries with WHERE clauses and JOINS, execute transactions with proper isolation, and recover from simulated crashes. Along the way we will include benchmarking utilities so you can measure performance characteristics and see the impact of optimizations firsthand.

Why does this matter for your career? Even if you never ship a database engine into production, the understanding you gain will make you a better engineer in every domain. You will understand why certain data structures excel under specific access patterns. You will develop intuition about memory hierarchies and cache behavior that applies to any performance-sensitive system. You will learn how to reason about concurrent execution and failure modes systematically. You will see how abstract correctness properties translate into concrete implementation requirements. These are transferable skills

that improve your ability to design, debug, and optimize any complex software system.

A Tour of the Architecture We Are Going to Construct

Before we start coding, let us take a high-level tour of the architecture. Understanding how pieces fit together will help you see why each component exists and what problems it solves. Think of a database engine as a layered stack where each layer provides abstractions that hide complexity from layers above.

At the lowest level sits physical storage: disks or SSDs where bytes are persisted across reboots. Storage hardware has specific characteristics that fundamentally shape database design. Reading data from disk is orders of magnitude slower than reading from memory, so databases minimize I/O through caching and batching. Sequential reads are dramatically faster than random reads, so databases organize data for scan efficiency. Writes have their own constraints: SSDs wear out with excessive writes, and fsync operations to guarantee durability introduce latency spikes. Every database architect must design around these hardware realities.

Above storage sits the storage manager layer, which organizes raw bytes into structured pages. A page is a fixed-size block of data, typically 4096 bytes, that serves as the fundamental unit of I/O. The storage manager handles reading pages from disk files into memory buffers and writing modified pages back to disk. It manages file formats, record layouts within pages, free space tracking, and basic operations like creating new pages or finding existing ones by page number. This layer knows nothing about SQL or transactions. It simply moves pages between disk and memory according to requests from above.

The buffer pool manager sits on top of the storage manager and provides a crucial optimization: caching. Instead of reading every page directly from disk for every query, the buffer pool keeps frequently accessed pages resident in memory. When a client requests a page, the buffer pool first checks its cache. If the page is present (a cache hit), it returns the in-memory copy immediately. If not (a cache miss), it reads from disk into an available buffer frame. The buffer pool must decide which pages to evict when memory fills up, track which pages have been modified and need flushing to disk, and coordinate access between concurrent readers and writers through pinning semantics

and latches. Buffer pool design is one of the most impactful areas for database performance because cache hit rates directly determine how many expensive disk I/O operations are needed.

The logging layer provides durability through write-ahead logging (WAL). The fundamental insight behind WAL is elegant: instead of writing every modification directly to data pages on disk, we first append a log record describing the change to an append-only log file and flush that log to stable storage. Only after the log is safely on disk do we consider the change durable. The actual data page can be updated in memory and flushed to disk later at our convenience. This approach transforms random writes scattered across many pages into sequential appends to a single log file, which is dramatically more efficient. More importantly, it enables crash recovery: if the system crashes, we replay the log on restart to reconstruct all committed changes that were not yet written to data pages.

The indexing layer provides fast access to data without scanning every page. A B-tree index maintains sorted keys with pointers to data, enabling logarithmic-time lookups instead of linear scans. For a table with one million rows stored across 250,000 pages, a full table scan requires reading all 250,000 pages while a B-tree lookup typically touches only three or four pages. Indexes trade write overhead for read speed: every insert, update, or delete must also modify affected indexes. Beyond B-trees we will explore hash indexes for equality lookups and discuss alternative structures like LSM trees that excel under different workload patterns.

The transaction manager coordinates concurrent access and provides ACID semantics. Atomicity ensures that either all operations in a transaction complete or none do. Consistency means transactions preserve database integrity constraints. Isolation prevents concurrent transactions from interfering with each other in harmful ways. Durability guarantees that committed changes survive crashes. The transaction manager implements these properties through concurrency control mechanisms like locking protocols or multiversion concurrency control, and coordinates with the logging layer for recovery.

The SQL processing layers handle query language parsing, planning, optimization, and execution. The parser transforms text into an abstract syntax tree representing the query structure. Semantic analysis validates that tables and columns exist and types are compatible. The planner converts the validated AST into a logical plan of relational operations. The optimizer applies

rewrite rules and cost models to select an efficient physical execution strategy among many alternatives. Finally, the execution engine runs the chosen plan using operators that scan data, apply filters, perform joins, sort results, and compute aggregations.

At the top sits the client interface: either a command-line tool or network server that accepts SQL statements from users and returns results. This is the only part users interact with directly, but all the complexity we have described operates beneath it to deliver correct, efficient, and reliable behavior.

Implementation Roadmap: From Empty Project to Working Database

The book follows a deliberate progression from low-level storage up through query execution. Each chapter builds on previous work so that code accumulates into a functioning system rather than remaining as isolated examples. Here is the roadmap:

Chapters 2 through 5 establish the foundation: project setup, storage manager, file format with slotted pages for variable-length records, and buffer pool management. After Chapter 5 you will be able to store records on disk and retrieve them efficiently through a caching layer.

Chapters 6 through 8 add reliability and concurrency: write-ahead logging with crash recovery, B-tree indexes for fast lookups, and transaction management with locking and MVCC. After Chapter 8 your database can handle concurrent transactions with proper isolation while surviving crashes without data loss.

Chapters 9 through 11 bring SQL processing to life: parsing text into abstract syntax trees, planning and optimizing queries using cost models and statistics, and executing plans through a pipeline of relational operators. After Chapter 11 you will have end-to-end SQL support for SELECT, INSERT, UPDATE, and DELETE with WHERE clauses, JOINS, GROUP BY, ORDER BY, and more.

Chapters 12 through 14 focus on advanced topics: performance engineering including profiling and optimization techniques, reliability engineering covering crash consistency and corruption recovery, and distributed database concepts including replication and sharding. Chapter 15 concludes by reviewing the complete architecture we built and discussing how it compares to production systems and where you might extend it further.

This progression mirrors how real databases are architected: storage first, then durability and concurrency, then query processing on top. Each layer depends on abstractions provided by layers below, so building bottom-up ensures a solid foundation.

How to Read This Book and Use the Code

This book assumes you are comfortable programming in at least one language and have basic familiarity with data structures and algorithms. Prior database experience is helpful but not required: if you know how to write SQL queries, that is enough background. Some chapters touch on operating system concepts like file I/O, memory mapping, threading, and synchronization, so comfort with systems-level thinking will help.

Every chapter includes complete, runnable Rust code designed to integrate with previous chapters into a single project. The code lives in a consistent repository structure with clear module boundaries. You can follow along by cloning the companion repository and building incrementally as we progress, or you can read through and study the implementations conceptually. Either approach works because every code example is self-contained enough to understand while also fitting into the larger architecture.

Some chapters present simplified versions of algorithms for educational clarity before showing how production systems handle additional complexity. I will always call these simplifications out explicitly so you know where our implementation diverges from what you would find in PostgreSQL or SQLite. This approach keeps individual chapters focused while giving you accurate mental models of real systems.

The book uses inline citations [n] to reference sources for factual claims, historical information, and externally derived material. A consolidated bibliography appears at the end with full references and URLs. You are encouraged to follow these links to primary sources when you want deeper understanding on any topic.

Now let us begin building. In the next chapter we will set up our Rust project, discuss why Rust is a strong choice for database implementation, and write our first code: initializing an empty database file with a proper header structure.

Chapter 2: Project Setup and Systems Foundations

Choosing Rust for Database Implementation

The language you choose for a database engine shapes everything: how you manage memory, how you handle errors, how you reason about concurrency, and ultimately how reliable the resulting system is. For mydb we will use Rust. This is a deliberate choice with trade-offs that deserve explicit discussion.

Rust provides memory safety without garbage collection through its ownership system and borrow checker. Database engines are among the most demanding consumers of memory in any software stack. They manage large buffer pools, maintain complex data structures across threads, and operate for weeks or months without restart. Memory bugs in this context are catastrophic: a use-after-free in a buffer pool can corrupt indexes silently, a double-free can crash the entire process, and dangling pointers can introduce security vulnerabilities that persist undetected for years. Rust's compile-time guarantees eliminate entire classes of these bugs before they reach production.

The absence of garbage collection is equally important. Garbage collectors introduce unpredictable pauses that are unacceptable in latency-sensitive database operations. When a user commits a transaction, they expect a deterministic response time, not one that depends on whether the GC happens to be running a major collection cycle. Rust's manual memory management through ownership gives you control over allocation lifetimes without the safety hazards of C or C++.

Rust also provides excellent support for systems-level programming. You have direct access to file I/O, memory mapping, threading primitives, and atomic operations. The standard library includes robust implementations of common data structures like hash maps and binary heaps. The type system catches many bugs at compile time that would require extensive testing in other languages. Error handling through Result types forces you to consider failure modes explicitly rather than relying on exceptions or silent failures.

The ecosystem is growing rapidly for database work specifically. DuckDB, one of the fastest analytical databases available, is written primarily in Rust and C++. PolarDB-X from Alibaba uses Rust for its distributed transaction coordinator. Many new storage engines and data processing systems choose Rust as their implementation language. The `sqlparser-rs` crate provides a mature SQL parser used by multiple projects. Crates like `tempfile`, `serde`, and `criterion` make testing, serialization, and benchmarking straightforward.

C++ remains the dominant language in the database world: PostgreSQL, MySQL/InnoDB, MongoDB, Elasticsearch, and many others are C++ codebases. C++ offers comparable performance with a richer history of database-specific libraries and patterns. However, C++ also gives you far more ways to shoot yourself in the foot: undefined behavior from uninitialized memory, data races from unsynchronized access, subtle lifetime bugs from raw pointer misuse. Modern C++ with smart pointers and RAII mitigates many issues but does not eliminate them. For an educational project where correctness matters as much as performance, Rust's stricter guarantees are worth the learning curve.

Go is another option that some database projects use, notably CockroachDB and etcd. Go simplifies concurrency with goroutines and channels, has fast compilation, and produces reliable code. However, its garbage collector introduces pauses that make it less suitable for latency-critical storage operations, and its lack of generics until recently (and still limited compared to Rust) makes implementing generic data structures more cumbersome. For a book focused on internals where memory behavior matters, Go's abstractions hide too much.

The trade-offs are real. Rust has a steeper learning curve than C or Go. The borrow checker can be frustrating when you first encounter it, especially for patterns involving shared mutable state that databases rely heavily on. Compilation times are slower than many alternatives. Some advanced database techniques like lock-free data structures require unsafe Rust blocks, which reintroduce the possibility of memory unsafety (though confined to auditable boundaries). We will address all these challenges as we encounter them and show how idiomatic Rust patterns handle common database engineering problems.

Project Structure and Build Configuration

Let us create our project structure. We will organize `mydb` as a Cargo workspace with clear module boundaries that mirror the architectural layers

described in Chapter 1. This structure will evolve as we add subsystems, but establishing it now prevents refactoring pain later.

Create a new Rust project:

```
1 cargo new mydb --lib
2 cd mydb
```

The initial Cargo.toml looks like this:

```
1 [package]
2 name = "mydb"
3 version = "0.1.0"
4 edition = "2021"
5
6 [dependencies]
7 thiserror = "1.0"
8 anyhow = "1.0"
9 log = "0.4"
10 env_logger = "0.10"
11
12 [dev-dependencies]
13 tempfile = "3.8"
14 criterion = "0.5"
15
16 [[bench]]
17 name = "storage_bench"
18 harness = false
```

We use `thiserror` for typed error enums, `anyhow` for ergonomic error handling in application code, `log` and `env_logger` for structured logging throughout the engine, `tempfile` for tests that need temporary database files, and `criterion` for benchmarking. As we add features we will include more dependencies: a tokenizer crate for SQL lexing, maybe an LRU cache implementation for the buffer pool, and serialization utilities.

Now let us set up the directory structure in `src/`:


```

1  src/
2  ├── lib.rs           # Library root, public API exports
3  ├── main.rs         # CLI entry point (in examples/ or as binary)
4  ├── storage/        # Storage layer: files, pages, I/O
5  │   ├── mod.rs
6  │   ├── file_manager.rs # Database file operations
7  │   └── page.rs      # Page abstraction and layout
8  ├── buffer_pool/    # Buffer pool manager
9  │   ├── mod.rs
10  │   └── replacer.rs  # Page replacement algorithms
11  ├── log/            # Write-ahead logging
12  │   ├── mod.rs
13  │   └── record.rs    # Log record format
14  ├── index/          # Index implementations
15  │   ├── mod.rs
16  │   ├── btree.rs     # B+ tree index
17  │   └── hash_index.rs # Hash index
18  ├── transaction/    # Transaction management
19  │   ├── mod.rs
20  │   ├── lock_manager.rs # Locking protocols
21  │   └── mvcc.rs      # Multiversion concurrency control
22  ├── types/          # Data types and serialization
23  │   ├── mod.rs
24  │   ├── value.rs     # Cell values (INTEGER, TEXT, etc.)
25  │   └── schema.rs    # Table schemas
26  ├── sql/            # SQL processing
27  │   ├── mod.rs
28  │   ├── lexer.rs     # Tokenizer
29  │   ├── parser.rs    # Recursive descent parser
30  │   └── ast.rs       # Abstract syntax tree
31  ├── planner/        # Query planning and optimization
32  │   ├── mod.rs
33  │   ├── logical_plan.rs # Logical operators
34  │   ├── physical_plan.rs # Physical execution strategies
35  │   └── optimizer.rs  # Cost-based optimization
36  ├── executor/       # Query execution engine
37  │   ├── mod.rs
38  │   ├── operator.rs   # Operator trait and base types
39  │   ├── scan.rs       # Table and index scans
40  │   ├── join.rs       # Join implementations
41  │   └── aggregate.rs  # Aggregation operators
42  └── engine.rs       # Top-level database engine API

```

This structure is intentionally layered. The storage module knows nothing about SQL or transactions. The buffer pool module depends only on storage for page I/O. The transaction manager depends on buffer pool and logging. The SQL modules depend on types and schema information but not directly on storage internals. These boundaries prevent circular dependencies and make

each component independently testable.

Create the examples directory for our CLI tool:

```
1 mkdir examples
2 touch examples/cli.rs
```

Add a binary target to Cargo.toml:

```
1 [[bin]]
2 name = "mydb-cli"
3 path = "examples/cli.rs"
```

Now let us write the library root. Create src/lib.rs:

```
1 /// mydb - A relational database engine built from scratch.
2 ///
3 /// This crate implements a complete database storage engine with support for
4 /// persistent storage, buffered I/O, write-ahead logging, B-tree indexing,
5 /// transactional concurrency control, and SQL query processing.
6
7 pub mod storage;
8 pub mod buffer_pool;
9 pub mod log;
10 pub mod index;
11 pub mod transaction;
12 pub mod types;
13 pub mod sql;
14 pub mod planner;
15 pub mod executor;
16 pub mod engine;
17
18 // Re-export commonly used types at the crate root for convenience.
19 pub use engine::DatabaseEngine;
20 pub use types::value::Value;
21 pub use types::schema::{Schema, ColumnType};
```

Each module starts empty and grows as we implement its functionality. This forward declaration ensures the project compiles throughout development even before every component is fully implemented. We will fill in each module incrementally across chapters.

Core Abstractions: Pages, Records, and Offsets

Before writing any storage code, let us define the core abstractions that will appear throughout the entire engine. These concepts are fundamental to how databases think about data organization.

A page is a fixed-size block of bytes that serves as the atomic unit of I/O between memory and disk. Our implementation uses 4096-byte pages, matching the typical filesystem block size and CPU cache line alignment. This choice is not arbitrary: reading or writing in multiples of the filesystem block size avoids partial-block operations that waste bandwidth. Every page has a unique page ID within a database file, starting from zero.

Why do databases think in pages rather than individual records or raw bytes? The answer lies in hardware characteristics. Disk I/O has enormous latency compared to memory access: spinning disks require mechanical seeking and rotation (milliseconds per operation), while even fast SSDs introduce tens of microseconds of overhead per read. Memory access takes nanoseconds. If you read one byte at a time from disk, the overhead dominates completely. By reading in fixed-size pages, you amortize the I/O cost across many bytes of useful data. Furthermore, operating systems and storage drivers are optimized for block-sized transfers, so aligning with page boundaries gives you better performance.

Within each page lives structured data organized according to the page type. Data pages store table records using a slotted format: a header describing the page state, an array of slot pointers indicating where each record begins, and the actual record payloads packed into available space. Index pages store keys and child pointers arranged for efficient binary search. Special pages hold metadata like database headers, free space maps, or transaction logs. Each page type defines its own internal layout but shares the same fundamental size and I/O semantics.

A record (also called a row or tuple) represents a single logical entity stored in a table. Records can have variable length because they contain fields of different sizes: a fixed-width integer column occupies four bytes while a text column might hold anywhere from zero to thousands of bytes depending on the actual value. Variable-length records create complexity for storage management because you cannot compute record positions by simple arithmetic. Instead, databases use directory structures within pages (slot

arrays) that track where each record starts and how much space it occupies.

Offsets are integer positions within a page or file that locate specific data. Page offsets measure bytes from the beginning of a page: offset 0 is the first byte, offset 4095 is the last byte of a 4096-byte page. File offsets measure bytes from the beginning of the database file and can be converted to page IDs by dividing by the page size. Understanding offsets is critical for implementing storage correctly because every read and write operation must specify exact positions.

Let us formalize these concepts in code. Create `src/storage/page.rs`:

```

1  /// Page abstraction for fixed-size blocks of data.
2  ///
3  /// Pages are the fundamental unit of I/O in mydb. Each page is PAGE_SIZE bytes,
4  /// identified by a unique PageId within a database file, and can hold
   ↳ structured
5  /// data depending on its type (data page, index page, metadata page, etc.).
6
7  use std::fmt;
8
9  /// Default page size: 4096 bytes, matching typical filesystem block size.
10 pub const PAGE_SIZE: usize = 4096;
11
12 /// Unique identifier for a page within a database file.
13 #[derive(Debug, Clone, Copy, PartialEq, Eq, PartialOrd, Ord, Hash)]
14 pub struct PageId(u64);
15
16 impl PageId {
17     /// Create a new PageId from a raw number.
18     pub fn new(id: u64) -> Self {
19         Self(id)
20     }
21
22     /// Get the raw numeric ID.
23     pub fn as_u64(&self) -> u64 {
24         self.0
25     }
26
27     /// Compute the byte offset of this page within a database file.
28     pub fn file_offset(&self) -> u64 {
29         self.0 * PAGE_SIZE as u64
30     }
31 }
32
33 impl fmt::Display for PageId {
34     fn fmt(&self, f: &mut fmt::Formatter<'_,_>) -> fmt::Result {
35         write!(f, "Page({})", self.0)

```

```

36     }
37 }
38
39 /// A page buffer holding PAGE_SIZE bytes of data in memory.
40 ///
41 /// This is the in-memory representation of a disk page. The buffer pool manager
42 /// allocates PageBuffer instances and fills them by reading from disk or
43 /// initializing with zeros for new pages. Higher-level components interpret
44 /// the bytes according to the appropriate page format (slotted page, B-tree
45 /// node, etc.).
46 #[derive(Clone)]
47 pub struct PageBuffer {
48     /// The page ID this buffer represents.
49     pub page_id: PageId,
50     /// Raw byte contents of the page.
51     pub data: [u8; PAGE_SIZE],
52 }
53
54 impl PageBuffer {
55     /// Create a new empty page buffer filled with zeros.
56     pub fn new(page_id: PageId) -> Self {
57         Self {
58             page_id,
59             data: [0u8; PAGE_SIZE],
60         }
61     }
62
63     /// Get a mutable slice of the page data.
64     pub fn as_mut_slice(&mut self) -> &mut [u8] {
65         &mut self.data
66     }
67
68     /// Get an immutable slice of the page data.
69     pub fn as_slice(&self) -> &[u8] {
70         &self.data
71     }
72 }
73
74 impl fmt::Debug for PageBuffer {
75     fn fmt(&self, f: &mut fmt::Formatter<'_>) -> fmt::Result {
76         f.debug_struct("PageBuffer")
77             .field("page_id", &self.page_id)
78             .field("data_len", &PAGE_SIZE)
79             .finish()
80     }
81 }
82
83 /// Types of pages that can exist in a database file.
84 #[derive(Debug, Clone, Copy, PartialEq, Eq)]
85 pub enum PageType {

```

```

86     /// Unknown or uninitialized page.
87     Unused,
88     /// Database header and metadata (page 0).
89     Header,
90     /// Data page storing table records using slotted format.
91     Data,
92     /// Internal B-tree node page (non-leaf index page).
93     IndexInternal,
94     /// Leaf B-tree node page (stores actual index entries).
95     IndexLeaf,
96 }
97
98 impl PageType {
99     /// Encode this page type as a single byte for storage in page headers.
100     pub fn as_byte(self) -> u8 {
101         match self {
102             Self::Unused => 0x00,
103             Self::Header => 0x01,
104             Self::Data => 0x10,
105             Self::IndexInternal => 0x20,
106             Self::IndexLeaf => 0x30,
107         }
108     }
109
110     /// Decode a page type from a stored byte value.
111     pub fn from_byte(byte: u8) -> Self {
112         match byte {
113             0x01 => Self::Header,
114             0x10 => Self::Data,
115             0x20 => Self::IndexInternal,
116             0x30 => Self::IndexLeaf,
117             _ => Self::Unused,
118         }
119     }
120 }

```

This code establishes the foundational types. `PageId` wraps a `u64` to prevent accidental misuse of raw integers throughout the codebase. `PageBuffer` holds exactly `PAGE_SIZE` bytes and tracks which page it represents. `PageType` enumerates the kinds of pages we support with byte-level encoding for on-disk storage. Every subsequent storage component will use these types consistently.

Create `src/storage/mod.rs`:

```

1  /// Storage layer: manages database files, pages, and disk I/O.
2  ///
3  /// This module provides the lowest level abstractions in mydb. It handles
4  /// opening and creating database files, reading and writing pages at specific
5  /// offsets, tracking file size, and managing free space for new pages.
6
7  pub mod file_manager;
8  pub mod page;
9
10 pub use page::{PageId, PageBuffer, PageType, PAGE_SIZE};
11 pub use file_manager::FileManager;

```

Error Handling Strategy Across the Engine

Database engines must handle errors carefully because failures can corrupt data or lose transactions if mishandled. Rust's type system gives us excellent tools for this: Result types force callers to consider error cases explicitly, and enum-based error types let us distinguish between recoverable and fatal conditions.

We will use a two-layer error strategy. At the library level, modules define specific error enums using thiserror that describe what can go wrong in that domain. At the application level (CLI tool, test harnesses), we use anyhow::Result for ergonomic error propagation without repeating boilerplate type annotations everywhere.

Create src/storage/file_manager.rs with our first concrete storage implementation:

```

1  /// Database file management: open, create, read, write pages to disk.
2  ///
3  /// The FileManager is responsible for all direct interaction with the database
4  /// file on disk. Higher-level components never open files themselves; they
5  /// request page reads and writes through this interface.
6
7  use std::fs::{File, OpenOptions};
8  use std::io::{Read, Seek, SeekFrom, Write};
9  use std::path::{Path, PathBuf};
10 use thiserror::Error;
11
12 use super::page::{PageBuffer, PageId, PageType, PAGE_SIZE};
13
14 /// Errors that can occur during file management operations.
15 #[derive(Error, Debug)]

```

```

16 pub enum FileManagerError {
17     /// A low-level I/O error occurred.
18     #[error("I/O error: {0}")]
19     Io(#[from] std::io::Error),
20
21     /// Attempted to access a page beyond the current file size.
22     #[error("Page {page_id} does not exist in database")]
23     PageNotFound { page_id: PageId },
24
25     /// Database file is corrupted or has invalid header.
26     #[error("Database file corrupted: {0}")]
27     Corrupted(String),
28
29     /// File path or permissions error.
30     #[error("File access error: {0}")]
31     FileAccess(String),
32 }
33
34 /// Manages the database file on disk.
35 ///
36 /// The FileManager maintains an open handle to the database file and provides
37 /// methods for reading and writing pages at specific offsets. It tracks the
38 /// current file size in pages so that new pages can be allocated by extending
39 /// the file. All I/O operations go through this struct to ensure consistency
40 /// and enable future optimizations like memory-mapped access or async I/O.
41 pub struct FileManager {
42     /// The underlying file handle.
43     file: File,
44     /// Path to the database file for error reporting.
45     path: PathBuf,
46     /// Current number of pages in the database file.
47     num_pages: u64,
48 }
49
50 impl FileManager {
51     /// Open an existing database file or create a new one if it does not exist.
52     ///
53     /// If creating a new file, initializes page 0 with a valid database header.
54     pub fn open_or_create<P: AsRef<Path>>>(path: P) -> Result<Self,
55     ↪ FileManagerError> {
56         let path = path.as_ref().to_path_buf();
57         let exists = path.exists();
58
59         let file = OpenOptions::new()
60             .read(true)
61             .write(true)
62             .create(true)
63             .truncate(false)
64             .open(&path)?;

```



```

65     let metadata = file.metadata()?;
66     let file_size = metadata.len();
67
68     // Calculate number of complete pages in the file.
69     let num_pages = if file_size == 0 {
70         0
71     } else {
72         (file_size + PAGE_SIZE as u64 - 1) / PAGE_SIZE as u64
73     };
74
75     let mut manager = Self {
76         file,
77         path,
78         num_pages,
79     };
80
81     if !exists || file_size == 0 {
82         // Initialize a new database with a header page.
83         manager.init_header()?;
84     } else {
85         // Validate existing header.
86         manager.validate_header()?;
87     }
88
89     Ok(manager)
90 }
91
92 /// Initialize page 0 with the database header structure.
93 fn init_header(&mut self) -> Result<(), FileManagerError> {
94     let mut buffer = PageBuffer::new(PageId::new(0));
95
96     // Write magic number: "MYDB" followed by version byte and format flags.
97     // This allows us to identify valid database files and detect corruption.
98     buffer.data[0..4].copy_from_slice(b"MYDB");
99     buffer.data[4] = 1; // Format version
100    buffer.data[5] = PageType::Header.as_byte();
101
102    // Store page size as big-endian u32 at offset 8.
103    let page_size_bytes = (PAGE_SIZE as u32).to_be_bytes();
104    buffer.data[8..12].copy_from_slice(&page_size_bytes);
105
106    // Initialize page count to 1 (just the header page).
107    let page_count_bytes: [u8; 8] = 1u64.to_be_bytes();
108    buffer.data[12..20].copy_from_slice(&page_count_bytes);
109
110    // Write header page to disk.
111    self.write_page(PageId::new(0), &buffer)?;
112    self.num_pages = 1;
113
114    Ok(())

```

```

115     }
116
117     /// Validate that the database file has a correct header.
118     fn validate_header(&mut self) -> Result<(), FileManagerError> {
119         let buffer = self.read_page(PageId::new(0))?;
120
121         // Check magic number.
122         if &buffer.data[0..4] != b"MYDB" {
123             return Err(FileManagerError::Corrupted(
124                 "Invalid magic number: not a mydb database file".to_string(),
125             ));
126         }
127
128         // Check version.
129         if buffer.data[4] != 1 {
130             return Err(FileManagerError::Corrupted(format!(
131                 "Unsupported format version: {}",
132                 buffer.data[4]
133             )));
134         }
135
136         // Verify page size matches our expectation.
137         let stored_page_size =
138     ↪ u32::from_be_bytes(buffer.data[8..12].try_into().unwrap());
139         if stored_page_size as usize != PAGE_SIZE {
140             return Err(FileManagerError::Corrupted(format!(
141                 "Page size mismatch: file uses {} bytes, expected {}",
142                 stored_page_size, PAGE_SIZE
143             )));
144         }
145
146         // Read current page count from header.
147         let stored_count =
148     ↪ u64::from_be_bytes(buffer.data[12..20].try_into().unwrap());
149         self.num_pages = stored_count.max(self.num_pages);
150
151         Ok(())
152     }
153
154     /// Read a page from disk into a PageBuffer.
155     /// Returns an error if the page ID is beyond the current file size.
156     pub fn read_page(&mut self, page_id: PageId) -> Result<PageBuffer,
157     ↪ FileManagerError> {
158         let offset = page_id.file_offset();
159         let max_offset = self.num_pages * PAGE_SIZE as u64;
160
161         if offset + PAGE_SIZE as u64 > max_offset {
162             return Err(FileManagerError::PageNotFound { page_id });
163         }
164     }

```

```

162
163     let mut buffer = PageBuffer::new(page_id);
164     self.file.seek(SeekFrom::Start(offset))?;
165     self.file.read_exact(&mut buffer.data)?;
166
167     Ok(buffer)
168 }
169
170 /// Write a page buffer to disk at its designated offset.
171 pub fn write_page(&mut self, page_id: PageId, buffer: &PageBuffer) ->
172 ↪ Result<(), FileManagerError> {
173     let offset = page_id.file_offset();
174     self.file.seek(SeekFrom::Start(offset))?;
175     self.file.write_all(&buffer.data)?;
176
177     // Track file growth if we wrote a new page at the end.
178     let new_page_count = page_id.as_u64() + 1;
179     if new_page_count > self.num_pages {
180         self.num_pages = new_page_count;
181         self.update_page_count_in_header()?;
182     }
183
184     Ok(())
185 }
186
187 /// Update the page count stored in the header.
188 fn update_page_count_in_header(&mut self) -> Result<(), FileManagerError> {
189     let mut buffer = self.read_page(PageId::new(0))?;
190     let count_bytes: [u8; 8] = self.num_pages.to_be_bytes();
191     buffer.data[12..20].copy_from_slice(&count_bytes);
192     self.write_page(PageId::new(0), &buffer)?;
193     Ok(())
194 }
195
196 /// Allocate a new page at the end of the file.
197 /// Returns the PageId of the newly allocated page and a zero-filled buffer.
198 pub fn allocate_page(&mut self) -> Result<(PageId, PageBuffer),
199 ↪ FileManagerError> {
200     let page_id = PageId::new(self.num_pages);
201     let buffer = PageBuffer::new(page_id);
202
203     // Extend file by writing the new zero-filled page.
204     self.write_page(page_id, &buffer)?;
205
206     Ok((page_id, buffer))
207 }
208
209 /// Get the current number of pages in the database.
210 pub fn num_pages(&self) -> u64 {

```

```

210         self.num_pages
211     }
212
213     /// Flush all buffered data to disk using fsync.
214     ///
215     /// This ensures that all writes are durable on stable storage, surviving
216     /// a system crash or power loss. Called at transaction commit and during
217     /// checkpoint operations.
218     pub fn sync(&self) -> Result<(), FileManagerError> {
219         self.file.sync_all()?;
220         Ok(())
221     }
222
223     /// Get the file path for this database.
224     pub fn path(&self) -> &Path {
225         &self.path
226     }
227 }

```

Let us walk through this implementation carefully because every design decision here matters.

The `FileManager` struct holds a single `File` handle that remains open for the lifetime of the database connection. Opening and closing files repeatedly would be prohibitively expensive. The `path` field stores the file location for error messages and future operations like WAL file naming. The `num_pages` field tracks how many pages exist in the file so we can allocate new pages sequentially without scanning or seeking.

The `open_or_create` method handles both cases: opening an existing database validates its header to detect corruption, while creating a new one initializes a proper header structure. This dual behavior mirrors real databases like SQLite that create files transparently when connecting to non-existent paths.

The header format encodes several critical pieces of information. The magic number “MYDB” identifies valid database files and immediately rejects corrupted or incompatible files. The version byte allows future format evolution: if we change the file layout, we bump the version and refuse to open older formats without migration. The page size field enables cross-platform compatibility: different databases could theoretically use different page sizes, and we verify that our expected size matches what the file declares. The page count lets us quickly determine where new pages should be allocated.

All multi-byte integers in the header use big-endian byte order. This is a

deliberate choice for portability: the same database file can be read correctly on both little-endian machines (most Intel and AMD CPUs) and big-endian systems without byte-swapping. SQLite makes the same choice for exactly this reason [4].

The `read_page` method seeks to the correct offset within the file and reads exactly `PAGE_SIZE` bytes using `read_exact`, which fails if fewer bytes are returned (indicating truncation or corruption). The `write_page` method performs the reverse operation and updates the page count if we extended the file. Both methods operate on exact byte offsets computed from `PageId.file_offset()`, ensuring no ambiguity about where data lives.

The `sync` method calls `fsync` through Rust's `sync_all`, which flushes both kernel page cache buffers and filesystem metadata to physical storage. This is essential for durability guarantees: without `fsync`, writes might appear successful in user space while still sitting in volatile RAM that would be lost on power failure. We will call this at critical points during transaction commits.

First Binary: Initializing an Empty Database File

Now let us write our first executable program: a simple CLI tool that creates and inspects database files. This gives us something runnable to test our storage layer immediately.

Create `examples/cli.rs`:

```

1  /// Command-line interface for mydb.
2  ///
3  /// Usage:
4  ///   mydb-cli init <path>           Create a new database file
5  ///   mydb-cli info <path>          Show database file information
6  ///   mydb-cli shell <path>         Start interactive SQL shell (future)
7
8  use std::env;
9  use std::process;
10
11 use mydb::storage::{FileManager, PAGE_SIZE};
12
13 fn main() {
14     env_logger::init();
15
16     let args: Vec<String> = env::args().collect();
17 
```

```

18     if args.len() < 2 {
19         eprintln!("Usage: mydb-cli <command> [arguments]");
20         eprintln!("Commands: init, info");
21         process::exit(1);
22     }
23
24     let command = &args[1];
25
26     match command.as_str() {
27         "init" => cmd_init(&args),
28         "info" => cmd_info(&args),
29         _ => {
30             eprintln!("Unknown command: {}", command);
31             process::exit(1);
32         }
33     }
34 }
35
36 /// Initialize a new database file at the given path.
37 fn cmd_init(args: &[String]) {
38     if args.len() < 3 {
39         eprintln!("Usage: mydb-cli init <path>");
40         process::exit(1);
41     }
42
43     let path = &args[2];
44
45     match FileManager::open_or_create(path) {
46         Ok(mut fm) => {
47             // Sync to ensure header is durable.
48             if let Err(e) = fm.sync() {
49                 eprintln!("Error syncing database: {}", e);
50                 process::exit(1);
51             }
52
53             println!("Database initialized at {}", path);
54             println!("Page size: {} bytes", PAGE_SIZE);
55             println!("Pages allocated: {}", fm.num_pages());
56         }
57         Err(e) => {
58             eprintln!("Error initializing database: {}", e);
59             process::exit(1);
60         }
61     }
62 }
63
64 /// Display information about an existing database file.
65 fn cmd_info(args: &[String]) {
66     if args.len() < 3 {
67         eprintln!("Usage: mydb-cli info <path>");

```

```

68         process::exit(1);
69     }
70
71     let path = &args[2];
72
73     match FileManager::open_or_create(path) {
74         Ok(fm) => {
75             println!("Database file: {}", path);
76             println!("Page size: {} bytes", PAGE_SIZE);
77             println!("Pages: {}", fm.num_pages());
78             println!(
79                 "File size: {} bytes {:.2} KB",
80                 fm.num_pages() * PAGE_SIZE as u64,
81                 (fm.num_pages() as f64 * PAGE_SIZE as f64) / 1024.0
82             );
83         }
84         Err(e) => {
85             eprintln!("Error reading database: {}", e);
86             process::exit(1);
87         }
88     }
89 }

```

Build and run the project:

```

1 cargo build --release
2 ./target/release/mydb-cli init test.db
3 ./target/release/mydb-cli info test.db

```

Expected output:

```

1 Database initialized at test.db
2 Page size: 4096 bytes
3 Pages allocated: 1
4
5 Database file: test.db
6 Page size: 4096 bytes
7 Pages: 1
8 File size: 4096 bytes (4.00 KB)

```

You now have a working database file with a valid header structure. If you inspect it with a hex editor, you will see the “MYDB” magic number at offset 0, followed by version and metadata fields exactly as specified in our format. This small program already demonstrates core database behaviors: file creation, header initialization, validation on open, and durability through fsync.

In the next chapter we will dive deeper into storage fundamentals: how disks and SSDs actually work, why page-based I/O matters for performance, binary serialization details including endianness and alignment, and building a complete storage manager interface that higher layers can rely on.

Chapter 3: Storage Fundamentals — Disks, SSDs, and Pages

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

How Disks and SSDs Actually Work

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Why Databases Think in Pages, Not Bytes

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Building Our Page Abstraction Layer

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Binary Data Layout: Endianness, Alignment, and Serialization

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Implementing the Storage Manager Interface

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 4: File Format and Slotted Page Records

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Designing the Database File Format

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Slotted Page Layout: Headers, Directories, and Records

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Variable-Length Record Serialization

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Free-Space Management Within Pages

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Implementing Page Reads and Writes with Verification

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 5: Buffer Pool Manager — Caching at Scale

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

The Purpose of Buffer Pools and Why They Are Hard

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Pinning Semantics and Reference Counting

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Dirty Pages, Write Policies, and Flush Behavior

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Handling Memory Pressure and Eviction Under Load

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 6: Logging, Durability, and Crash Recovery

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

The Durability Problem and Why fsync Matters

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Write-Ahead Logging: Protocol and Invariants

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Designing Our Log Record Format

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Checkpointing Strategies and Implementation

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Crash Recovery: Redo and Undo in Practice

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 7: Indexing — B-Trees, Hash Indexes, and Beyond

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Why Indexes Exist and What They Trade Away

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Implementing a B+ Tree from Scratch

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Page Splits, Merges, and Concurrency Considerations

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Hash Indexes for Equality Lookups

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

LSM Trees and Alternative Index Architectures

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 8: Transactions and Concurrency Control

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

ACID Properties: What They Mean in Implementation Terms

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Locking Protocols and Two-Phase Locking

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Deadlock Detection and Prevention

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Latches vs Locks: Protecting Data Structures

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Implementing MVCC with Snapshot Isolation

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 9: SQL Parsing and Semantic Analysis

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Designing Our SQL Subset and Grammar

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Lexical Analysis: Tokenizing SQL Input

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Recursive Descent Parsing and AST Construction

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Symbol Tables and Schema Resolution

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 10: Query Planning and Cost-Based Optimization

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

From AST to Logical Plan

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Physical Operators and Execution Strategies

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Building a Cost Model from Statistics

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Cardinality Estimation: Histograms and Selectivity

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Join Ordering and Query Rewrite Rules

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 11: Query Execution Engine

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Operator Architecture: Iterators and Pipelining

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Table Scans and Index Scans

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Filtering, Projection, and Expression Evaluation

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Join Implementations: Nested-Loop, Hash, and Merge

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Aggregation and Grouping with Hash Tables

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 12: Performance Engineering and Optimization

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Understanding Database Performance Characteristics

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

I/O Patterns: Sequential Versus Random Access

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Cache Efficiency and Data Locality

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Vectorized Execution vs Tuple-at-a-Time Processing

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Profiling, Benchmarking, and Workload Generation

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Systematic Optimization: Finding and Fixing Bottlenecks

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 13: Reliability, Correctness, and Failure Handling

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Crash Consistency: Guarantees and Invariants

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Torn Pages and Partial Writes

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Corruption Detection with Checksums and Validation

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Disk-Full Conditions and Safe Failure Modes

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Backup, Restore, and Integrity Verification

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 14: Beyond Single Node — Distributed Database Concepts

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Why Single-Node Databases Eventually Hit Limits

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Replication Architectures: Leader-Follower and Multi-Leader

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Log Replication and Consistency Models

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Sharding and Distributed Query Execution

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Consensus, Fault Tolerance, and Rebalancing

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Chapter 15: Conclusion — From Educational Engine to Production System

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Architecture Review: How All Subsystems Fit Together

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

What We Simplified and Why Production Systems Differ

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Comparing Our Implementation to Real Databases

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

Paths Forward: Extending Your Database Engine

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

The Engineer's Mindset for Building Reliable Systems

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.

References

This content is not available in the sample book. The book can be purchased on Leanpub at <https://leanpub.com/buildingdatabaseenginesfromscratch>.