

Operating Systems Project: Topic 11

Asynchronous I/O Benchmark and Enhancement (io_uring)

Liangsen Wang

224040364@link.cuhk.edu.cn

2026.1.29

Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O
- 3 Theory: Advanced Mechanisms
- 4 Implementation Guide
- 5 Evaluation Strategy
- 6 Conclusion

Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O
- 3 Theory: Advanced Mechanisms
- 4 Implementation Guide
- 5 Evaluation Strategy
- 6 Conclusion

The Mission: Topic 11 Requirements

Objective: Compare the performance of the modern `io_uring` interface against traditional Synchronous I/O.

Requirement 1: Benchmark (User Space)

- **Task:** Write a C program to perform file I/O using both read/write and `io_uring`.
- **Metrics:** Measure **Latency**, **Throughput (IOPS)**, and **Syscall Counts**.
- **Goal:** Prove that `io_uring` is faster for high-concurrency workloads.

The Question:
Can we do I/O
without paying for
System Calls?

Requirement 2: Kernel Enhancement (Advanced)

- **Target:** `fs/io_uring.c`.
- **Task:** Modify submission logic or optimize queue contention.

Advanced Options

Option A: New Submission Strategy

- **Current:** `io_uring` processes requests FIFO.
- **Idea:** Implement "Priority Submission". Add a flag to urgent requests so the kernel processes them before others in the batch.

Option B: Lock Contention Analysis

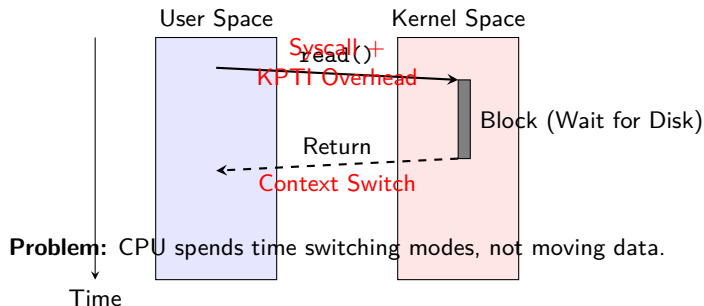
- **Problem:** Multiple threads sharing one Ring Buffer fight for the SQ lock.
- **Task:** Analyze lock contention using `perf` and propose optimizations (e.g., per-CPU rings).

Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O**
- 3 Theory: Advanced Mechanisms
- 4 Implementation Guide
- 5 Evaluation Strategy
- 6 Conclusion

Theory 1: The Cost of Synchronous I/O

Why is `read()` slow? Because of **Context Switches** and **Spectre/Meltdown Mitigations**.



Theory 2: Why not Legacy AIO?

Linux already had `libaio` (`io_submit`). Why create `io_uring`?

Legacy AIO Problems:

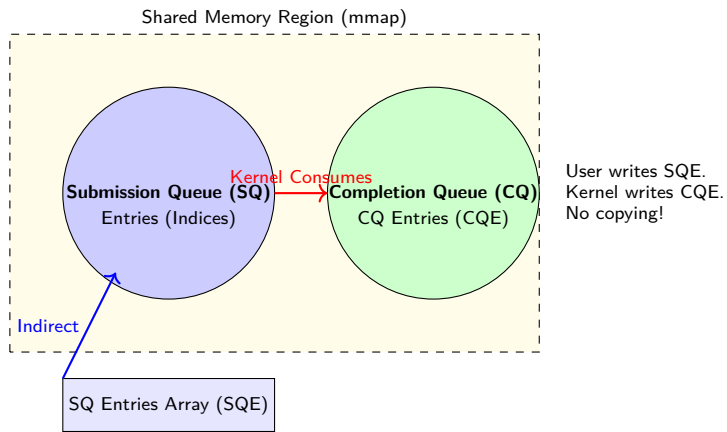
- **Direct I/O Only:** Only worked with `O_DIRECT` (bypassing cache). Buffered I/O would still block!
- **Complex API:** Hard to use correctly.
- **Data Copy:** Still required copying data structures for every submission.

io_uring Solutions:

- **Universal:** Works for Buffered I/O, Network Sockets, etc.
- **Zero Copy:** Shared memory ring buffers.
- **Extensible:** Designed to support any syscall asynchronously.

Theory 3: The io_uring Architecture

Core Concept: Two circular ring buffers mapped into user space.



Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O
- 3 Theory: Advanced Mechanisms**
- 4 Implementation Guide
- 5 Evaluation Strategy
- 6 Conclusion

Advanced 1: Lockless Coordination

How do User and Kernel share a ring without locks? **Memory Barriers.**

The Producer-Consumer Contract

- **User (Producer):** Writes SQE $\rightarrow$ `smp_store_release(tail)`
- **Kernel (Consumer):** `smp_load_acquire(tail)` $\rightarrow$ Reads SQE

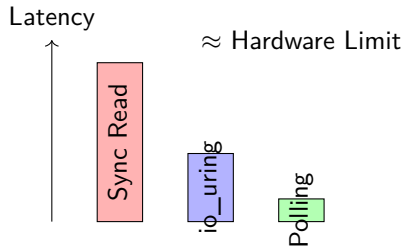
Why this matters: Explicit locking ('mutex') puts threads to sleep. Memory barriers just ensure ordering, allowing simpler wait-free progress in the fast path.

Advanced 2: IOPOLL (Zero Syscalls)

Standard `io_uring` still needs 1 syscall (`io_uring_enter`) to wake the kernel. **Polling Mode** (`IORING_SETUP_IOPOLL`) eliminates even that.

Mechanism:

- Kernel creates a specific kernel thread (`io-wq`).
- This thread **spin-loops**, checking the SQ ring for new entries.
- User simply writes to memory. Kernel picks it up instantly.



Trade-off: High CPU usage (100% on one core) for ultra-low latency.

Advanced 3: Registered Buffers & Files

Problem: For every I/O, the kernel must: 1. Map user virtual address to physical pages (`get_user_pages`). 2. Lookup file descriptor references (`fget`).

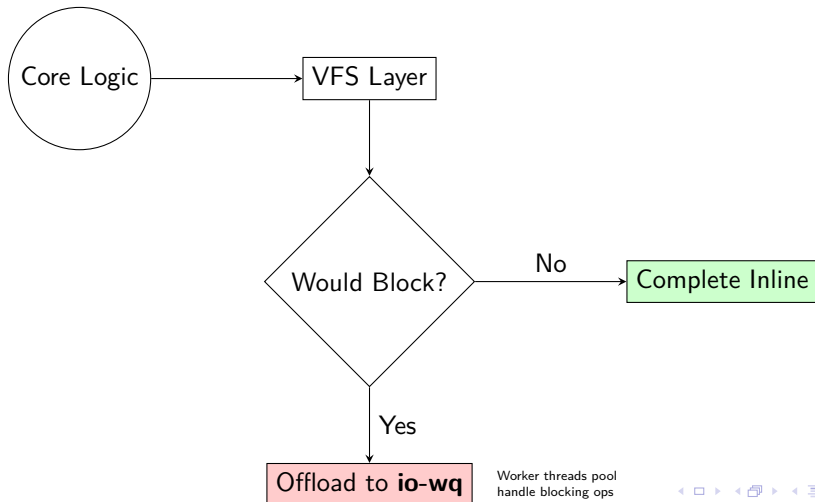
Solution: Pre-Registration

- User maps buffers **once** at startup.
- Kernel pins pages in RAM permanently.
- During I/O, skip mapping and reference counting.

Result: Another 10-20% performance gain.

Advanced 4: What if it blocks?

If an operation cannot be non-blocking (e.g., buffered read not in cache), `io_uring` cannot just block the whole ring.



Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O
- 3 Theory: Advanced Mechanisms
- 4 Implementation Guide**
- 5 Evaluation Strategy
- 6 Conclusion

Implementation 1: liburing Setup

Don't use raw syscalls. Use **liburing**.

```
1 #include <liburing.h>
2
3 #define QUEUE_DEPTH 4096
4 struct io_uring ring;
5
6 // Setup
7 int ret = io_uring_queue_init(QUEUE_DEPTH, &ring, 0);
8 if (ret < 0) {
9     fprintf(stderr, "queue_init failed: %d\n", ret);
10    return 1;
11 }
12
```

Implementation 2: The Submission Loop

```
1 struct io_uring_sqe *sqe;
2 struct iovec iov;
3
4 // 1. Get an entry from the ring
5 sqe = io_uring_get_sqe(&ring);
6
7 // 2. Prepare the operation (PREP macro)
8 io_uring_prep_readv(sqe, fd, &iov, 1, offset);
9
10 // Optional: Set User Data (to identify this request later)
11 io_uring_sqe_set_data(sqe, my_data_ptr);
12
13 // 3. Submit (Batched!)
14 // This only does the syscall if the ring is full or you force it
15 io_uring_submit(&ring);
16
```

Implementation 3: The Completion Loop

```
1 struct io_uring_cqe *cqe;
2 unsigned head;
3 int count = 0;
4
5 // Efficiently iterate over completed events
6 io_uring_for_each_cqe(&ring, head, cqe) {
7     count++;
8
9     // Check result
10    if (cqe->res < 0) {
11        // Handle Error
12    }
13
14    // Process data...
15 }
16
17 // Mark events as seen (Advance the Ring)
18 io_uring_cq_advance(&ring, count);
19
```

Kernel 1: The Battlefield (fs/io_uring.c)

For the advanced requirement, you must edit the kernel.

Key Functions

- `io_uring_enter()`: The system call entry point.
- `io_submit_sqes()`: Iterates the user's SQ ring.
- `io_issue_sqe()`: Attempts to issue a single request.

Task: Insert a `printk` or logic inside `io_submit_sqes` to count how many requests are processed in one batch.

Kernel 2: Implementing Priority

```
1 // In fs/io_uring.c :: io_submit_sqes(...)
2
3 // Loop through SQEs
4 io_for_each_link(req, head) {
5
6     // YOUR MODIFICATION: Check a flag
7     // Assuming you added a flag IOSQE_URGENT to include/uapi/linux/io_uring.h
8
9     if (req->flags & IOSQE_URGENT) {
10         // Handle High Priority: Maybe dispatch to a special workqueue?
11         trace_printk("High Prio Req detected!\n");
12     }
13
14     io_submit_sqe(req, ...);
15 }
16
```

Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O
- 3 Theory: Advanced Mechanisms
- 4 Implementation Guide
- 5 Evaluation Strategy**
- 6 Conclusion

Evaluation 1: Designing the Test

You need to stress the system to see the difference.

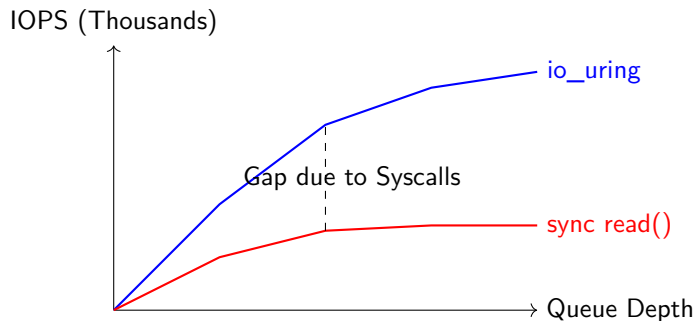
Test A: Throughput (IOPS)

- 4KB Random Reads.
- Queue Depth = 128.
- **Target:** 500k+ IOPS on NVMe.

Test B: Latency

- Queue Depth = 1.
- Measure time per operation.
- **Target:** Compare Polling mode vs. Interrupt mode.

Evaluation 2: What success looks like



If the lines overlap, you are likely limited by the Disk, not the Software. Use a RAM-disk or faster NVMe for testing.

Outline

- 1 Project Goals & Requirements
- 2 Theory: Evolution of Linux I/O
- 3 Theory: Advanced Mechanisms
- 4 Implementation Guide
- 5 Evaluation Strategy
- 6 Conclusion**

Summary

- ① **Problem:** Synchronous I/O wastes CPU cycles on system calls and context switches.
- ② **Solution:** `io_uring` uses shared memory rings to batch submission and completion.
- ③ **Implementation:** Use `liburing` for the user-space benchmark. Modify `fs/io_uring.c` for kernel enhancements.
- ④ **Result:** Order-of-magnitude improvement in IOPS for high-speed devices.

Roadmap & Resources

Action Plan:

- 1 **Read:** `man io_uring_enter`.
- 2 **Baseline:** Write the `read()` benchmark loop.
- 3 **Experiment:** Write the `io_uring` version. Enable Polling.
- 4 **Hack:** Compile kernel and inject counters into the submission loop.

Resources:

- *"Efficient IO with io_uring"* (Jens Axboe, Kernel Maintainer).
- `liburing` GitHub repository (Examples folder is gold).