

Enabling Faster NGS Analysis on Optane-based Heterogeneous Memory

Jiaolin Luo, Luanzheng Guo, Jie Ren, Kai Wu and Dong Li

University of California, Merced

Email: {jluo38,lguo4,jren6,kwu42,dli35}@ucmerced.edu

Abstract—Next-Generation Sequencing (NGS) analysis technologies are a pioneer approach for genome sequencing. The computation of NGS analysis exhibits a unique pattern, in which the execution requests a high density of small I/Os in the process *de novo* genome assembly. The small I/Os can have a huge impact on performance and delineate sequencing performance. To solve the problem caused by small I/Os, we leverage the byte-addressable feature of emerging persistent memory, which has largely transformed the computation architectures of HPC. We first conduct experiments to study the impact of persistent memory on NGS analysis. Furthermore, we propose an optimization mechanism, which converts POSIX read/write calls to pure memory LOAD/STORE instructions at runtime, to significantly advance I/O efficiency. Our evaluation demonstrates the effectiveness of the devised optimization mechanism, in which we achieve a performance improvement by 31%.

I. INTRODUCTION

NGS analysis, e.g., the state-of-the-art sequencing software Sentieon [1], drives scientific researches in multifarious disciplines, in particular for biology and clinical medicine researches. In the process *de novo* assembly of NGS analysis such as WGS and WES, the execution issues a large number of consecutive and small I/Os with the sizes of a few cache lines (as illustrated in Figure 1), in which the I/O performance poses a bottleneck of the execution of NGS analysis pipelines.

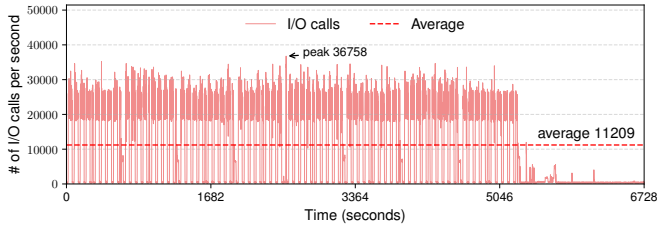


Fig. 1: The number of POSIX I/O calls in NGS analysis.

One fundamental design assumption of Linux I/O stack – consisting of page cache, vfs layer, file systems, block layer and block drivers – is that backing storage devices are block-addressable, in which the I/O granularity is known as sector and the data accesses are through the DMA mechanism transferring sectors to or from page cache. Such assumption is incorrect if backing devices are byte-addressable without the need of transferring data through DMA mechanism. In order to request DMA I/Os to or from devices and maintain file system metadata, existing Linux file I/O mechanisms (e.g., POSIX APIs *read* and *write*) involve in switching processor ring level from user context to kernel context, which can further incur extra overhead on critical I/O execution path. As a result, the design of existing I/O stack will obscure the

advantages of persistent memory and poses a bottleneck while employing persistent memory into HPC architectures.

Emerging persistent memory has profound impact on established HPC architectures and revolutionized the design of HPC software from multiple aspects. We seek to leverage the benefits of emerging persistent memory (e.g., Intel’s Optane DC) to optimize the efficiency of I/O stack in NGS analysis. First, persistent memory is byte-addressable and can transfer data without the involvement of DMA operations. Second, in terms of per GB price, the price of Intel’s Optane DC is merely 50% of the price of server level DRAM memory. Third, the maximum available capacity of persistent memory can easily reach up to tens of times of maximum DRAM capacity. Fourth, the memory access latency of persistent memory remains comparable to DRAM memory.

The benefits of persistent memory enable HPC applications to run with larger and cheaper extended memory. Nevertheless, the performance improvement of NGS analysis from simply configuring persistent memory as faster local storage is infinitesimal, in particular for the applications of their executions are I/O intensive (as illustrated in Figure 1). The overhead of the I/O stack in operating system hides the benefits of persistent memory.

We first conduct experiments to answer the following concerns in NGS analysis:

- Q1: Is I/O a bottleneck?
- Q2: What is the impact of persistent memory?
- Q3: Why I/O is slow on persistent memory?

In addition, we exploit the byte-addressable feature of persistent memory to shorten the length of I/O stack in NGS analysis through converting POSIX read/write calls to memory LOAD/STORE instructions at runtime.

II. CHARACTERIZATION

We conduct experiments – denoted as A, B and C – to answer the three concerns (Q1-Q3) from Section I, in which the experiment environment is illustrated in Table I.

Experiment A: Is I/O a bottleneck? We use Linux *perf* to collect the statistics of POSIX I/O calls in state-of-the-art NGS analysis software Sentieon.

TABLE I: Experiment Platform Specifications

Software	Ubuntu 18.04 with Linux kernel version 5.4.24
Sentieon	Version 201911
WGS fastq reads	Project RM8398 of NA12878
WGS reference genome	HG38
Processor	Intel® Xeon® Gold 6252N CPU @ 2.30GHz
DRAM	six 16-GB DDR4 DIMMs × 2 sockets (192 GB in total)
PM	six 128-GB Optane DC NVDIMMs × 2 sockets (1.5 TB in total)
Storage	Intel® D3-S4510 SSD 480GB

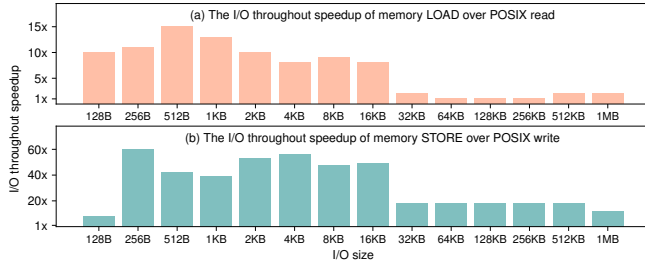


Fig. 2: The I/O throughput speedup of memory LOAD/STORE instructions over POSIX read/write calls.

Figure 1 shows the trends of the number of POSIX read/write calls across a sequencing execution. The dense I/O spikes range from 20000 to 30000 calls per second and approach to the maximum random IOPS of local SSD.

The high I/O numbers in Figure 1 and prior analysis demonstrate existing I/O stack poses a bottleneck in the execution of NGS analysis.

Experiment B: The impact of Optane on execution time. To mitigate the performance impact from I/O stack as shown in experiment A, will that make any difference if we simply store files on faster storage such as DRAM memory disk or Optane memory disk? The answer is negative.

We evaluate the Sentieon performance on DRAM-only and Optane-only respectively, where DRAM and Optane are configured as main memory accordingly. In addition, we also evaluate the Sentieon performance while sequencing data and intermediate files are on local SSD, DRAM memory disk and Optane memory disk respectively.

Table II shows the performance impact of persistent memory Intel’s Optane DC on the execution of NGS analysis. We have two findings: (1) The execution of Optane as main memory is 2.7 times slower than the execution of DRAM as main memory; (2) The execution time of Optane as local storage, DRAM as local storage and SSD as local storage has no patent performance difference.

One reasonable explanation might justify the findings is that the Linux I/O stack is inefficient as we have analyzed in Section I. We will further explore the causes in next experiment.

Experiment C: Why I/O is slow on persistent memory? We devise a micro benchmark to compare the I/O throughput speedups of pure memory LOAD/STORE instructions over traditional POSIX read/write calls with respect to various I/O sizes.

Figure 2 shows the I/O throughput speedups of pure memory LOAD/STORE instructions over traditional POSIX read/write. The I/O throughput speedups are evaluated with I/O sizes from 128B to 1MB. We have two observations: (1) For read, the I/O throughput speedup is up to 15 times; (2) For write, the I/O throughput speedup is up to 60 times.

TABLE II: The impact of Optane on the performance of Sentieon.

Main memory	Data storage	Average execution time (minutes)
DRAM	SSD	120
DRAM	DRAM memory disk	120
DRAM	Optane memory disk	120
Optane	SSD	320
Optane	DRAM memory disk	320
Optane	Optane memory disk	320

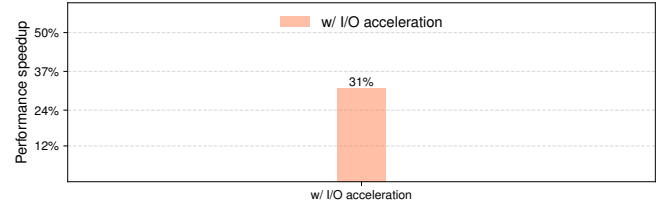


Fig. 3: The performance speedup with Optane based I/O acceleration, in which the baseline is the performance without Optane based I/O acceleration.

The results suggest the inefficient I/O stack can hide the benefits of persistent memory and justify the findings in Experiment B.

III. OPTIMIZATION

As we have known from Section I and Section II, the overhead of operating system I/O stack hides the benefits of persistent memory. To shorten the I/O stack length, we introduce an approach of converting POSIX read/write calls to pure memory LOAD/STORE instructions at runtime.

In the design, file-related calls (e.g., *open*, *close*, *lseek*, *read* and *write*) are re-interposed via the technique LD_PRELOAD and handled in our optimization library at runtime. When file I/O calls are from those files stored on persistent memory, the dynamic library will remap those files to the address space of processes and convert POSIX read/write calls to memory LOAD/STORE instructions.

IV. EVALUATION

The evaluation environment is shown in Table I. We compare the performance speedup of the execution time of Sentieon sequencing pipeline with I/O acceleration, in which the baseline is the execution time without I/O acceleration.

Figure 3 demonstrates the effectiveness of our Optane-based I/O acceleration technique with a performance speedup by 31%.

V. CONCLUSIONS

We first explore the impact of persistent memory (Intel’s Optane DC) on I/O intensive NGS analysis software (Sentieon) and understand its performance characterization via experiments. Our study reveals the inefficiency of Linux I/O stack can pose a bottleneck in NGS analysis.

In addition, we design an optimization method by shortening the length of Linux I/O stack via converting POSIX read/write calls to pure memory LOAD/STORE instructions. The evaluation demonstrates such optimization mechanism can make a huge performance improvement in NGS analysis.

ACKNOWLEDGMENTS

This work was partially supported by U.S. National Science Foundation (CNS-1617967, CCF-1553645 and CCF1718194). We thank Intel, Sentieon and reviewers for their constructive comments.

REFERENCES

- [1] K. I. Kendig, S. Baheti, M. A. Bockol, T. M. Drucker, S. N. Hart, J. R. Heldenbrand, M. Hernaez, M. E. Hudson, M. T. Kalmbach, E. W. Klee et al., “Sentieon dnaseq variant calling workflow demonstrates strong computational performance and accuracy,” *Frontiers in genetics*, vol. 10, 2019.