

DITRON: A FLEXIBLE AND VERSATILE DISTRIBUTED TENSOR COMPILER FOR LLM

Anonymous authors

Paper under double-blind review

ABSTRACT

As the performance scaling of individual devices slows down, distributed system acceleration has emerged as the mainstream. With the increasing maturity of parallel optimizations (e.g., tensor parallelism, sequence parallelism, expert parallelism) in recent years, researchers have found that the key to scalability lies in optimizing the overlap between computation and communication. Developing sophisticated overlapping kernels is challenging and exceeds the capabilities of most researchers, hindering the development of model architectures and parallel strategies. To address this issue, we present DITRON, a flexible and versatile distributed compiler that offers high-level programming interfaces for overlapping kernels. DITRON provides programming abstractions at three levels: (1) fine-grained tile-level programming within the scale-up domain; (2) chunk-level data transfer for the scale-out domain; (3) task-level distributed MegaKernel generation for entire LLM. DITRON inherits Triton’s programming model, enabling the transformation of original Triton kernels into parallel kernels with minimal modifications to the source code. Evaluation results show that overlapping kernels developed with DITRON are $1.27 \times -19.18 \times$ faster than non-overlapping versions, and even outperform expert-tuned CUDA libraries by 6%–30%. End-to-end inference yields a 5%–30% speedup over vLLM. Moreover, DITRON has been validated for training tasks on over 10,000 GPUs, proving robust capabilities for large-scale industrial deployment and saving millions of GPU hours monthly.

1 INTRODUCTION

The speed of large language model (LLM) training depends on the running speed of distributed programs; only with the support of high-speed training can models with diverse architectures be put into practical application (Dubey et al., 2024; Comanici et al., 2025; Rivière et al., 2024; DeepSeek-AI et al., 2024b; Qwen-Team, 2024; OpenAI, 2023; Anthropic, 2024). In the past, compiler work (such as Triton (Tillet et al., 2019) and TileLang (TileLang-Team, 2025)) and libraries (such as CUTLASS/CuTeDSL (Nvidia, 2022) and ThunkderKitten (Spector et al., 2024)) focus on finding ways to improve kernel performance on a single device, while communication in distributed clusters relied on separate libraries (NVIDIA, 2024) or compilers (Shah et al., 2025). As the scale of the cluster expands, communication overhead gradually surpasses computation overhead (according to Chang et al. (2024), communication can take 20% – 80% runtime in training and inference), diluting the benefits gained from improved computation kernel performance. As a result, it is necessary to jointly optimize computation and communication and overlap their latency with each other.

Although previous work (Jangda et al., 2022; Google, 2025; Wang et al., 2023; Chang et al., 2024; Zheng et al., 2025; Zhang et al., 2025; Zhao et al., 2025; Gond et al., 2025) provides domain-specific languages (DSL) or hand-tuned kernels for overlapping, users still struggle to develop new kernels based on these tools. This difficulty arises from rigid DSL constraints and the high level of expertise required for low-level kernel programming. The root cause behind this challenge is that these frameworks lack suitable programming abstractions that align well with both the algorithmic characteristics of LLMs and the hardware attributes of distributed clusters.

The major difference between distributed programming and single-device programming lies in the existence of non-uniform memory access (NUMA), which can be viewed as two additional levels of

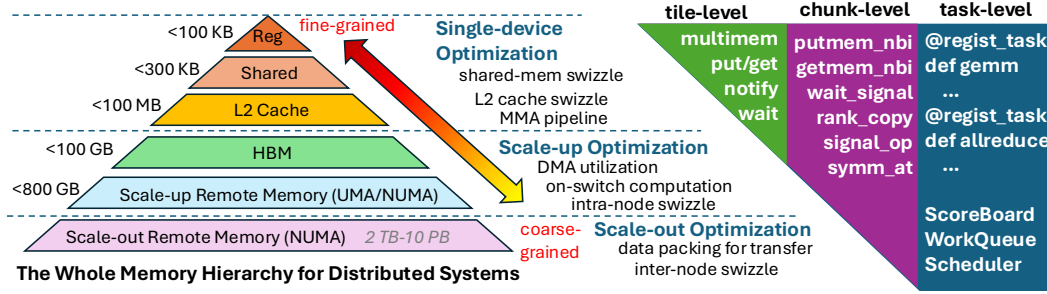


Figure 1: DITRON provides three levels of programming abstractions, which corresponds to the hardware hierarchy of distributed systems.

the memory hierarchy: scale-up remote memory and scale-out remote memory (shown in Figure 1). We summarize distributed NUMA programming model as the following:

1. Memory consistency by signals: Global synchronization in distributed GPU cluster is expensive in terms of latency. As a result, in-HBM signals are used for peer-to-peer synchronization. Signals are constructed and accessible by all the GPUs.

2. Non-blocking P2P copy and blocking multi-memory operations: On one hand, due to the inconsistency and long latency of remote memory access, peer-to-peer (P2P) data transfer between GPUs are asynchronous, making it possible to overlap with other computation and communication operations. On the other hand, modern GPUs equipped with NVLink Sharp support on-switch reduce computation and broadcast for acceleration, which is called multi-memory operations.

3. Resource partition for task pipeline: Communication and computation compete with each other for hardware resources, including streaming multiprocessors (SM), L2 cache, and HBM bandwidth. Efficient programming requires separate dedicated SMs for communication. The GPU is logically partitioned into multiple parts; each part is responsible for one task (computation or communication), forming a producer-consumer pipeline.

We aim to provide a new programming interface for distributed programming. The principle of our design is to offer minimal programming abstractions while maintaining consistency with existing compilers (Tillet et al., 2019) for computation when falling back to single-device programming, ensuring that the programming complexity is manageable for most researchers. Through extensive practice and refinement, we have developed the DITRON compiler, with its key abstractions shown in Figure 1:

1. Tile-level: For computation, DITRON inherits all the tile-level programming primitives from existing compilers (Tillet et al., 2019), thus retaining standard optimizations for single devices. For communication, we first introduce multi-memory primitives to leverage on-switch communication acceleration; second, we add tile-level remote memory access primitives; third, we incorporate wait/notify signal operations for intra-node synchronization.

2. Chunk-level: For data transfer, DITRON supports non-blocking put memory (to remote) and get memory (from remote) operations as well as chunk-level signal operations for inter-node synchronization. We also introduce new optimizations including low-latency protocol and intra/inter-node swizzling.

3. Task-level: Whether for computation or communication, distinct tasks ought to be mapped to separate hardware resources and pipelined through sophisticated scheduling. DITRON offers a task-level interface that supports task registration, dependency control, and task scheduling (all with minimal manual intervention) while generating efficient MegaKernels in real time.

The contributions of DITRON are twofold. First, we provide a simple yet efficient programming interface via compiler techniques, enabling researchers to add flexible parallel acceleration to their single-device code with minimal manual effort. This capability can potentially encourage more exploration of model architectures and the design of new parallel strategies. Second, through extensive validation and evaluation, DITRON achieves a $1.27\times - 19.18\times$ average speedup over non-overlapped

kernels across a wide range of workloads, and even outperforms expert-tuned overlapping libraries implemented in CUDA by 6% – 30%. It also delivers a 5% – 30% end-to-end performance improvement for vLLM (Kwon et al., 2023) under large batch sizes. Validation across more than 10,000 GPUs demonstrates DITRON’s scalability and robustness for practical industrial deployment, saving millions of GPU hours per month. Additionally, DITRON supports AMD GPUs and PCIe-connected GPUs: on AMD GPUs, it achieves a 2% – 38% speedup over RocmBLAS+RCCL; on L20 GPUs, it delivers an average speedup of $8.33\times$ over CuBLAS+NCCL. DITRON is publicly available at [https://github.com/\[blinded-for-review\]](https://github.com/[blinded-for-review])¹.

2 HARDWARE AND SOFTWARE IN DISTRIBUTED SYSTEM

A distributed system consists of multiple levels of hardware, the basic building block of which include GPUs, host CPUs, and switches. The overview of a GPU cluster is shown in Figure 2.

Computation and Communication Hardware:

In distributed systems, the primary computing hardware is the GPU, which can be divided into three levels: warp, block, and grid. The warp level controls computing units such as CUDA Cores (for scalar and vector calculations) and Tensor Cores (for matrix multiplication). Block-level is responsible for warp partition and shared memory control. The grid level manages workload mapping across streaming multiprocessors (SMs) via scheduling and controls L2 cache access order through swizzling.

The exploitation of single-GPU hardware resources to achieve high performance has been well studied in prior works (Nvidia, 2022; Dao et al., 2022; DeepSeek-AI, 2025; Spector et al., 2024), and we recommend readers refer to these for further details.

The communication aspect can be categorized into two domains: the intra-node scale-up domain and the inter-node scale-out domain. Within the scale-up domain, there are two memory access modes: UMA and NUMA. In UMA, the address space of each GPU is shared among all GPUs, enabling data copying between any two GPUs via standard load/store instructions or DMA engines (dedicated hardware units for direct memory access in GPUs). In NUMA, however, each GPU’s memory space is symmetric. To access data from another GPU, local data pointers in HBM must first be mapped to the symmetric address space and then translated to the remote GPU’s address.

UMA access via DMA engines is triggered by host CPUs rather than GPU SMs, thereby freeing up more resources for computation. In contrast, UMA access through load/store instructions (multi-memory access) occupies SMs, with reduction and broadcast operations achievable via dedicated hardware known as NVLink Sharp (NVIDIA, 2018). By contrast, NUMA access is always initiated by GPU-side SMs, leaving fewer SMs available for computation.

For inter-node domain, only NUMA mode is supported. GPU’s memory space is mapped to symmetric address space and data copy is done by network interface cards (NICs) and network switches, allowing asynchronous (non-blocking, or *nbi* for short) data copy from local GPU to remote GPU (putmem) or from remote to local (getmem).

Overlapping Library and Compiler: DITRON draws inspiration from three lines of research: compute-communication overlapping libraries (Chang et al., 2024; Zhang et al., 2025; Gond et al., 2025; Hong et al., 2025), single-device compilers (Tillet et al., 2019), and compute-communication overlapping compilers (Wang et al., 2023; Jangda et al., 2022; Google, 2025; Zheng et al., 2025). FLUX (Chang et al., 2024), COMET (Zhang et al., 2025) are libraries developed using CUTLASS (Nvidia, 2022), which enable overlapping by inlining wait/notify PTX assembly into CUTLASS matrix multiplication code and hand-crafting communication kernels using NVSHMEM (NVIDIA, 2025). The lack of programming abstraction makes it hard to develop new kernels

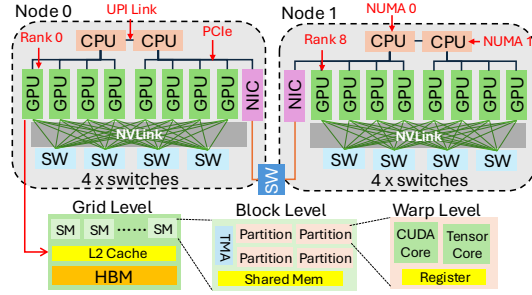


Figure 2: GPU Cluster Overview.

¹The code is open-sourced, and the URL is anonymized to comply with double-blind review requirements.

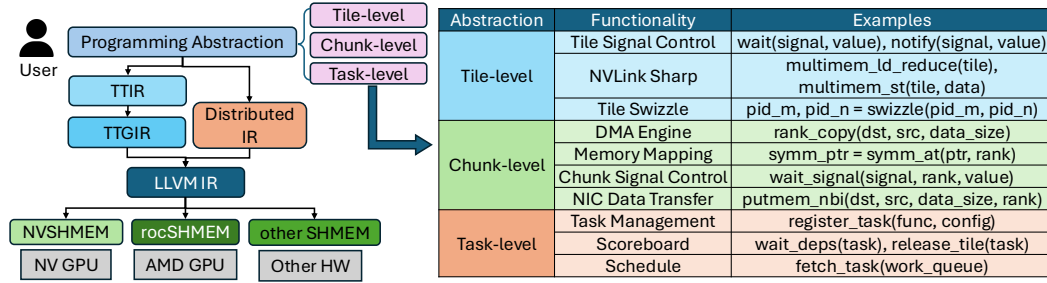


Figure 3: DITRON compiler stack includes (1) programming abstraction; (2) distributed IR; (3) LLVM IR and hardware-specific OpenSHMEM implementation. The right side shows examples for three levels of programming abstractions.

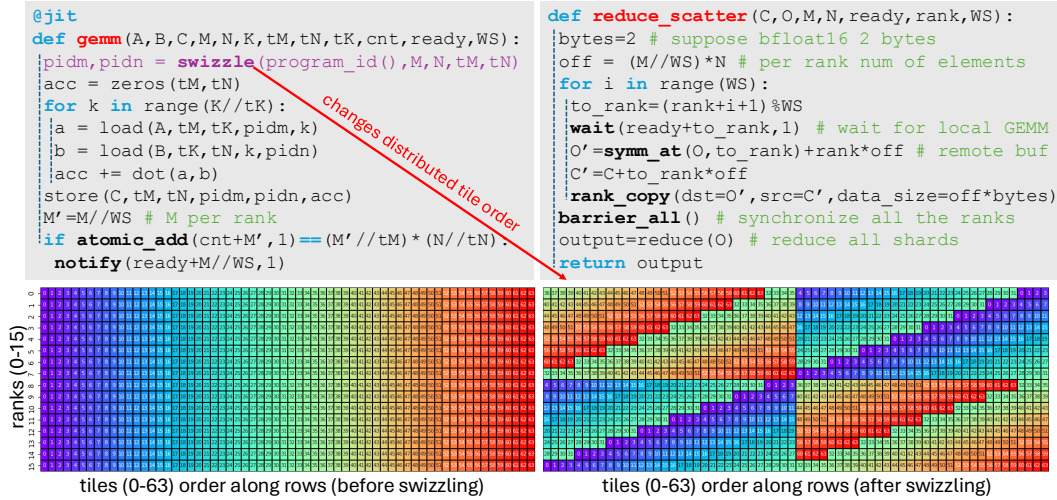


Figure 4: Simplified GEMM+ReduceScatter code example and distributed swizzling. The ReduceScatter requires that each GPU rank i takes 4 data chunks ($[4i, 4i + 4)$) eventually. Tiles are processed from left to right. Without swizzling, all ranks start from data tiles for rank 0 and the other ranks are blocked, while with swizzling, all the ranks can start without blocking. A clearer enlarged swizzle view is presented in the Appendix B and the full code example is in Appendix E.2.

using these libraries. Triton (Tillet et al., 2019) compiler provides a set of NumPy-like primitives to reduce programming complexity. TileLink (Zheng et al., 2025) extends Triton’s primitives to intra-node communication and enables overlapping kernels such as AllGather-GEMM and GEMM-ReduceScatter. The programming abstraction of TileLink only supports the *tile-level* in Figure 1. Similarly, other compilers such as CoCoNet (Jangda et al., 2022), Dist-Einsum (Wang et al., 2023), Pallas (Google, 2025), and Mirage (Wu et al., 2025) design new DSLs that only cover part of the abstractions in Figure 1. But in this work, we aim to support all the three levels of abstractions.

3 DITRON

We present DITRON, a flexible and versatile compiler designed for compute-communication overlapping kernel programming. An overview of DITRON is provided in Figure 3, which also includes primitive examples for the three levels of programming abstractions. In what follows, we first elaborate on the programming abstractions. Figure 4 presents a simplified code example for GEMM+ReduceScatter, upon which we then explain the swizzle optimizations. Finally, we detail the implementation specifics of DITRON.

3.1 THREE LEVELS OF PROGRAMMING ABSTRACTIONS

Tile-level Abstraction: In addition to the tile-level computation abstraction inherited from Triton (Tillet et al., 2019), DITRON further provides abstraction for tile-level signal control and distributed data transfer:

(1) Scale-up domain signals are viewed as tensors located in HBM. The signals can be used for spin-locks (through *wait* and *notify* primitive) to control producer-consumer execution. One signal can be shared by a tile of data, indicating that the whole tile of data is controlled by this signal.

(2) Distributed data is treated as tiles that reside in the HBM of remote GPUs. Although the data is physically located in remote HBM, local kernels can still hold a view (or pointers) of remote tiles. DITRON provides primitives to fetch the remote tiles and perform reduction/broadcast via NVLink Sharp (through *multimem_ld_reduce* and *multimem_st* primitives).

(3) Distributed data tiles can be swizzled like regular tiles to enhance performance. For distributed communication, local HBM is treated as a cache for remote memory, and data from remote memory arrives in local HBM at distinct steps. Consequently, to leverage temporal locality, tiles corresponding to earlier steps should be scheduled before other tiles, which is achieved via the *swizzle* primitive.

Chunk-level Abstraction: To leverage the hardware features of DMA engines and NICs, DITRON further provides chunk-level abstraction. Tile-level data transfer in scale-out domain is inefficient because the low-level hardware is designed for 1-D contiguous data transfer instead of N-D tile transfer. N-D tile transfer should be translated to multiple small 1-D data transfer at a low level, hampering the achieved bandwidth. Our chunk-level abstraction design includes:

(1) Data chunk is composed of multiple data tiles. The data tiles are packed together to form a larger data chunk before communication. For example, in our GEMM+AllToAll kernel (designed for Ulysses sequence parallel (Jacobs et al., 2023)), data should be packed along head dimension before AllToAll communication and concatenated along sequence dimension after communication.

(2) Distributed data chunks are symmetric, the address of which can be mapped to remote GPU’s address space. GPUs across nodes can also share the symmetric view of data chunks (through *symm_at* primitive).

(3) Distributed data chunk transfer is asynchronous (through *rank_copy* and *put/getmem_nbi*) and can be synchronized by dedicated signal operations (through *wait_signal*). As a result, the data chunks can also be swizzled to exploit temporal locality, where scale-up remote memory becomes cache for scale-out remote memory.

Task-level Abstraction: Computation and communication kernels require dedicated hardware resources to achieve best performance. Manually scheduling the resource partition and execution order is error-prone and inefficient. DITRON provides lightweight task-level abstraction to help combine the computation and communication into one optimized MegaKernel:

(1) Both computation kernel and communication kernel can be converted to tasks. DITRON uses a unified interface to wrap kernels as tasks. The interface encapsulates (through *register_task*) the kernel’s input/output tensors, the number of tiles within the kernel, and the tile dependency to its producer kernels.

(2) Tasks should be executed in producer-consumer order, which is controlled by a software scoreboard. DITRON provides scoreboard implementation and corresponding *wait_deps* and *release_tile* primitives. The dependency of tasks is captured during task registration.

(3) The mapping from task to physical SM can be scheduled and optimized through a scheduler. DITRON provides a scheduler interface that supports compile-time round-robin and zig-zag scheduling. Customized scheduling can be also added by users. Besides, DITRON also supports runtime scheduling (via *fetch_task* primitive) that chooses tasks for each SM dynamically.

3.2 SWIZZLING OPTIMIZATION

As noted in previous sections, the tile-level and chunk-level abstraction designs incorporate distributed tile/chunk swizzle optimizations for distributed scenarios to enhance locality in remote memory. DITRON supports two swizzle modes: gather mode and scatter mode. The key insight

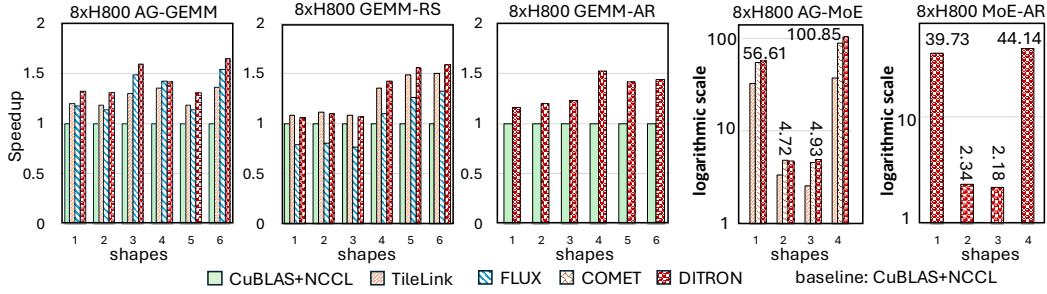


Figure 5: Evaluation of single workload on $8\times$ H800 GPUs. Workloads are AllGather+GEMM (AG-GEMM), GEMM+ReduceScatter (GEMM-RS), GEMM+AllReduce (GEMM-AR), AllGather+MoE (AG-MoE), MoE+AllReduce (MoE-AR). A logarithmic scale is used for the speedup of MoE due to its wide span, and the specific values are listed in the Appendix D.

is that gather mode aims to delay long-distance data transfers since computation depends on communication, the latter must be initiated as early as possible. In contrast, scatter mode prioritizes long-distance data transfers because communication can overlap with subsequent computation if started early, it is given precedence.

For the simplified GEMM+ReduceScatter example in Figure 4, we adopt a global perspective of all data tiles across all parallel devices. The heatmaps at the bottom illustrate 16 GPUs, each tasked with computing 64 result tiles. Each rank processes tiles from left to right. Without swizzling, only the data tiles required by rank 0 are processed initially, while other tiles remain unread, which blocks the execution of ranks dependent on these tiles. With swizzling, the order is rescheduled such that each rank has at least one data tile computed initially.

Non-perfect Shape Problem: For real-world workloads, input shapes may not be perfectly divisible by tile sizes, and cross-rank tiles/chunks exist. Swizzling for such workloads differs slightly. In gather mode, cross-rank tiles/chunks are postponed because they depend on multiple communication steps. In scatter mode, by contrast, they are prioritized: since multiple ranks rely on these tiles/chunks, processing them earlier allows subsequent computations to start sooner. We put the swizzling views for non-perfect tiling in Appendix B.

3.3 IMPLEMENTATION DETAILS OF DITRON

DITRON is implemented with around 59,000 lines of Python code and 7,000 lines of C++ code, consisting primarily of a compiler stack and an optimized kernel library.

Compiler Stack: The compiler stack of DITRON is illustrated in Figure 3. Users write NumPy-like programs using our three levels of abstractions, which cover both computation and communication. We have implemented a distributed IR for compilation: single-device semantics (e.g., *dot*, *load/store*) in the programs are lowered to standard Triton IR (TTIR) and Triton GPU IR (TTGIR), while distributed semantics are lowered to our distributed IR. This distributed IR adheres to the OpenSHMEM standard, which is a communication standard designed for distributed systems. OpenSHMEM has been implemented in efficient low-level communication libraries by NVIDIA (NVSHMEM) and AMD (roC SHMEM), with support for other hardware achievable via custom SHMEM libraries. The distributed IR is then lowered to LLVM IR, where we leverage LLVM’s *CallExtern* capability to invoke external symbols from the SHMEM library.

Other Optimizations for Kernels:

(1) Low-latency (LL) Protocol: State-of-the-art communication libraries NCCL (NVIDIA, 2024) supports the LL and LL128 protocols, which are synchronization-free but incur approximately 50% bandwidth waste. DITRON also supports the LL protocol for latency-oriented kernels (e.g., batch size 1).

(2) Device-to-device (D2D) copy elimination: We observe that data movement via default driver APIs (*cudaMemcpy*) or PyTorch APIs (*copy*) can lead to stragglers and unexpected SM utilization,

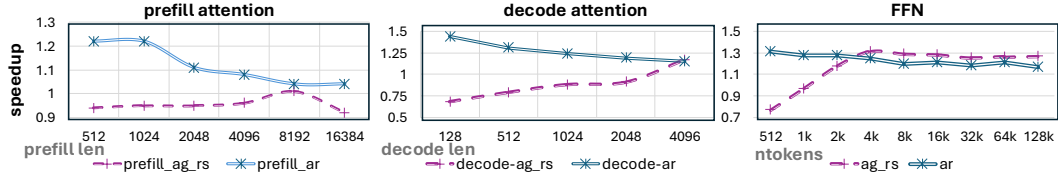


Figure 6: Module-level TP evaluation for Qwen3-32B: Results for prefill and decode attention show that DITRON with AllReduce achieves better speedups over CuBLAS+NCCL compared to using AllGather and ReduceScatter. For MLP, results indicate that the combination of AllGather and ReduceScatter performs better for large shapes, while AllReduce is superior for small shapes.

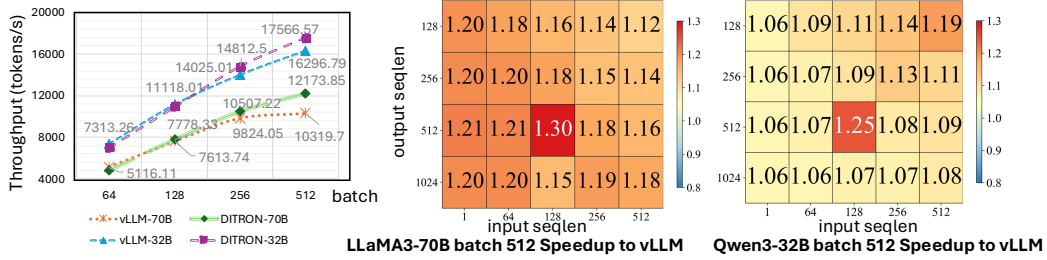


Figure 7: End-to-end TP evaluation results for Qwen3-32B and LLaMA3-70B on 8 H800s: DITRON demonstrates advantages over vLLM for large batch sizes and large models.

rendering overlap ineffective. Thus, we integrate these copies into our communication or computation kernels to prevent slowdowns.

(3) PCIe non-atomic synchronization and ring communication: The PCIe link provides no hardware guarantee for data transfer order, thus not supporting atomic operations. To deploy DITRON on PCIe-based GPUs such as the RTX 4090 and L20, we implement barriers using volatile load/store operations, eliminating the need for atomic operations. Additionally, PCIe-based communication requires a ring topology, and cross-UPI communication should be offloaded to the NIC. DITRON supports ring-based communication for AllGather and ReduceScatter, and implements cross-UPI optimizations to achieve high performance.

4 EXPERIMENTS

DITRON supports kernel implementation and optimization for various parallel strategies (TP, SP, EP). In our experiments, we evaluate DITRON under different parallel configurations. For inference, where the primary parallel strategy is TP within a single node, we first test 5 distinct workloads for GEMM and MoE. Second, we integrate these into attention modules and FFN modules to assess module-level performance. Third, we evaluate end-to-end inference using Qwen3-32B and LLaMA3-70B. For training, we test TP, SP, and EP kernels implemented by DITRON across 8 to 128 GPUs, examining both weak and strong scaling. Finally, we use DITRON’s kernels to scale the training of a Qwen3-like 640B MoE model to 11,520 GPUs. DITRON also supports AMD GPUs and PCIe GPUs, which is discussed at the end of the evaluation section.

4.1 EVALUATION FOR INFERENCE

Single Workload Evaluation: We evaluated 5 distinct workloads on 8 H800 GPUs across 26 configurations, with each workload featuring multiple shape configurations derived from real-world models such as LLaMA (Dubey et al., 2024), Mixtral (Jiang et al., 2024), GPT (OpenAI, 2023), Qwen (Yang et al., 2025), and DeepSeek (DeepSeek-AI et al., 2024a). Detailed shape configurations are provided in Appendix C.2. Our baselines include CuBLAS (NVIDIA,

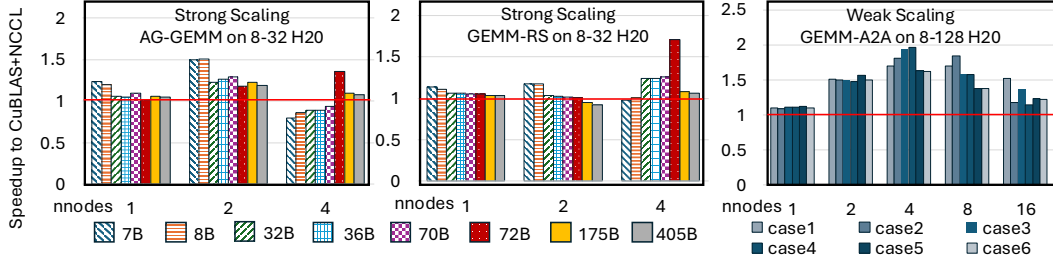


Figure 8: Scaling TP and SP workloads on H20 clusters: Results are presented as speedups over CuBLAS+NCCL, with shapes derived from various real-world LLMs (details in Appendix C.2). DITRON maintains speedups for AG-GEMM across up to 16 GPUs. For 32 GPUs, speedups are achievable only with shapes from large LLMs. The speedup of GEMM-RS remains consistent across 8–32 GPUs. GEMM-A2A exhibits the best weak scaling due to its scalable AllToAll.

2022)+NCCL (NVIDIA, 2024) (non-overlapping), TileLink (Zheng et al., 2025), FLUX (Chang et al., 2024), and COMET (Zhang et al., 2025), with results presented in Figure 5.

Overall, for AG-GEMM, the geometric speedup of DITRON is $1.43\times$ to CuBLAS+NCCL, $1.13\times$ to TileLink, and $1.09\times$ to FLUX. For GEMM-RS, the geometric speedup of DITRON is $1.27\times$ to CuBLAS+NCCL, $1.02\times$ to TileLink, and $1.30\times$ to FLUX. TileLink and FLUX have no support for AllReduce. For GEMM-AR, the geometric speedup to CuBLAS+NCCL is $1.32\times$. The speedup of AG-GEMM and GEMM-RS mainly comes from the overlapping of communication instead of faster GEMM (actually, we use Triton’s GEMM and the GEMM is slightly slower than that of CuBLAS and FLUX), nearly 87.5% communication latency is hidden by computation for large input shapes. The speedup of GEMM-AR comes from faster AllReduce implemented using DITRON, which supports both one-shot algorithm and two-shot algorithm via multi-memory reduction/broadcast provided by NVLink Sharp. For AG-MoE, the geometric speedup of DITRON is $19.18\times$ to CuBLAS+NCCL, $1.89\times$ to TileLink, and $1.06\times$ to COMET. For MoE-AR, the average speedup to CuBLAS+NCCL is $13.89\times$.

Module-level Evaluation: We incorporate the aforementioned workloads into attention modules and FFN modules. The QKV projection and output projection in attention modules are replaced by DITRON kernels. Prefill and decode performance is presented in Figure 6. Prefill input lengths range from 512 to 16k, and decode output lengths range from 128 to 4k. The two lines represent the difference in speedups over CuBLAS+NCCL between using AllGather+ReduceScatter and using AllReduce. The results indicate that for attention modules, both prefill and decode prefer AllReduce communication because the reduction dimension for GEMM (head dimension size in attention) is small (128), making GEMM latency not enough to hide communication latency of AllGather or ReduceScatter. The geometric mean speedup of DITRON’s attention module (using AllReduce) to CuBLAS+NCCL is $1.12\times$ for prefill and $1.26\times$ for decode. On the other hand, GEMMs in FFN modules use large intermediate size and the latency of GEMMs can hide communication latency when given enough input tokens (large batch or sequence length), so overlapping AllGather and ReduceScatter is better than using AllReduce for sequence length larger than 2k. The speedup for 128k tokens of DITRON is $1.17\times$ using AllReduce and $1.27\times$ using AllGather+ReduceScatter.

End-to-end Evaluation: We evaluate DITRON using end-to-end models LLaMA3-70B (Dubey et al., 2024) and Qwen3-32B (Yang et al., 2025). The results are shown in Figure 7. Our baseline is vLLM (Kwon et al., 2023), which is a strong baseline for TP inference as vLLM employs efficient AllReduce kernels designed by experts. For batch size less than 128, vLLM is slightly better than DITRON in throughput. But for batch sizes larger than 128, DITRON achieves 5% – 30% speedup

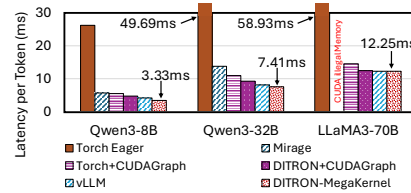


Figure 9: Batch 1 inference latency for Qwen3-8B, Qwen3-32B, LLaMA3-70B using DITRON’s distributed Megakernel on $8\times$ H800.

to vLLM. Specially, for batch size 512, we show the detailed speedup under different input lengths and output lengths for LLaMA3-70B and Qwen3-32B in Figure 7. The speedup translates to 12k tokens/s throughput for LLaMA3-70B and 17k tokens/s throughput for Qwen3-32B.

Distributed MegaKernel Evaluation: For single batch inference, the hardware resources are usually underutilized, task-level scheduling that produces MegaKernel can eliminate kernel launch overhead, increase SM activity, and improve end-to-end performance. We compare DITRON’s distributed MegaKernel with PyTorch (Eager mode and CUDAGraph mode), DITRON+CUDAGraph, and Mirage (Wu et al., 2025) as shown in Figure 9. The geometric speedup is $6.28\times$ to Torch Eager, $1.73\times$ to Mirage, $1.33\times$ to Torch+CUDAGraph, $1.11\times$ to DITRON+CUDAGraph, and $1.10\times$ to vLLM. We also add MegaKernel code examples in Appendix E.4.

4.2 EVALUATION FOR TRAINING

Scale Single Workload: For training, we first evaluate both strong scaling and weak scaling performance for TP (AllGather+GEMM, GEMM+ReduceScatter), SP (GEMM+AllToAll), and EP (MoE Dispatch and MoE Combine) using DITRON. The results for TP and SP are shown in Figure 8 and the results for EP are shown in Figure 10. The detailed shape configurations are put in Appendix C.2. More results (on both H800 and H20) are put in Appendix D.

TP workloads renders heavy communication among nodes and the low bandwidth of network communication makes scaling non-beneficial. As a result, we only observe speedups of AllGather+GEMM and GEMM+ReduceScatter for 8-32 GPUs (strong scaling, number of total tokens is 32768), which ranges from $0.80\times$ to $1.71\times$ to CuBLAS+NCCL. For cases where the speedup is lower than 1, the main cause is GEMM becomes too small for each GPU after sharding, which cannot hide communication latency. As a comparison, GEMM+AllToAll (GEMM-A2A) gives consistent speedups for weak scaling (sequence length per rank remains unchanged) from 8 GPUs to 128 GPUs, where the shape of GEMM keeps constant. As for EP (we set 8192 tokens per rank, topk 8, hidden size 7168), dispatch and combine performance remains similar for strong scaling (globally 512 experts in total) and weak scaling (8 experts per rank). The speedup to PyTorch+NCCL implementations ranges from $1.04\times$ to $4.70\times$.

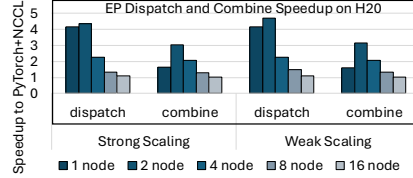


Figure 10: Scaling EP dispatch and combine using DITRON.

End-to-end Training of MoE Models: DITRON has been validated in industrial training settings on 11, 520 H20 GPUs using a training task that scales a Qwen3-like² MoE model to around 640B (33B activated). The achieved MFU is around 0.29. Compared to native PyTorch implementations, DITRON helps saving millions of GPU hours per month.

4.3 SUPPORT FOR OTHER PLATFORMS

DITRON can be transferred to more hardware platforms due to the flexibility of our distributed IR and the compatibility of dependent Triton compiler. For now, we manage to support AMD GPUs and PCIe GPUs (L20). We put our preliminary results in Appendix D. Overall, on AMD GPUs the speedup to RocmBLAS+RCCL ranges from 2% – 38%; on L20 GPUs the average speedup to CuBLAS+NCCL is $8.33\times$.

5 CONCLUSION

Distributed inference and training becomes necessary and requires more researchers to be able to program distributed kernels. This work presents DITRON, a flexible and versatile distributed compiler with tile-level, chunk-level, and task-level interfaces for overlapping kernels. Researchers can use DITRON to program efficient distributed kernels for tensor parallel, sequence parallel, and expert parallel with performance comparable to or better than expert-tuned kernels.

²Details omitted to protect commercial secrets

ETHICS STATEMENT

This research focuses on compiler optimization techniques for distributed computing systems, aiming to enhance the computational efficiency of training and inference for LLMs. Our work does not directly involve human subjects, personally identifiable information, or sensitive datasets. We believe that DITRON has positive ethical implications. By significantly improving the efficiency of computation-communication overlap, our tool can effectively reduce the total compute time and energy consumption required for large-scale models, which helps mitigate the environmental impact of AI development.

REPRODUCIBILITY STATEMENT

We are committed to ensuring the reproducibility of our research. To this end, we have made the following efforts:

1. The complete source code for the DITRON compiler has been made publicly available upon submission. For the consideration of anonymity, we choose not to include any effective link or code package with this submission. However, we are willing to supplement the code package at any time, provided that the paper is accepted for publication or with the permission of the reviewers and the organizing committee during review process.
2. The appendix of the paper provides several simplified code examples, including All-Gather GEMM, GEMM Reduce-Scatter, GEMM All-Reduce, and MegaKernel, which clearly demonstrate the practical usage of DITRON’s programming abstractions.

REFERENCES

- AMD. URL <https://github.com/ROCm/rocSHMEM>.
- Anthropic. Claude 3.5 sonnet. <https://www.anthropic.com/news/claude-3-5-sonnet>, 2024.
- Li-Wen Chang, Wenlei Bao, Qi Hou, Chengquan Jiang, Ningxin Zheng, Yinmin Zhong, Xuanrun Zhang, Zuquan Song, Ziheng Jiang, Haibin Lin, Xin Jin, and Xin Liu. FLUX: fast software-based communication overlap on gpus through kernel fusion. *CoRR*, abs/2406.06858, 2024. doi: 10.48550/ARXIV.2406.06858. URL <https://doi.org/10.48550/arXiv.2406.06858>.
- Gheorghe Comanici, Eric Bieber, Mike Schaekermann, Ice Pasupat, Noveen Sachdeva, Inderjit Dhillon, Marcel Blistein, Ori Ram, Dan Zhang, Evan Rosen, et al. Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. *arXiv preprint arXiv:2507.06261*, 2025.
- Tri Dao, Daniel Y. Fu, Stefano Ermon, Atri Rudra, and Christopher Ré. Flashattention: Fast and memory-efficient exact attention with io-awareness. *CoRR*, abs/2205.14135, 2022. doi: 10.48550/arXiv.2205.14135. URL <https://doi.org/10.48550/arXiv.2205.14135>.
- DeepSeek-AI. Deepseek deepgemm. <https://github.com/deepseek-ai/DeepGEMM>, 2025.
- DeepSeek-AI, Aixin Liu, Bei Feng, Bin Wang, Bingxuan Wang, Bo Liu, Chenggang Zhao, Chengqi Deng, Chong Ruan, Damai Dai, Daya Guo, Dejian Yang, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, Hao Zhang, Hanwei Xu, Hao Yang, Haowei Zhang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Li, Hui Qu, J. L. Cai, Jian Liang, Jianzhong Guo, Jiaqi Ni, Jiashi Li, Jin Chen, Jingyang Yuan, Junjie Qiu, Junxiao Song, Kai Dong, Kaige Gao, Kang Guan, Lean Wang, Lecong Zhang, Lei Xu, Leyi Xia, Liang Zhao, Liyue Zhang, Meng Li, Miaojun Wang, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Mingming Li, Ning Tian, Panpan Huang, Peiyi Wang, Peng Zhang, Qihao Zhu, Qinyu Chen, Qiusi Du, R. J. Chen, R. L. Jin, Ruiqi Ge, Ruizhe Pan, Runxin Xu, Ruyi Chen, S. S. Li, Shanghao Lu, Shangyan Zhou, Shanhuang Chen, Shaoqing Wu, Shengfeng Ye, Shirong Ma, Shiyu Wang, Shuang Zhou, Shuiping Yu, Shunfeng Zhou, Size Zheng, Tao Wang, Tian Pei, Tian Yuan, Tianyu Sun, W. L.

- Xiao, Wangding Zeng, Wei An, Wen Liu, Wenfeng Liang, Wenjun Gao, Wentao Zhang, X. Q. Li, Xiangyue Jin, Xianzu Wang, Xiao Bi, Xiaodong Liu, Xiaohan Wang, Xiaojin Shen, Xiaokang Chen, Xiaosha Chen, Xiaotao Nie, and Xiaowen Sun. Deepseek-v2: A strong, economical, and efficient mixture-of-experts language model. *CoRR*, abs/2405.04434, 2024a. doi: 10.48550/ARXIV.2405.04434. URL <https://doi.org/10.48550/arXiv.2405.04434>.
- DeepSeek-AI, Aixin Liu, Bei Feng, Bing Xue, Bingxuan Wang, Bochao Wu, Chengda Lu, Cheng-gang Zhao, Chengqi Deng, Chenyu Zhang, Chong Ruan, Damai Dai, Daya Guo, Dejian Yang, Deli Chen, Dongjie Ji, Erhang Li, Fangyun Lin, Fucong Dai, Fuli Luo, Guangbo Hao, Guanting Chen, Guowei Li, H. Zhang, Han Bao, Hanwei Xu, Haocheng Wang, Haowei Zhang, Honghui Ding, Huajian Xin, Huazuo Gao, Hui Li, Hui Qu, J. L. Cai, Jian Liang, Jianzhong Guo, Jiaqi Ni, Jiashi Li, Jiawei Wang, Jin Chen, Jingchang Chen, Jingyang Yuan, Junjie Qiu, Junlong Li, Junxiao Song, Kai Dong, Kai Hu, Kaige Gao, Kang Guan, Kexin Huang, Kuai Yu, Lean Wang, Lecong Zhang, Lei Xu, Leyi Xia, Liang Zhao, Litong Wang, Liyue Zhang, Meng Li, Miaojun Wang, Mingchuan Zhang, Minghua Zhang, Minghui Tang, Mingming Li, Ning Tian, Panpan Huang, Peiyi Wang, Peng Zhang, Qiancheng Wang, Qihao Zhu, Qinyu Chen, Qiushi Du, R. J. Chen, R. L. Jin, Ruiqi Ge, Ruisong Zhang, Ruizhe Pan, Runji Wang, Runxin Xu, Ruoyu Zhang, Ruyi Chen, S. S. Li, Shanghao Lu, Shangyan Zhou, Shanhua Chen, Shaoqing Wu, Shengfeng Ye, Shengfeng Ye, Shirong Ma, Shiyu Wang, Shuang Zhou, Shuiping Yu, Shunfeng Zhou, Shuting Pan, T. Wang, Tao Yun, Tian Pei, Tianyu Sun, W. L. Xiao, and Wangding Zeng. Deepseek-v3 technical report. *CoRR*, abs/2412.19437, 2024b. doi: 10.48550/ARXIV.2412.19437. URL <https://doi.org/10.48550/arXiv.2412.19437>.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, Aurélien Rodriguez, Austen Gregerson, Ava Spataru, Baptiste Rozière, Bethany Biron, Binh Tang, Bobbie Chern, Charlotte Caucheteux, Chaya Nayak, Chloe Bi, Chris Marra, Chris McConnell, Christian Keller, Christophe Touret, Chunyang Wu, Corinne Wong, Cristian Canton Ferrer, Cyrus Nikolaidis, Damien Allonsius, Daniel Song, Danielle Pintz, Danny Livshits, David Esiobu, Dhruv Choudhary, Dhruv Mahajan, Diego Garcia-Olano, Diego Perino, Dieuwke Hupkes, Egor Lakomkin, Ehab AlBadawy, Elina Lobanova, Emily Dinan, Eric Michael Smith, Filip Radenovic, Frank Zhang, Gabriel Synnaeve, Gabrielle Lee, Georgia Lewis Anderson, Graeme Nail, Grégoire Mialon, Guan Pang, Guillem Cucurell, Hailey Nguyen, Hannah Korevaar, Hu Xu, Hugo Touvron, Iliyan Zarov, Imanol Arrieta Ibarra, Isabel M. Kloumann, Ishan Misra, Ivan Evtimov, Jade Copet, Jaewon Lee, Jan Geffert, Jana Vranes, Jason Park, Jay Mahadeokar, Jeet Shah, Jelmer van der Linde, Jennifer Billock, Jenny Hong, Jenya Lee, Jeremy Fu, Jianfeng Chi, Jianyu Huang, Jiawen Liu, Jie Wang, Jiecao Yu, Joanna Bitton, Joe Spisak, Jongsoo Park, Joseph Rocca, Joshua Johnstun, Joshua Saxe, Junteng Jia, Kalyan Vasuden Alwala, Kartikeya Upasani, Kate Plawiak, Ke Li, Kenneth Heafield, Kevin Stone, and et al. The llama 3 herd of models. *CoRR*, abs/2407.21783, 2024. doi: 10.48550/ARXIV.2407.21783. URL <https://doi.org/10.48550/arXiv.2407.21783>.
- Raja Gond, Nipun Kwatra, and Ramachandran Ramjee. Tokenweave: Efficient compute-communication overlap for distributed llm inference, 2025. URL <https://arxiv.org/abs/2505.11329>.
- Google. Pallas, 2025. URL <https://docs.jax.dev/en/latest/pallas/index.html>.
- Ke Hong, Xiuhong Li, Minxu Liu, Qiuli Mao, Tianqi Wu, Zixiao Huang, Lufang Chen, Zhong Wang, Yichong Zhang, Zhenhua Zhu, Guohao Dai, and Yu Wang. Flashoverlap: A lightweight design for efficiently overlapping communication and computation. *CoRR*, abs/2504.19519, 2025. doi: 10.48550/ARXIV.2504.19519. URL <https://doi.org/10.48550/arXiv.2504.19519>.
- Sam Ade Jacobs, Masahiro Tanaka, Chengming Zhang, Minjia Zhang, Shuaiwen Leon Song, Samyam Rajbhandari, and Yuxiong He. DeepSpeed Ulysses: System optimizations for enabling training of extreme long sequence transformer models. *CoRR*, abs/2309.14509, 2023. doi: 10.48550/ARXIV.2309.14509. URL <https://doi.org/10.48550/arXiv.2309.14509>.

- Abhinav Jangda, Jun Huang, Guodong Liu, Amir Hossein Nodehi Sabet, Saeed Maleki, Youshan Miao, Madanlal Musuvathi, Todd Mytkowicz, and Olli Saarikivi. Breaking the computation and communication abstraction barrier in distributed machine learning workloads. In Babak Falsafi, Michael Ferdman, Shan Lu, and Thomas F. Wenisch (eds.), *ASPLOS '22: 27th ACM International Conference on Architectural Support for Programming Languages and Operating Systems*, Lausanne, Switzerland, 28 February 2022 - 4 March 2022, pp. 402–416. ACM, 2022. doi: 10.1145/3503222.3507778. URL <https://doi.org/10.1145/3503222.3507778>.
- Albert Q. Jiang, Alexandre Sablayrolles, Antoine Roux, Arthur Mensch, Blanche Savary, Chris Bamford, Devendra Singh Chaplot, Diego de Las Casas, Emma Bou Hanna, Florian Bressand, Gianna Lengyel, Guillaume Bour, Guillaume Lample, L  lio Renard Lavaud, Lucile Saulnier, Marie-Anne Lachaux, Pierre Stock, Sandeep Subramanian, Sophia Yang, Szymon Antoniak, Teven Le Scao, Th  ophile Gervet, Thibaut Lavril, Thomas Wang, Timoth  e Lacroix, and William El Sayed. Mixtral of experts. *CoRR*, abs/2401.04088, 2024. doi: 10.48550/ARXIV.2401.04088. URL <https://doi.org/10.48550/arXiv.2401.04088>.
- Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention. In Jason Flinn, Margo I. Seltzer, Peter Druschel, Antoine Kaufmann, and Jonathan Mace (eds.), *Proceedings of the 29th Symposium on Operating Systems Principles, SOSP 2023, Koblenz, Germany, October 23-26, 2023*, pp. 611–626. ACM, 2023. doi: 10.1145/3600006.3613165. URL <https://doi.org/10.1145/3600006.3613165>.
- NVIDIA. Nvidia nvswitch: Technical overview. Technical report, NVIDIA, 2018. URL <https://images.nvidia.com/content/pdf/nvswitch-technical-overview.pdf>.
- NVIDIA. cuBLAS, 2022. URL <https://developer.nvidia.com/cublas>.
- Nvidia. Cutlass, 2022. URL <https://github.com/NVIDIA/cutlass>.
- NVIDIA. Nvidia collective communications library. <https://developer.nvidia.com/nccl>, 2024.
- NVIDIA. NVSHMEM, 2025. URL <https://docs.nvidia.com/nvshmem/api/using.html>.
- OpenAI. GPT-4 technical report. *CoRR*, abs/2303.08774, 2023. doi: 10.48550/ARXIV.2303.08774. URL <https://doi.org/10.48550/arXiv.2303.08774>.
- Qwen-Team. Qwen2.5 technical report. *arXiv preprint arXiv:2412.15115*, 2024.
- Morgane Riv  re, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, L  onard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ram  , Johan Ferret, Peter Liu, Pouya Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, Anton Tsitsulin, Nino Vieillard, Piotr Stanczyk, Sertan Girgin, Nikola Momchev, Matt Hoffman, Shantanu Thakoor, Jean-Bastien Grill, Behnam Neyshabur, Olivier Bachem, Alanna Walton, Aliaksei Severyn, Alicia Parrish, Aliya Ahmad, Allen Hutchison, Alvin Abdagic, Amanda Carl, Amy Shen, Andy Brock, Andy Coenen, Anthony Laforge, Antonia Paterson, Ben Bastian, Bilal Piot, Bo Wu, Brandon Royal, Charlie Chen, Chintu Kumar, Chris Perry, Chris Welty, Christopher A. Choquette-Choo, Danila Sinopalnikov, David Weinberger, Dimple Vijaykumar, Dominika Rogozinska, Dustin Herbison, Elisa Bandy, Emma Wang, Eric Noland, Erica Moreira, Evan Senter, Evgenii Eltyshev, Francesco Visin, Gabriel Rasskin, Gary Wei, Glenn Cameron, Gus Martins, Hadi Hashemi, Hanna Klimczak-Plucinska, Harleen Batra, Harsh Dhand, Ivan Nardini, Jacinda Mein, Jack Zhou, James Svensson, Jeff Stanway, Jetha Chan, Jin Peng Zhou, Joana Carrasqueira, Joana Iljazi, Jocelyn Becker, Joe Fernandez, Joost van Amersfoort, Josh Gordon, Josh Lipschultz, Josh Newlan, Ju-yeong Ji, Kareem Mohamed, Kartikeya Badola, Kat Black, Katie Millican, Keelin McDonell, Kelvin Nguyen, Kiranbir Sodhia, Kish Greene, Lars Lowe S  sund, Lauren Usui, Laurent Sifre, Lena Heuermann, Leticia Lago, and Lilly McNealus. Gemma 2: Improving open language models at a practical size. *CoRR*, abs/2408.00118, 2024. doi: 10.48550/ARXIV.2408.00118. URL <https://doi.org/10.48550/arXiv.2408.00118>.

- Aashaka Shah, Abhinav Jangda, Binyang Li, Caio Rocha, Changho Hwang, Jithin Jose, Madan Musuvathi, Olli Saarikivi, Peng Cheng, Qinghua Zhou, Roshan Dathathri, Saeed Maleki, and Ziyue Yang. Msccl++: Rethinking gpu communication abstractions for cutting-edge ai applications, 2025. URL <https://arxiv.org/abs/2504.09014>.
- Benjamin Spector, Simran Arora, Aaryan Singhal, Daniel Y. Fu, and Christopher Ré. Thunderkitens: Simple, fast, and adorable AI kernels. *CoRR*, abs/2410.20399, 2024. doi: 10.48550/ARXIV.2410.20399. URL <https://doi.org/10.48550/arXiv.2410.20399>.
- TileLang-Team. Tilelang, 2025. URL <https://github.com/tile-ai/tilelang>.
- Philippe Tillet, Hsiang-Tsung Kung, and David D. Cox. Triton: an intermediate language and compiler for tiled neural network computations. In Tim Mattson, Abdullah Muzahid, and Armando Solar-Lezama (eds.), *Proceedings of the 3rd ACM SIGPLAN International Workshop on Machine Learning and Programming Languages, MAPL@PLDI 2019, Phoenix, AZ, USA, June 22, 2019*, pp. 10–19. ACM, 2019. doi: 10.1145/3315508.3329973. URL <https://doi.org/10.1145/3315508.3329973>.
- Shibo Wang, Jinliang Wei, Amit Sabne, Andy Davis, Berkin Ilbeyi, Blake Hechtman, Dehao Chen, Karthik Srinivasa Murthy, Marcello Maggioni, Qiao Zhang, Sameer Kumar, Tongfei Guo, Yuanzhong Xu, and Zongwei Zhou. Overlap communication with dependent computation via decomposition in large deep learning models. In Tor M. Aamodt, Natalie D. Enright Jerger, and Michael M. Swift (eds.), *Proceedings of the 28th ACM International Conference on Architectural Support for Programming Languages and Operating Systems, Volume 1, ASPLOS 2023, Vancouver, BC, Canada, March 25-29, 2023*, pp. 93–106. ACM, 2023. doi: 10.1145/3567955.3567959. URL <https://doi.org/10.1145/3567955.3567959>.
- Mengdi Wu, Xinhao Cheng, Shengyu Liu, Chunan Shi, Jianan Ji, Kit Ao, Praveen Velliengiri, Xupeng Miao, Oded Padon, and Zhihao Jia. Mirage: A multi-level superoptimizer for tensor programs. In *19th USENIX Symposium on Operating Systems Design and Implementation (OSDI 25)*, Boston, MA, July 2025. USENIX Association.
- An Yang, Anfeng Li, Baosong Yang, Beichen Zhang, Binyuan Hui, Bo Zheng, Bowen Yu, Chang Gao, Chengen Huang, Chenxu Lv, Chujie Zheng, Dayiheng Liu, Fan Zhou, Fei Huang, Feng Hu, Hao Ge, Haoran Wei, Huan Lin, Jialong Tang, Jian Yang, Jianhong Tu, Jianwei Zhang, Jian Yang, Jiayi Yang, Jingren Zhou, Jingren Zhou, Junyang Lin, Kai Dang, Keqin Bao, Kexin Yang, Le Yu, Lianghao Deng, Mei Li, Mingfeng Xue, Mingze Li, Pei Zhang, Peng Wang, Qin Zhu, Rui Men, Ruize Gao, Shixuan Liu, Shuang Luo, Tianhao Li, Tianyi Tang, Wenbiao Yin, Xingzhang Ren, Xinyu Wang, Xinyu Zhang, Xuancheng Ren, Yang Fan, Yang Su, Yichang Zhang, Yinger Zhang, Yu Wan, Yuqiong Liu, Zekun Wang, Zeyu Cui, Zhenru Zhang, Zhipeng Zhou, and Zihan Qiu. Qwen3 technical report. *CoRR*, abs/2505.09388, 2025. doi: 10.48550/ARXIV.2505.09388. URL <https://doi.org/10.48550/arXiv.2505.09388>.
- Shulai Zhang, Ningxin Zheng, Haibin Lin, Ziheng Jiang, Wenlei Bao, Chengquan Jiang, Qi Hou, Weihao Cui, Size Zheng, Li-Wen Chang, Quan Chen, and Xin Liu. Comet: Fine-grained computation-communication overlapping for mixture-of-experts. *CoRR*, abs/2502.19811, 2025. doi: 10.48550/ARXIV.2502.19811. URL <https://doi.org/10.48550/arXiv.2502.19811>.
- Chenggang Zhao, Shangyan Zhou, Liyue Zhang, Chengqi Deng, Zhean Xu, Yuxuan Liu, Kuai Yu, Jiashi Li, and Liang Zhao. DeepEP: an efficient expert-parallel communication library. <https://github.com/deepseek-ai/DeepEP>, 2025.
- Size Zheng, Jin Fang, Xuegui Zheng, Qi Hou, Wenlei Bao, Ningxin Zheng, Ziheng Jiang, Dongyang Wang, Jianxi Ye, Haibin Lin, Li-Wen Chang, and Xin Liu. Tilelink: Generating efficient compute-communication overlapping kernels using tile-centric primitives, 2025. URL <https://arxiv.org/abs/2503.20313>.

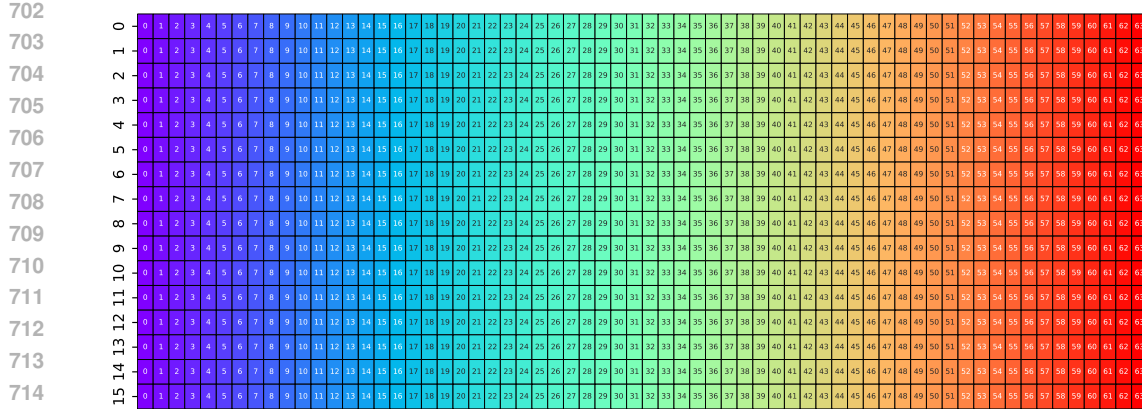


Figure 11: Global view of ranks and data tiles without swizzling. Corresponding to Figure 4. Each row corresponds to one rank; each column corresponds to one tile. Each rank computes GEMM from the left tile to the right tile. This example is for GEMM+ReduceScatter.

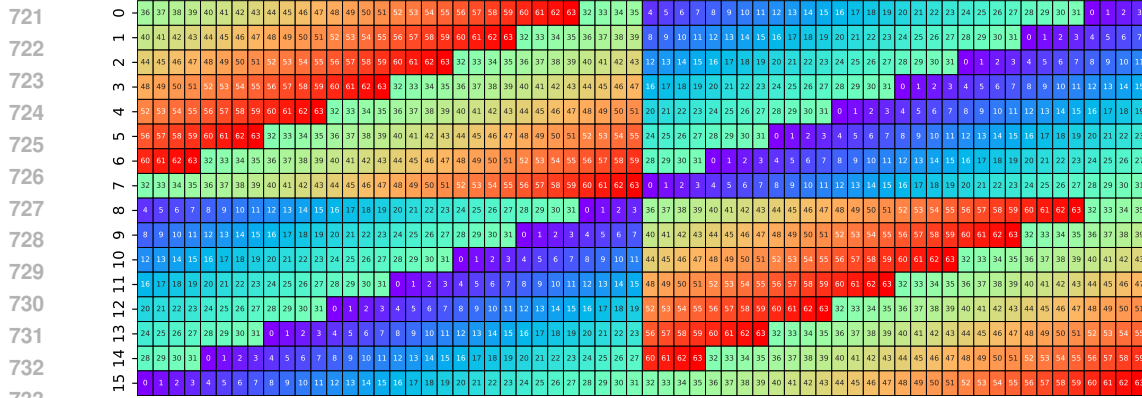


Figure 12: Global view of ranks and data tiles with swizzling. Corresponding to Figure 4. This example is for GEMM+ReduceScatter.

A USE OF LARGE LANGUAGE MODELS

In preparing this manuscript, LLM were used as a general-purpose writing assistance tool, primarily to improve the clarity and phrasing of certain sentences. All authors have reviewed and edited the final manuscript and take full responsibility for the entire content of the paper.

B SWIZZLE

Swizzling optimization has been mentioned in previous work (Jangda et al., 2022; Chang et al., 2024; Zheng et al., 2025). But their swizzling methods are limited to single node and there lack a clear explanation of the swizzling for non-perfect tiling scenarios (which is the most common for LLM training and inference).

Figure 11 and Figure 12 show an enlarged view of Figure 4, which correspond to swizzle optimization for 16 GPUs GEMM+ReduceScatter with perfect shapes (no cross-rank tiles). For this case, we use $M_{per_rank} = 1024$, $N_{per_rank} = 256$, $tile_M = tile_N = 256$.

For non-perfect shapes, some tiles are needed by multiple ranks for subsequent ReduceScatter. These tiles should be permuted to the left so that they can be computed earlier and transferred

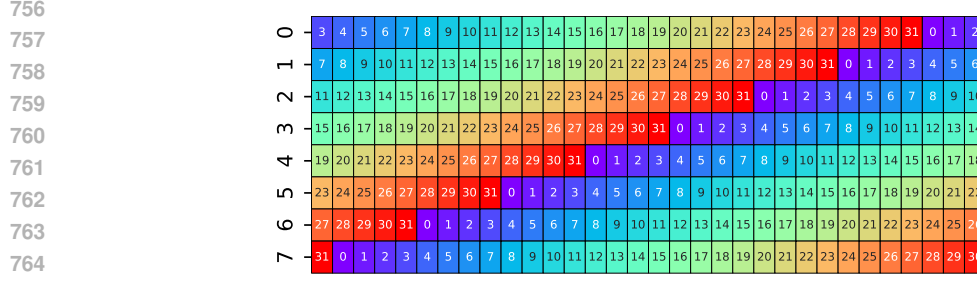


Figure 13: Single node global view of swizzling for non-perfect tiling. The tile orders are different from that of perfect tiling. Cross-rank tiles are permuted to left. This example is for GEMM+ReduceScatter.

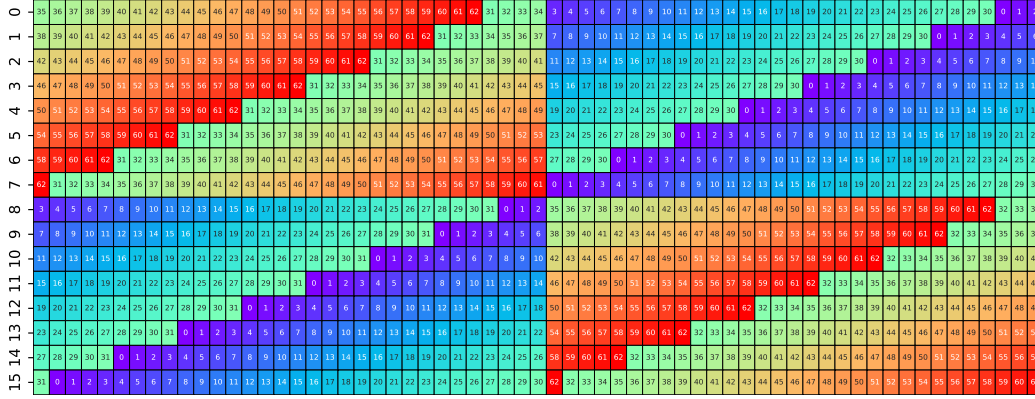


Figure 14: Two nodes global view of swizzling for non-perfect tiling. This example is for GEMM+ReduceScatter.

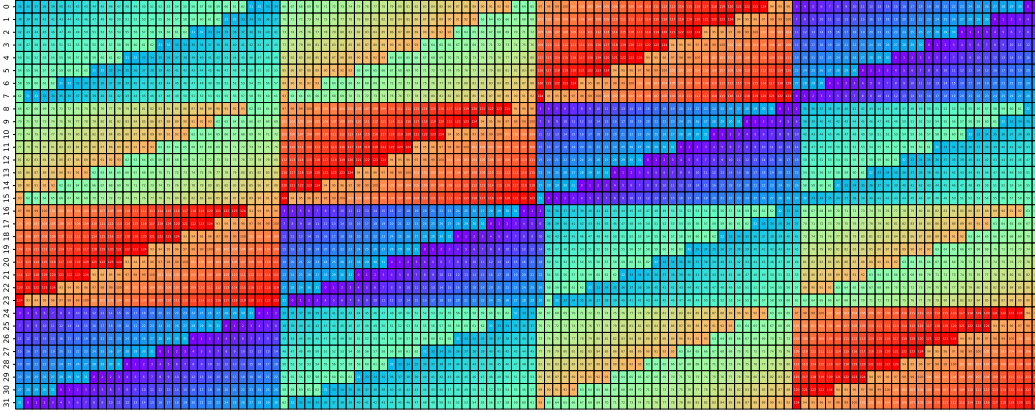


Figure 15: Four nodes global view of swizzling for non-perfect tiling. This example is for GEMM+ReduceScatter.

to the ranks that need them. Figure 13 shows swizzle results for 8 GPUs; Figure 14 shows swizzling for 16 GPUs; and Figure 15 shows swizzling for 32 GPUs. For this case, we use $M_{per_rank} = 997$, $N_{per_rank} = 256$, $tile_M = tile_N = 256$.

Besides GEMM+ReduceScatter, we also implement similar swizzling for AllGather+GEMM, and AllGather+MoE. The swizzle logic is implemented in DITRON through efficient kernels. We show the core code for swizzle as follows.

First, swizzle code for AllGather+GEMM. DITRON is implemented based on Triton, so we still use *triton.jit* as compile entry. We extend programming interface with multiple primitives for tile-level or even thread-level control, such as *laneid*, *warp_prefix_sum_kernel*, *__ballot_sync*, *__shfl_sync_i32*.

```

1  @triton.jit(do_not_specialize=["rank"])
2  def threadblock_swizzle_allgather_gemm_kernel(
3      tiled_m,
4      M,
5      rank,
6      WORLD_SIZE: tl.constexpr,
7      NNODES: tl.constexpr,
8      BLOCK_SIZE_M: tl.constexpr,
9      DEBUG: tl.constexpr = False,
10 ):
11     LOCAL_WORLD_SIZE = WORLD_SIZE // NNODES
12     node_id = rank // LOCAL_WORLD_SIZE
13     M_per_rank = M // WORLD_SIZE
14     M_per_node = M // NNODES
15     node_start = node_id
16
17     lane_id = laneid()
18
19     if lane_id < NNODES:
20         n = (lane_id + node_start) % NNODES
21         M_node_start = M_per_node * n
22         M_node_end = M_per_node * (n + 1)
23         tiled_m_node_start = M_node_start // BLOCK_SIZE_M
24         prev_tiled_m_node_end = (M_node_start - 1) // BLOCK_SIZE_M
25         tiled_m_node_end = (M_node_end - 1) // BLOCK_SIZE_M
26         next_tiled_m_node_start = M_node_end // BLOCK_SIZE_M
27
28         if lane_id == 0 and M_node_start != 0:
29             if prev_tiled_m_node_end == tiled_m_node_start:
30                 tiled_m_node_start += 1
31
32         if lane_id == 0 and M_node_end != M:
33             if next_tiled_m_node_start == tiled_m_node_end:
34                 tiled_m_node_end -= 1
35
36         if lane_id != NNODES - 1 and M_node_end != M:
37             if next_tiled_m_node_start == tiled_m_node_end:
38                 tiled_m_node_end -= 1
39         if DEBUG and lane_id == NNODES - 1:
40             print("tiled_m_node_end", tiled_m_node_end)
41
42         swizzled_tiled_m_size = tiled_m_node_end - tiled_m_node_start + 1
43     else:
44         swizzled_tiled_m_size = 0
45
46     if DEBUG and lane_id < NNODES:
47         print("swizzled_tiled_m_size", swizzled_tiled_m_size, lane_id)
48     swizzled_tiled_m_size_accum = (warp_prefix_sum_kernel(swizzled_tiled_m_size, lane_id, NNODES) -
49                                   swizzled_tiled_m_size)
50     if DEBUG and lane_id < NNODES:
51         print("swizzled_tiled_m_size_accum", swizzled_tiled_m_size_accum)
52
53     tiled_m_size_l = __shfl_down_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size, NNODES - node_start)
54     tiled_m_size_r = __shfl_up_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size, node_start)
55     tiled_m_size = 0
56     if lane_id < node_start:
57         tiled_m_size = tiled_m_size_l
58     elif lane_id < NNODES:
59         tiled_m_size = tiled_m_size_r
60
61     if DEBUG and lane_id < NNODES:
62         print("tiled_m_size", tiled_m_size)
63
64     tiled_m_size_accum = warp_prefix_sum_kernel(tiled_m_size, lane_id, NNODES) - tiled_m_size
65     mask = __ballot_sync(0xFFFFFFFF, tiled_m < swizzled_tiled_m_size_accum)
66     n = ffs(mask) - 1 - 1
67     if DEBUG and lane_id < NNODES:
68         print("tiled_m_size_accum", tiled_m_size_accum)
69         print("n", n, tiled_m, swizzled_tiled_m_size_accum, mask)
70
71     # map node
72     nid = (n + node_start) % NNODES
73     node_offset = __shfl_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size_accum, n)
74
75     tile_size = __shfl_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size, n)
76
77     tiled_m_intra_node = tiled_m - node_offset
78     local_rank = rank % LOCAL_WORLD_SIZE
79     m_start = M_per_node * nid + M_per_rank * local_rank
80     tiled_m_start = tl.cdiv(m_start, BLOCK_SIZE_M)
81     swizzled_node_offset = __shfl_sync_i32(0xFFFFFFFF, tiled_m_size_accum, nid)

```

```

82     rank_offset = max(0, tiled_m_start - swizzled_node_offset) # this may < 0, bad
83
84     # map rank
85     tiled_m_intra_node_new = (tiled_m_intra_node + rank_offset) % tile_size
86     return swizzled_node_offset + tiled_m_intra_node_new

```

Then, swizzle code for GEMM+ReduceScatter.

```

1  @triton.jit(do_not_specialize=["rank"])
2  def threadblock_swizzle_gemm_reduce_scatter_kernel(
3      tiled_m,
4      M,
5      rank,
6      WORLD_SIZE: tl.constexpr,
7      NNODES: tl.constexpr,
8      BLOCK_SIZE_M: tl.constexpr,
9      DEBUG: tl.constexpr = False,
10 ):
11     LOCAL_WORLD_SIZE = WORLD_SIZE // NNODES
12     node_id = rank // LOCAL_WORLD_SIZE
13     M_per_rank = M // WORLD_SIZE
14     M_per_node = M // NNODES
15     node_start = node_id + 1
16
17     lane_id = laneid()
18
19     if lane_id < NNODES:
20         n = (lane_id + node_start) % NNODES
21         M_node_start = M_per_node * n
22         M_node_end = M_per_node * (n + 1)
23         tiled_m_node_start = M_node_start // BLOCK_SIZE_M
24         # if tiled_m_start_node overlaps with previous node, then we need to add 1 to
25         # tiled_m_start_node
26         prev_tiled_m_node_end = (M_node_start - 1) // BLOCK_SIZE_M
27         if lane_id != 0 and M_node_start != 0:
28             if prev_tiled_m_node_end == tiled_m_node_start:
29                 tiled_m_node_start += 1
30
31         tiled_m_node_end = (M_node_end - 1) // BLOCK_SIZE_M
32         next_tiled_m_node_start = M_node_end // BLOCK_SIZE_M
33         if lane_id == NNODES - 1 and M_node_end != M:
34             if next_tiled_m_node_start == tiled_m_node_end:
35                 tiled_m_node_end -= 1
36
37         swizzled_tiled_m_size = tiled_m_node_end - tiled_m_node_start + 1
38     else:
39         swizzled_tiled_m_size = 0
40
41     if DEBUG and lane_id < NNODES:
42         print("swizzled_tiled_m_size", swizzled_tiled_m_size, lane_id)
43     swizzled_tiled_m_size_accu = (warp_prefix_sum_kernel(swizzled_tiled_m_size, lane_id, NNODES) -
44                                swizzled_tiled_m_size)
45     if DEBUG and lane_id < NNODES:
46         print("swizzled_tiled_m_size_accu", swizzled_tiled_m_size_accu)
47     # thread 0 hold node 'node_start' size
48     # thread 1 hold node 'node_start + 1' size
49     # ...
50     # thread NNODES - node_start hold node '0' size
51     # thread NNODES - 1 hold NNODES 'node_start - 1' size
52     # => thread 0 want to hold node 0 size
53
54     tiled_m_size_l = __shfl_down_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size, NNODES - node_start)
55     tiled_m_size_r = __shfl_up_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size, node_start)
56     tiled_m_size = 0
57     if lane_id < node_start:
58         tiled_m_size = tiled_m_size_l
59     elif lane_id < NNODES:
60         tiled_m_size = tiled_m_size_r
61
62     if DEBUG and lane_id < NNODES:
63         print("tiled_m_size", tiled_m_size)
64
65     tiled_m_size_accu = warp_prefix_sum_kernel(tiled_m_size, lane_id, NNODES) - tiled_m_size
66     mask = __ballot_sync(0xFFFFFFFF, tiled_m < swizzled_tiled_m_size_accu)
67     n = ffs(mask) - 1 - 1
68     if DEBUG and lane_id < NNODES + 1:
69         print("tiled_m_size_accu", tiled_m_size_accu)
70         print("n", n, tiled_m, swizzled_tiled_m_size_accu, mask)
71
72     # map node
73     nid = (n + node_start) % NNODES
74     node_offset = __shfl_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size_accu, n)
75
76     tile_size = __shfl_sync_i32(0xFFFFFFFF, swizzled_tiled_m_size, n)
77
78     tiled_m_intra_node = tiled_m - node_offset
79     local_rank = rank % LOCAL_WORLD_SIZE
80     m_start = M_per_node * nid + M_per_rank * (local_rank + 1)
81     tiled_m_start = m_start // BLOCK_SIZE_M

```

```

82 swizzled_node_offset = __shfl_sync_i32(0xFFFFFFFF, tiled_m_size_accum, nid)
83 rank_offset = max(0, tiled_m_start - swizzled_node_offset) # this may < 0, bad
84
85 # map rank
86 tiled_m_intra_node_new = (tiled_m_intra_node + rank_offset) % tile_size
87 return swizzled_node_offset + tiled_m_intra_node_new

```

At last, the swizzle code for AllGather+MoE.

```

1 @triton.jit(do_not_specialize=["rank"])
2 def threadblock_swizzle_ag_moe_kernel(
3     # input
4     ntokens_by_rank_by_expert_ptr,
5     # output
6     expert_id_ptr,
7     tiled_m_ptr,
8     segment_start_ptr,
9     segment_end_ptr,
10    ntiles_ptr,
11    # workspace buffer
12    ntokens_by_expert_by_rank_acc_ptr,
13    ntiles_by_expert_acc_ptr,
14    ntiles_by_expert_by_stage_ptr,
15    ntiles_by_expert_by_stage_acc_ptr,
16    rank,
17    N_EXPERTS: tl.constexpr,
18    TP_SIZE: tl.constexpr,
19    LOCAL_TP_SIZE: tl.constexpr,
20    NTILES_NEXT_POW_OF_2: tl.constexpr,
21    BLOCK_SIZE_M: tl.constexpr,
22    DEBUG: tl.constexpr = False,
23 ):
24     """
25     tile_index = g(expert_id, stage, index): if tile_index is grouped by expert then stage
26     but how to map (stage, index) => tile_index_in_expert?
27     first map tile_index_in_expert => (stage, off_in_expert_by_stage, index_in_expert_in_stage)
28     tile_index => (expert_id, tile_index_in_expert) : ntokens grouped by expert_id by
29     ↪ rank
30     => (expert_id, stage) : get_stage_id
31     => (expert_id, off_in_expert_by_stage) : cumsum by stage
32     => (expert_id, off_in_expert_by_stage, index_in_expert_in_stage) : atomic_add for
33     ↪ index_in_expert_in_stage
34     => tile_index_new : by function g
35
36     tile_index -> (expert_id, rank_offset) grouped by (expert_id, rank)
37     rank_offset -> segment_start, segment_end, stage
38
39     so we can remap tile_index as:
40     tile_index_new = g(f(tile_index))
41     """
42
43     thread_idx = tid(0)
44     N_EXPERTS_NEXT_POW_OF_2: tl.constexpr = next_power_of_2(N_EXPERTS)
45     TP_SIZE_NEXT_POW_OF_2: tl.constexpr = next_power_of_2(TP_SIZE)
46     offs_by_expert = tl.arange(0, N_EXPERTS_NEXT_POW_OF_2)
47     mask_by_expert = offs_by_expert < N_EXPERTS
48     offs_by_rank = tl.arange(0, TP_SIZE_NEXT_POW_OF_2)
49     mask_by_rank = (offs_by_rank < TP_SIZE)
50     offs_by_expert_by_rank = offs_by_expert[:, None] * TP_SIZE + offs_by_rank[None, :]
51     mask_by_expert_by_rank = mask_by_expert[:, None] & mask_by_rank[None, :]
52     offs_by_rank_by_expert = offs_by_rank[:, None] * N_EXPERTS + offs_by_expert[None, :]
53     mask_by_rank_by_expert = mask_by_rank[:, None] & mask_by_expert[None, :]
54     ntokens_by_rank_by_expert = tl.load(ntokens_by_rank_by_expert_ptr + offs_by_rank_by_expert,
55                                         mask=mask_by_rank_by_expert)
56
57     ntokens_by_expert_by_rank = ntokens_by_rank_by_expert.T
58     ntokens_by_expert_by_rank_acc = tl.cumsum(ntokens_by_expert_by_rank, axis=1)
59     ntokens_by_expert = tl.sum(ntokens_by_rank_by_expert, axis=0)
60     ntiles_by_expert = tl.cdiv(ntokens_by_expert, BLOCK_SIZE_M)
61     ntiles_by_expert_acc = tl.cumsum(ntiles_by_expert, axis=0)
62     ntiles = tl.sum(ntiles_by_expert)
63
64     tl.store(ntokens_by_expert_by_rank_acc_ptr + offs_by_expert_by_rank,
65             ↪ ntokens_by_expert_by_rank_acc,
66             mask=mask_by_expert_by_rank)
67     tl.store(ntiles_by_expert_acc_ptr + offs_by_expert, ntiles_by_expert_acc)
68
69     # # for each tiled_m in expert eid => stage id / segment_start / segment_end / tiled_m
70     tile_index = tl.arange(0, NTILES_NEXT_POW_OF_2)
71     mask_tile_idx = tile_index < ntiles
72
73     __syncthreads()
74     # tile_index -> (expert_id, offset_by_expert, tile_index_in_expert) -> stage, segment_start,
75     ↪ segment_end
76     expert_id, off_by_expert, off_in_expert =
77     ↪ bisect_right_with_offset_kernel(ntiles_by_expert_acc_ptr, tile_index,
78                                     N_EXPERTS)

```

```

77     stage, segment_start, segment_end = get_tile_stage(
78         off_in_expert,
79         rank,
974         TP_SIZE,
81         LOCAL_TP_SIZE,
82         BLOCK_SIZE_M,
83         ntokens_by_expert_by_rank_acc_ptr + expert_id * TP_SIZE,
84     )
977
85     # histogram by expert by stage
86     tl.store(ntiles_by_expert_by_stage_ptr + offs_by_expert_by_rank, 0, mask=mask_by_expert_by_rank)
978
87     __syncthreads()
88     off_in_expert_in_stage = tl.atomic_add(ntiles_by_expert_by_stage_ptr + expert_id * TP_SIZE +
979     ↪ stage, 1,
980                                         sem="relaxed", scope="gpu", mask=mask_tile_idx)
981
982     __syncthreads()
983
984     # do some cumsum
985     ntiles_by_expert_by_stage = tl.load(ntiles_by_expert_by_stage_ptr + offs_by_expert_by_rank,
986                                         mask=mask_by_expert_by_rank, other=0)
987     ntiles_by_expert_by_stage_acc = tl.cumsum(ntiles_by_expert_by_stage, axis=1)
988     __syncthreads()
989     tl.store(
990         ntiles_by_expert_by_stage_acc_ptr + offs_by_expert_by_rank,
991         ntiles_by_expert_by_stage_acc,
992         mask=mask_by_expert_by_rank,
993     )
994     __syncthreads()
995
996     off_in_expert_by_stage = tl.where(
997         stage == 0,
998         0,
999         tl.load(ntiles_by_expert_by_stage_acc_ptr + expert_id * TP_SIZE + stage - 1,
1000         ↪ mask=mask_tile_idx, other=0),
1001     )
1002
1003     tile_index_by_expert_by_stage = off_by_expert + off_in_expert_by_stage + off_in_expert_in_stage
1004     __syncthreads()
1005
1006     tl.store(expert_id_ptr + tile_index_by_expert_by_stage, expert_id, mask=mask_tile_idx)
1007     tl.store(tiled_m_ptr + tile_index_by_expert_by_stage, tile_index, mask=mask_tile_idx)
1008     tl.store(segment_start_ptr + tile_index_by_expert_by_stage, segment_start, mask=mask_tile_idx)
1009     tl.store(segment_end_ptr + tile_index_by_expert_by_stage, segment_end, mask=mask_tile_idx)
1010     thread_idx = tid(0)
1011     if thread_idx == 0:
1012         tl.store(ntiles_ptr, ntiles)
1013     if DEBUG and thread_idx < ntiles:
1014         print("expert_id", expert_id)
1015

```

C EVALUATION SETUP AND WORKLOAD CONFIGURATIONS

C.1 EVALUATION SETUP

The H800 clusters used in our evaluation are equipped with NVLink of uni-directional bandwidth of 200 GB/s. There are 8 NICs in one node and each GPU can communicate with another GPU in other nodes at a uni-directional bandwidth of 50 GB/s. The H20 clusters used in our evaluation are equipped with NVLink of uni-directional bandwidth of 450 GB/s. There are 4 NICs in one node and each GPU can communicate with another GPU in other nodes at a uni-directional bandwidth of 25 GB/s.

Table 1: Shapes for AllGather+GEMM on H800.

No.	From Model	M	N	K
1	LLaMA3-7B	8192	11008	4096
2	LLaMA3.1-8B	8192	14336	4096
3	LLaMA3.1-70B	8192	28672	8192
4	LLaMA3.1-405B	8192	53248	16384
5	Mistral-7B	8192	14336	4096
6	Qwen2-72B	8192	29568	8192

C.2 WORKLOAD CONFIGURATION

The input shapes for AllGather+GEMM used in Figure 5 are listed in Table 1. The input shapes for GEMM+ReduceScatter and GEMM+AllReduce used in Figure 5 are listed in Table 2. The input shapes for AllGather+MoE and MoE+AllReduce used in Figure 5 are listed in Table 3.

Table 2: Shapes for GEMM+ReduceScatter and GEMM+AllReduce on H800.

No.	From Model	M	N	K
1	LLaMA3-7B	8192	4096	11008
2	LLaMA3.1-8B	8192	4096	14336
3	LLaMA3.1-70B	8192	8192	28672
4	LLaMA3.1-405B	8192	16384	53248
5	Mistral-7B	8192	4096	14336
6	Qwen2-72B	8192	8192	29568

Table 3: Shapes for AllGather+MoE and MoE+AllReduce on H800.

No.	From Model	num_tokens	hidden_dim	intermediate_size	num_experts	topk
1	Qwen1.5-MoE-A2.7B	8192	2048	1408	60	4
2	Mixtral-8x7B	8192	14336	4096	8	2
3	Mixtral-8x22B	8192	16384	6144	8	2
4	DeepSeek-MoE	8192	1408	2048	64	6

Table 4: Shapes for strong scaling AllGather+GEMM.

Label	From Model	M	N	K
7B	LLaMA3-7B	32768	11008	4096
8B	LLaMA3.1-8B	32768	14336	4096
32B	Qwen3-32B	32768	25600	5120
36B	Seed-OSS-36B	32768	27648	5120
70B	LLaMA3.1-70B	32768	28672	8192
72B	Qwen2-72B	32768	29568	8192
175B	GPT-3-175B	32768	49152	12288
405B	LLaMA3.1-405B	32768	53248	16384

Table 5: Shapes for strong scaling GEMM+ReduceScatter.

1	From Model	M	N	K
7B	LLaMA3-7B	32768	4096	11008
8B	LLaMA3.1-8B	32768	4096	14336
32B	Qwen3-32B	32768	5120	25600
36B	Seed-OSS-36B	32768	5120	27648
70B	LLaMA3.1-70B	32768	8192	28672
72B	Qwen2-72B	32768	8192	29568
175B	GPT-3-175B	32768	12288	49152
405B	LLaMA3.1-405B	32768	16384	53248

Table 6: Shapes for strong scaling GEMM+AllToAll.

Label	seqlen	heads	head_size
case 1	262144	64	128
case 2	491520	64	128
case 3	262144	32	128
case 4	491520	32	128
case 5	262144	16	128
case 6	491520	16	128

For scaling experiments, the shapes for strong scaling AllGather+GEMM are shown in Table 4. The shapes remain the same from 8 GPUs to 32 GPUs. The shapes for weak scaling are the same as those in strong scaling except that the value of M varies with the number of GPUs. We use $M = 8192$ for 8 GPUs, $M = 16384$ for 16 GPUs, and $M = 32768$ for 32 GPUs. The shapes for strong scaling GEMM+ReduceScatter are shown in Table 5. The shapes remain the same for 8 GPUs and 32 GPUs. For weak scaling, the configurations for M are the same as weak scaling AllGather+GEMM.

The shapes for strong scaling GEMM+AllToAll are shown in Table 6. The shapes remain the same from 8 GPUs to 128 GPUs. The shapes for weak scaling are the same as those in strong scaling except that the seqlen varies with the number of GPUs. For 8 GPUs, seqLens are 65536 and 122880; for 16 GPUs, they are 131072 and 245760; for 32 GPUs, they are 262144 and 491520; for 64 GPUs, they are 524288 and 983040; for 128 GPUs, they are 1048576 and 1966080.

D DETAILED EXPERIMENT RESULTS

D.1 DETAILED EVALUATION RESULTS FOR MOE

The detailed results for AllGather+MoE in Figure 5 are shown in Table 7. The detailed results for MoE+AllReduce in Figure 5 are shown in Table 8. The shapes for each case are shown in Table 3.

Table 7: Latency of AllGather+MoE (ms).

	CuBLAS+NCCL	TileLink	COMET	DITRON
1	18.12	0.56	0.33	0.32
2	7.12	2.14	1.48	1.51
3	8.49	3.30	1.85	1.72
4	24.27	0.65	0.28	0.24

Table 8: Latency of MoE+AllReduce (ms).

	CuBLAS+NCCL	TileLink	DITRON
1	22.04	1.44	0.445
2	7.99	3.70	2.603
3	8.97	4.32	3.289
4	28.64	0.84	0.441

Table 9: Strong scaling results (speedup to CuBLAS+NCCL) for AllGather+GEMM on H800.

From Model	1 Node	2 Node	4 Node
LLaMA3-7B	1.31	1.11	0.63
LLaMA3.1-8B	1.35	1.12	0.61
Qwen3-32B	1.59	1.32	0.65
Seed-OSS-36B	1.64	1.32	0.64
LLaMA3.1-70B	1.67	1.36	0.77
Qwen2-72B	1.67	1.39	1.03
GPT-3-175B	1.46	1.68	1.00
LLaMA3.1-405B	1.44	1.76	1.02

D.2 DETAILED EVALUATION RESULTS FOR SCALING

In Figure 8 we have shown strong scaling results for AllGather+GEMM on H20. Here we show strong scaling results on H800 and weak scaling results on both H20 and H800 in Table 9, Table 10, and Table 11.

Table 10: Weak scaling results (speedup to CuBLAS+NCCL) for AllGather+GEMM on H20

From Model	1 Node	2 Node	4 Node
LLaMA3-7B	1.23	1.50	0.80
LLaMA3.1-8B	1.19	1.50	0.86
Qwen3-32B	1.06	1.22	0.89
Seed-OSS-36B	1.05	1.26	0.89
LLaMA3.1-70B	1.10	1.29	0.93
Qwen2-72B	1.01	1.18	1.36
GPT-3-175B	1.05	1.23	1.09
LLaMA3.1-405B	1.05	1.19	1.08

Table 11: Weak scaling results (speedup to CuBLAS+NCCL) for AllGather+GEMM on H800

From Model	1 Node	2 Node	4 Node
LLaMA3-7B	1.30	0.98	0.63
LLaMA3.1-8B	1.30	1.02	0.61
Qwen3-32B	1.53	1.14	0.65
Seed-OSS-36B	1.57	1.15	0.64
LLaMA3.1-70B	1.61	1.25	0.77
Qwen2-72B	1.66	1.29	1.03
GPT-3-175B	1.43	1.58	1.00
LLaMA3.1-405B	1.34	1.63	1.02

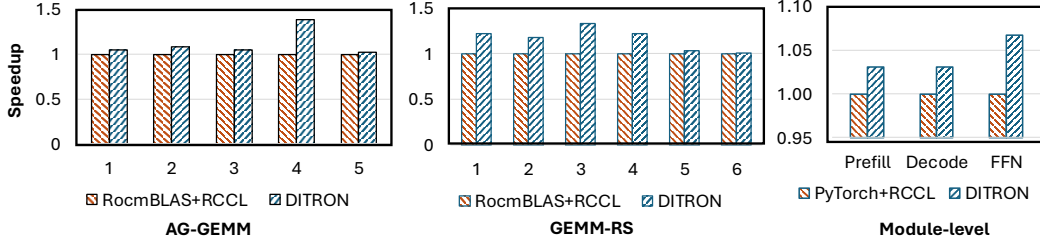


Figure 16: Evaluation results on 8x AMD MI308X GPUs. AllGather+GEMM and GEMM+ReduceScatter shapes are in Table 1 and Table 2. Module-level evaluation uses configurations from Qwen3-32B.

D.3 EVALUATION RESULTS ON AMD GPU

DITRON can be deployed to AMD GPUs through RocSHMEM (AMD) and Triton for AMD. For evaluation, we use shapes from Table 1 and Table 2. Our baseline is RocmBLAS+RCCL. Overall, the speedup ranges from $1.02\times$ to $1.38\times$. The geometric speedup to baseline is $1.11\times$ for AllGather+GEMM and $1.16\times$ to GEMM+ReduceScatter. For module-level evaluation, we use Qwen3-32B. The speedup for prefill attention is $1.03\times$, the speedup for decode attention is $1.03\times$, and the speedup for FFN is $1.07\times$.

D.4 EVALUATION RESULTS ON L20 GPU

We evaluate 6 different workloads on L20, including GEMM+ReduceScatter (Table 12), AllToAll (Table 13), AllGather+MoE (Table 14), MoE+ReduceScatter (Table 15), MoE+AllReduce (Table 16), and AllToAll+GEMM (Table 17). The geometric mean speedup of all the workloads is $8.33\times$.

Table 12: GEMM+RS performance on 8×L20 GPUs.

M	N	K	CuBLAS+NCCL	DITRON	Speedup
8192	8192	29568	18.63 ms	7.34 ms	2.54×

Table 13: All2All performance on 8×L20 GPU.

num_tokens	hidden_size	num_experts	topk	NCCL	DITRON	speedup
8	3584	128	8	0.42 ms	0.09 ms	4.51×

Table 14: AG+MoE performance on 8×L20 GPU (shape configurations see Table 3).

No.	num_tokens	CuBLAS+NCCL	DITRON	speedup
1	2048	21.10 ms	0.58 ms	36.26×
2	2048	43.91 ms	2.67 ms	16.47×
3	2048	51.11 ms	3.44 ms	14.86×
4	2048	21.43 ms	0.43 ms	49.84×

Table 15: MoE+RS performance on 8×L20 GPU.

num_tokens	hidden_size	intermediate_size	num_experts	topk	CuBLAS+NCCL	DITRON	speedup
8192	2048	1536	32	2	15.80 ms	1.71 ms	9.26×

Table 16: MoE+AR performance on 8×L20 GPU.

num_tokens	hidden_size	intermediate_size	num_experts	topk	CuBLAS+NCCL	DITRON	speedup
8192	2048	1536	32	2	18.19 ms	6.54 ms	2.78×

Table 17: A2A+GEMM performance on 8×L20 GPU.

DType	M	N	K	CuBLAS+NCCL	DITRON	speedup
Int8	7168	9216	3072	7.11 ms	2.15 ms	3.31×
FP8	7168	9216	3072	8.06 ms	2.17 ms	3.71×

E CODE EXAMPLES

In this section, we provide some simplified code examples, including All-Gather GEMM, GEMM Reduce-Scatter and GEMM All-Reduce. For each example, the code is generally divided into three main parts: the host-side launcher, the producer kernel, and the consumer kernel.

E.1 ALL-GATHER GEMM

Here we provide the code example for intra-node All-Gather GEMM. The following is the host-side Python function for launching the All-Gather GEMM operation. It orchestrates the producer (communication) and consumer (computation) kernels.

```

1 def ag_gemm(a, b, ctx: AllGatherGEMMTensorParallelContext, use_cooperative=True):
2     """allgather gemm
3     Allgather global matrix A and do matmul with local matrix B, produces local matrix C
4
5     Args:
6         a (torch.Tensor<float>): local matmul A matrix. shape: [M_per_rank, K]
7         b (torch.Tensor<float>): local matmul B matrix. shape: [N_per_rank, K]
8         ctx: (AllGatherGEMMTensorParallelContext, Optional): if not provided, created immediately
9
10    Returns:
11        C (torch.Tensor<float>): local matmul C matrix. shape: [M, N_per_rank]
12    """
13
14    assert a.shape[1] == b.shape[
15        0], f"tensor_B should has shape (col_major) [{b.shape[0]}, {a.shape[1]}], but get
16        ↳ [{b.shape[0]}]"
17    assert a.dtype == b.dtype, f"Dtype of input and weight must be same: tensor_A dtype {a.dtype},
18        ↳ tensor_B dtype {b.dtype}"
19
20    M_per_rank, K = a.shape
21    N_per_rank, _ = b.shape
22
23    assert a.shape[0] * ctx.num_ranks <= ctx.max_M and a.shape[1] == ctx.K, f"Shape of tensor_A must
24        ↳ not exceed the maxsize M of ctx: tensor_A shape [{a.shape}], ctx shape [{ctx.max_M}, {ctx.K}]"
25    assert b.shape[0] == ctx.N_per_rank, f"N_per_rank of tensor_B must match that of ctx: tensor_B
26        ↳ shape [{b.shape[0]}], ctx shape [{ctx.N_per_rank}]"
27    assert ctx.tensor_dtype == a.dtype, f"dtype of ctx must match that of ctx: tensor_A dtype
28        ↳ {a.dtype}, ctx dtype {ctx.tensor_dtype}"
29
30    C = torch.empty([ctx.num_ranks * M_per_rank, N_per_rank], dtype=a.dtype, device=a.device)
31
32    local_copy_and_barrier_all(ctx.local_rank, ctx.rank, ctx.num_ranks, a, ctx.symm_workspace,
33        ↳ ctx.symm_comm_buf, ctx.symm_barrier, M_per_rank, K, ctx.phase, is_internode=False,
34        ↳ use_cooperative=use_cooperative)
35    ctx.phase += 2
36
37    current_stream = torch.cuda.current_stream()
38    ctx.ag_intranode_stream.wait_stream(current_stream)
39
40    cp_engine_producer_all_gather_full_mesh_pull(ctx.rank, ctx.num_ranks, a, ctx.symm_workspaces,
41        ↳ ctx.symm_barriers, ctx.ag_intranode_stream, for_correctness=ctx.for_correctness,
42        ↳ all_gather_method=ctx.all_gather_method)
43
44    M_per_rank, K = a.shape
45    M = M_per_rank * ctx.num_ranks
46    grid = lambda META: (triton.cdiv(M, META["BLOCK_SIZE_M"]) * triton.cdiv(ctx.N_per_rank,
47        ↳ META["BLOCK_SIZE_N"]), )
48    kernel_consumer_gemm[grid](ctx.symm_workspace[:M], b, C, M, ctx.N_per_rank, ctx.K,
49        ↳ ctx.symm_workspace.stride(0), ctx.symm_workspace.stride(1), b.stride(1), b.stride(0),
50        ↳ c.stride(0), C.stride(1), ctx.rank, ctx.num_ranks, ctx.symm_barrier, ctx.BLOCK_M,
51        ↳ ctx.BLOCK_N, ctx.BLOCK_K, ctx.GROUP_SIZE_M, num_stages=ctx.stages, num_warps=ctx.warps)
52    current_stream.wait_stream(ctx.ag_intranode_stream)
53
54    return C

```

The producer kernel below is responsible for the All-Gather communication. It copies the local data tensor to the symmetric memory (remote_tensor_buffers) of other GPUs and then sets a signal (_set_signal_cuda) to notify the consumer kernels that the data is ready.

```

1 def cp_engine_producer_all_gather_full_mesh_pull(
2     rank,
3     num_ranks,
4     local_tensor: torch.Tensor,
5     remote_tensor_buffers: List[torch.Tensor],
6     barrier_buffers: List[torch.Tensor],
7     stream: torch.cuda.Stream,
8     for_correctness=False,
9 ):
10     M_per_rank, N = local_tensor.shape

```

```

11 rank_orders = [(rank + i) % num_ranks for i in range(num_ranks)]
12
13
14 with torch.cuda.stream(stream):
15     if for_correctness:
16         # fake a slow communication case
17         # test if the computation is waiting for the correct communication
18         _add_noise_workload_debug()
19     for src_rank in rank_orders:
20         if src_rank == rank:
21             continue
22         dst = remote_tensor_buffers[rank][src_rank * M_per_rank:(src_rank + 1) * M_per_rank, :]
23         src = remote_tensor_buffers[src_rank][src_rank * M_per_rank:(src_rank + 1) * M_per_rank,
24         ↪ :]
25         dst.copy_(src)
26         _set_signal_cuda(barrier_buffers[rank][src_rank], 1, stream)

```

The consumer kernel first waits on a signal (`dl.wait`) indicating that the required data are available. Once the data is ready, it proceeds with the GEMM computation on the data.

```

1 @triton.jit(do_not_specialize=["rank"])
2 def kernel_consumer_gemm(
3     a_ptr, b_ptr, c_ptr,
4     M, N, K,
5     stride_am, stride_ak,
6     stride_bk, stride_bn,
7     stride_cm, stride_cn, rank, WORLD_SIZE: tl.constexpr, barrier_ptr,
8     BLOCK_SIZE_M: tl.constexpr, BLOCK_SIZE_N: tl.constexpr, BLOCK_SIZE_K: tl.constexpr, #
9     GROUP_SIZE_M: tl.constexpr, #
10 ):
11     a_dtype = a_ptr.dtype.element_ty
12     b_dtype = b_ptr.dtype.element_ty
13     c_dtype = c_ptr.dtype.element_ty
14     tl.static_assert(a_dtype == b_dtype, "A and B must have the same dtype")
15
16     pid = tl.program_id(axis=0)
17     num_pid_m = tl.cdiv(M, BLOCK_SIZE_M)
18     num_pid_n = tl.cdiv(N, BLOCK_SIZE_N)
19     num_pid_in_group = GROUP_SIZE_M * num_pid_n
20     group_id = pid // num_pid_in_group
21     first_pid_m = group_id * GROUP_SIZE_M
22     group_size_m = min(num_pid_m - first_pid_m, GROUP_SIZE_M)
23     pid_m = first_pid_m + (pid % num_pid_in_group) % group_size_m
24     pid_n = (pid % num_pid_in_group) // group_size_m
25
26     m_per_rank = M // WORLD_SIZE
27     m_offset = m_per_rank * rank
28     pid_m_offset = tl.cdiv(m_offset, BLOCK_SIZE_M)
29     pid_m = (pid_m + pid_m_offset) % num_pid_m
30
31     offs_am = pid_m * BLOCK_SIZE_M
32     rank_beg = offs_am // m_per_rank
33     rank_end = (min(offs_am + BLOCK_SIZE_M, M) - 1) // m_per_rank
34     token = dl.wait(barrier_ptr + rank_beg, rank_end - rank_beg + 1, "gpu", "acquire", waitValue=1)
35
36     offs_am = (pid_m * BLOCK_SIZE_M + tl.arange(0, BLOCK_SIZE_M)) % M
37     offs_bn = (pid_n * BLOCK_SIZE_N + tl.arange(0, BLOCK_SIZE_N)) % N
38     offs_k = tl.arange(0, BLOCK_SIZE_K)
39     a_ptrs = a_ptr + (offs_am[:, None] * stride_am + offs_k[None, :] * stride_ak)
40     b_ptrs = b_ptr + (offs_k[:, None] * stride_bk + offs_bn[None, :] * stride_bn)
41
42     a_ptrs = dl.consume_token(a_ptrs, token)
43
44     if a_dtype == tl.int8:
45         accumulator = tl.zeros((BLOCK_SIZE_M, BLOCK_SIZE_N), dtype=tl.int32)
46     else:
47         accumulator = tl.zeros((BLOCK_SIZE_M, BLOCK_SIZE_N), dtype=tl.float32)
48     for k in range(0, tl.cdiv(K, BLOCK_SIZE_K)):
49         a = tl.load(a_ptrs, mask=offs_k[None, :] < K - k * BLOCK_SIZE_K, other=0.0)
50         b = tl.load(b_ptrs, mask=offs_k[:, None] < K - k * BLOCK_SIZE_K, other=0.0)
51         accumulator += tl.dot(a, b)
52         a_ptrs += BLOCK_SIZE_K * stride_ak
53         b_ptrs += BLOCK_SIZE_K * stride_bk
54
55     offs_cm = pid_m * BLOCK_SIZE_M + tl.arange(0, BLOCK_SIZE_M)
56     offs_cn = pid_n * BLOCK_SIZE_N + tl.arange(0, BLOCK_SIZE_N)
57     c_ptrs = c_ptr + stride_cm * offs_cm[:, None] + stride_cn * offs_cn[None, :]
58     c_mask = (offs_cm[:, None] < M) & (offs_cn[None, :] < N)
59
60     tl.store(c_ptrs, accumulator.to(c_dtype), mask=c_mask)

```

E.2 GEMM REDUCE-SCATTER

Here we provide the code example for GEMM Reduce-Scatter. The following is the host-side Python function for launching the GEMM Reduce-Scatter operation.

```

1  def gemm_rs_op(input, weight, ctx: GEMMReduceScatterTensorParallelContext, persistent: bool = True,
2      fuse_scatter: bool = False):
3      world_size = ctx.rs_ctx.world_size
4      local_world_size = ctx.rs_ctx.local_world_size
5      rs_stream = ctx.rs_stream
6      output_dtype = ctx.output_dtype
7      num_gemm_sms = ctx.num_gemm_sms
8
9      M, local_K = input.shape
10     N = weight.shape[0]
11     assert N == ctx.rs_ctx.N
12
13     assert M % world_size == 0
14     assert weight.shape[1] == local_K
15     M_per_rank = M // world_size
16     current_stream = torch.cuda.current_stream()
17     rs_stream.wait_stream(current_stream)
18
19     output = torch.empty((M_per_rank, N), dtype=output_dtype, device=input.device)
20     workspace = torch.zeros((world_size, ), dtype=torch.int32, device=input.device)
21     gemm_out = ctx.get_gemm_out_buf(input)
22     scatter_signal = ctx.rs_ctx.scatter_signal_buf
23
24     triton_config = triton.Config(
25         {
26             "BLOCK_SIZE_M": ctx.BLOCK_M, "BLOCK_SIZE_N": ctx.BLOCK_N, "BLOCK_SIZE_K": ctx.BLOCK_K,
27             "GROUP_SIZE_M": ctx.GROUP_M
28         }, num_stages=ctx.stages, num_warps=8)
29     triton_config = update_triton_config(M, N, local_K, input.dtype, world_size, local_world_size,
30         triton_config)
31     gemm_rs_producer(input, weight, gemm_out, scatter_signal, workspace, world_size,
32         local_world_size, fuse_scatter, triton_config)
33
34     if not fuse_scatter:
35         with torch.cuda.stream(rs_stream):
36             reduce_scatter_2d_op(gemm_out, ctx.rs_ctx, output)
37         current_stream.wait_stream(rs_stream)
38     else:
39         nvshmem_barrier_all_on_stream(current_stream)
40         ring_reduce(gemm_out, output, ctx.rs_ctx.local_rank, local_world_size)
41         nvshmem_barrier_all_on_stream(current_stream)
42     return output

```

The producer kernel below is responsible for the GEMM computation. When FUSE_SCATTER is enabled, it directly writes the output to the symmetric memory of the destination rank.

```

1  @triton.jit(launch_metadata=_matmul_launch_metadata, repr=_gemm_rs_non_persistent_repr)
2  def kernel_gemm_rs_producer(
3      # Pointers to matrices
4      a_ptr, # [M, K]_Ti
5      b_ptr, # [K, N]_Ti
6      c_ptr, # [M, N]_To
7      # Matrix dimensions
8      M,
9      N,
10     K,
11     # The stride variables represent how much to increase the ptr by when moving by 1
12     # element in a particular dimension. E.g. 'stride_am' is how much to increase 'a_ptr'
13     # by to get the element one row down (A has M rows).
14     stride_am,
15     stride_ak, #
16     stride_bk,
17     stride_bn, #
18     stride_cm,
19     stride_cn,
20     barrier_ptr,
21     counter_ptr,
22     FUSE_SCATTER: tl.constexpr,
23     LOCAL_WORLD_SIZE: tl.constexpr,
24     WORLD_SIZE: tl.constexpr,
25     # Meta-parameters
26     BLOCK_SIZE_M: tl.constexpr,
27     BLOCK_SIZE_N: tl.constexpr,
28     BLOCK_SIZE_K: tl.constexpr, #
29     GROUP_SIZE_M: tl.constexpr,
30 ):
31     """Kernel for computing the matmul C = A x B.
32     A has shape (M, K), B has shape (K, N) and C has shape (M, N)
33     """
34     tl.static_assert(a_ptr.dtype.is_ptr(), "A should be a pointer")
35     tl.static_assert(b_ptr.dtype.is_ptr(), "B should be a pointer")
36     tl.static_assert(c_ptr.dtype.is_ptr(), "C should be a pointer")
37     a_dtype = a_ptr.dtype.element_ty
38     b_dtype = b_ptr.dtype.element_ty
39     tl.static_assert(a_dtype == b_dtype, "A and B should have the same dtype")
40
41     rank = dl.rank()
42     NNODES = WORLD_SIZE // LOCAL_WORLD_SIZE
43
44     pid = tl.program_id(axis=0)

```

```

1404 45     num_pid_m = tl.cdiv(M, BLOCK_SIZE_M)
1405 46     num_pid_n = tl.cdiv(N, BLOCK_SIZE_N)
1406 47
1407 48     M_per_rank = M // WORLD_SIZE
1408 49     pid_m, pid_n = swizzle_2d(pid, num_pid_m, num_pid_n, GROUP_SIZE_M)
1409 50
1410 51     if NNODES != 1: # with complex threadblock swizzle logic
1411 52         pid_m = threadblock_swizzle_gemm_reduce_scatter_kernel(pid_m, M, rank, WORLD_SIZE, NNODES,
1412 53             ↪ BLOCK_SIZE_M)
1413 54     else:
1414 55         pid_m_offset = (rank + 1) * M_per_rank // BLOCK_SIZE_M
1415 56         pid_m = (pid_m + pid_m_offset) % num_pid_m
1416 57
1417 58         # Create pointers for the first blocks of A and B.
1418 59         offs_am = (pid_m * BLOCK_SIZE_M + tl.arange(0, BLOCK_SIZE_M)) % M
1419 60         offs_bn = (pid_n * BLOCK_SIZE_N + tl.arange(0, BLOCK_SIZE_N)) % N
1420 61         offs_k = tl.arange(0, BLOCK_SIZE_K)
1421 62         a_ptrs = a_ptr + (offs_am[:, None] * stride_am + offs_k[None, :] * stride_ak)
1422 63         b_ptrs = b_ptr + (offs_k[:, None] * stride_bk + offs_bn[None, :] * stride_bn)
1423 64
1424 65         # Iterate to compute a block of the C matrix.
1425 66         if a_ptr.dtype.element_ty == tl.int8:
1426 67             accumulator = tl.zeros((BLOCK_SIZE_M, BLOCK_SIZE_N), dtype=tl.int32)
1427 68         else:
1428 69             accumulator = tl.zeros((BLOCK_SIZE_M, BLOCK_SIZE_N), dtype=tl.float32)
1429 70
1430 71         for k in range(0, tl.cdiv(K, BLOCK_SIZE_K)):
1431 72             # Load the next block of A and B, generate a mask by checking the K dimension.
1432 73             a = tl.load(a_ptrs, mask=offs_k[None, :] < K - k * BLOCK_SIZE_K, other=0.0)
1433 74             b = tl.load(b_ptrs, mask=offs_k[:, None] < K - k * BLOCK_SIZE_K, other=0.0)
1434 75             # We accumulate along the K dimension.
1435 76             accumulator += tl.dot(a, b)
1436 77             # Advance the ptrs to the next K block.
1437 78             a_ptrs += BLOCK_SIZE_K * stride_ak
1438 79             b_ptrs += BLOCK_SIZE_K * stride_bk
1439 80
1440 81         # Write back the block of the output matrix C with masks.
1441 82         offs_cm = pid_m * BLOCK_SIZE_M + tl.arange(0, BLOCK_SIZE_M)
1442 83         offs_cn = pid_n * BLOCK_SIZE_N + tl.arange(0, BLOCK_SIZE_N)
1443 84         c_ptrs = c_ptr + stride_cm * offs_cm[:, None] + stride_cn * offs_cn[None, :]
1444 85         out_mask = (offs_cm[:, None] < M) & (offs_cn[None, :] < N)
1445 86
1446 87         if not FUSE_SCATTER:
1447 88             tl.store(c_ptrs, accumulator, mask=out_mask)
1448 89
1449 90             # inc barrier
1450 91             segment_start = pid_m * BLOCK_SIZE_M // M_per_rank
1451 92             segment_end = (min((pid_m + 1) * BLOCK_SIZE_M, M) - 1) // M_per_rank
1452 93             __syncthreads()
1453 94             segment = segment_start + tid(axis=0)
1454 95             if segment <= segment_end:
1455 96                 m_start = M_per_rank * segment
1456 97                 m_end = M_per_rank * (segment + 1) - 1
1457 98                 tiled_m_start = m_start // BLOCK_SIZE_M
1458 99                 tiled_m_end = m_end // BLOCK_SIZE_M
1459 100                 tiled_m_size = tiled_m_end - tiled_m_start + 1
1460 101                 val = atomic_add(counter_ptr + segment, 1, semantic="release", scope="gpu")
1461 102                 if val == num_pid_n * tiled_m_size - 1:
1462 103                     # or use other signal op semantic
1463 104                     st(barrier_ptr + segment, 1, scope="gpu", semantic="release")
1464 105             else:
1465 106                 rank_start = pid_m * BLOCK_SIZE_M // M_per_rank
1466 107                 rank_end = (min((pid_m + 1) * BLOCK_SIZE_M, M) - 1) // M_per_rank
1467 108                 for cur_rank in range(rank_start, rank_end + 1):
1468 109                     m_start = max(M_per_rank * cur_rank, pid_m * BLOCK_SIZE_M)
1469 110                     m_end = min(M_per_rank * (cur_rank + 1) - 1, (min((pid_m + 1) * BLOCK_SIZE_M, M) - 1))
1470 111                     remote_c_ptr = dl.symm_at(c_ptr, cur_rank)
1471 112                     mask_offset = m_start - pid_m * BLOCK_SIZE_M
1472 113                     remote_offs_cm = m_start % M_per_rank + rank * M_per_rank + tl.arange(0, BLOCK_SIZE_M) -
1473 114                         ↪ mask_offset
1474 115                     remote_c_ptrs = remote_c_ptr + stride_cm * remote_offs_cm[:, None] + stride_cn *
1475 116                         ↪ offs_cn[None, :]
1476 117                     remote_mask = (offs_cm[:, None] <= m_end) & (offs_cm[:, None] >= m_start) &
1477 118                         ↪ (offs_cn[None, :] < N)
1478 119                     tl.store(remote_c_ptrs, accumulator, mask=remote_mask)

```

The consumer kernel performs the reduce-scatter communication operation. After the GEMM computation is complete, this kernel reduces the partial results across ranks and scatters them to produce the final output.

```

1452 1 def reduce_scatter_multi_node(input: torch.Tensor, ctx: ReduceScatter2DContext, output:
1453 2     ↪ Optional[torch.Tensor] = None):
1454 3     """
1455 4     A hierarchical reduce-scatter implementation that overlaps the intra-node scatter
1456 5     with the local reduce and the inter-node p2p(after reduce). It also provides a rank-wise
1457 6     signal and supports overlap with gemm.
1458 7     """
1459 8     M, N = input.shape
1460 9     M_per_rank = M // ctx.world_size

```

```

9
10 current_stream = torch.cuda.current_stream()
11 ctx.reduction_stream.wait_stream(current_stream)
12
13 # directly reduce_scatter to output if nnodes == 1
14 out_each_node = output if ctx.nnodes == 1 else None
15 if not has_fullmesh_nvlink():
16     rs_result_per_node = reduce_scatter_for_each_node_ring(input, ctx, out_each_node)
17 else:
18     rs_result_per_node = reduce_scatter_for_each_node(input, ctx, out_each_node)
19
20 if ctx.nnodes == 1:
21     return rs_result_per_node
22
23 nvshmem_barrier_all_on_stream(current_stream)
24 if output is None:
25     output = torch.empty((M_per_rank, N), dtype=input.dtype, device=input.device)
26 ring_reduce(rs_result_per_node, output, ctx.node_id, ctx.nnodes)
27 return output
28
29
30 def reduce_scatter_2d_op(input: torch.Tensor, ctx: ReduceScatter2DContext, output:
    ↪ Optional[torch.Tensor] = None):
31     M, N = input.shape
32     assert input.dtype == ctx.dtype
33     assert ctx.max_M >= M and ctx.N == N
34     assert M % ctx.world_size == 0
35
36     nvshmem_barrier_all_on_stream(torch.cuda.current_stream())
37     output = reduce_scatter_multi_node(input, ctx, output)
38     ctx.reset_barriers()
39     return output

```

E.3 GEMM ALL-REDUCE

Here we provide the example for GEMM All-Reduce. The following is the host-side Python function for launching the All-Reduce GEMM operation.

```

1 def gemm_allreduce_op(ctx: GemmARContext, a, b, gemm_config: triton.Config, copy_to_local=True,
    ↪ USE_MULTIMEM_ST=True):
2
3     current_stream = torch.cuda.current_stream()
4     ar_stream = ctx.ar_stream
5     ar_stream.wait_stream(current_stream)
6
7     M, N = a.shape[0], b.shape[0]
8     # Check constraints.
9     assert a.shape[1] == b.shape[1], "Incompatible dimensions"
10    assert a.dtype == b.dtype, "Incompatible dtypes"
11    symm_c = ctx.get_gemm_out_buf(a, b)
12    symm_ar_out = ctx.symm_ar_out_buf
13    gemm_barrier = ctx.gemm_barrier_buf
14    multi_st_barrier = ctx.multi_st_barrier_buf
15    NUM_COMM_SMS = ctx.NUM_COMM_SMS
16    ar_out = torch.empty((M, N), dtype=a.dtype, device=a.device)
17    BLOCK_SIZE_M = gemm_config.all_kargs()["BLOCK_SIZE_M"]
18    BLOCK_SIZE_N = gemm_config.all_kargs()["BLOCK_SIZE_N"]
19    # add mask in 'consumer_all_reduce' can remove this constraint
20    assert N % BLOCK_SIZE_N == 0
21    persistent_gemm_notify(a, b, symm_c, gemm_barrier, gemm_config)
22    with torch.cuda.stream(ar_stream):
23        consumer_all_reduce(symm_c, symm_ar_out, ar_out, gemm_barrier, multi_st_barrier,
            ↪ BLOCK_SIZE_M=BLOCK_SIZE_M,
24            BLOCK_SIZE_N=BLOCK_SIZE_N, NUM_COMM_SMS=NUM_COMM_SMS,
            ↪ USE_MULTIMEM_ST=USE_MULTIMEM_ST)
25    current_stream.wait_stream(ar_stream)
26    if USE_MULTIMEM_ST and copy_to_local:
27        ar_out.copy_(symm_ar_out.reshape(-1)[:M * N].reshape(M, N))
28    nvshmem_barrier_all_on_stream(current_stream)
29    if USE_MULTIMEM_ST and not copy_to_local:
30        return symm_ar_out.reshape(-1)[:M * N].reshape(M, N)
31    return ar_out

```

The producer kernel below is responsible for the GEMM computation. It performs computation and notifies the consumer kernels when tiles are ready for all-reduce.

```

1 @triton.jit(do_not_specialize=[])
2 def kernel_persistent_gemm_notify(a_ptr, b_ptr, c_ptr, gemm_barrier_ptr, #
3     M, N, K, #
4     stride_am, stride_ak, #
5     stride_bn, stride_bk, #
6     stride_cm, stride_cn, #
7     BLOCK_SIZE_M: tl.constexpr, #
8     BLOCK_SIZE_N: tl.constexpr, #

```

```

9          BLOCK_SIZE_K: tl.constexpr, #
10         GROUP_SIZE_M: tl.constexpr, #
11         NUM_GEMM_SMS: tl.constexpr, #
12     ):
13         start_pid = tl.program_id(axis=0)
14         num_pid_m = tl.cdiv(M, BLOCK_SIZE_M)
15         num_pid_n = tl.cdiv(N, BLOCK_SIZE_N)
16         k_tiles = tl.cdiv(K, BLOCK_SIZE_K)
17         num_tiles = num_pid_m * num_pid_n
18
19         tile_id_c = start_pid - NUM_GEMM_SMS
20
21         offs_k_for_mask = tl.arange(0, BLOCK_SIZE_K)
22         num_pid_in_group = GROUP_SIZE_M * num_pid_n
23
24         for tile_id in tl.range(start_pid, num_tiles, NUM_GEMM_SMS):
25             pid_m, pid_n = _compute_pid(tile_id, num_pid_in_group, num_pid_m, GROUP_SIZE_M, NUM_GEMM_SMS)
26             start_m = pid_m * BLOCK_SIZE_M
27             start_n = pid_n * BLOCK_SIZE_N
28             offs_am = start_m + tl.arange(0, BLOCK_SIZE_M)
29             offs_bn = start_n + tl.arange(0, BLOCK_SIZE_N)
30             offs_am = tl.where(offs_am < M, offs_am, 0)
31             offs_bn = tl.where(offs_bn < N, offs_bn, 0)
32             offs_am = tl.max_contiguous(tl.multiple_of(offs_am, BLOCK_SIZE_M), BLOCK_SIZE_M)
33             offs_bn = tl.max_contiguous(tl.multiple_of(offs_bn, BLOCK_SIZE_N), BLOCK_SIZE_N)
34
35             accumulator = tl.zeros((BLOCK_SIZE_M, BLOCK_SIZE_N), dtype=tl.float32)
36             for ki in range(k_tiles):
37                 offs_k = ki * BLOCK_SIZE_K + tl.arange(0, BLOCK_SIZE_K)
38                 a_ptrs = a_ptr + (offs_am[:, None] * stride_am + offs_k[None, :] * stride_ak)
39                 b_ptrs = b_ptr + (offs_bn[:, None] * stride_bn + offs_k[None, :] * stride_bk)
40
41                 a = tl.load(a_ptrs, mask=offs_k_for_mask[None, :] < K - ki * BLOCK_SIZE_K, other=0.0)
42                 b = tl.load(b_ptrs, mask=offs_k_for_mask[None, :] < K - ki * BLOCK_SIZE_K, other=0.0)
43                 accumulator = tl.dot(a, b.T, accumulator)
44
45             tile_id_c += NUM_GEMM_SMS
46             pid_m, pid_n = _compute_pid(tile_id_c, num_pid_in_group, num_pid_m, GROUP_SIZE_M,
47                                     NUM_GEMM_SMS)
48             offs_cm = pid_m * BLOCK_SIZE_M + tl.arange(0, BLOCK_SIZE_M)
49             offs_cn = pid_n * BLOCK_SIZE_N + tl.arange(0, BLOCK_SIZE_N)
50             c_ptrs = c_ptr + stride_cm * offs_cm[:, None] + stride_cn * offs_cn[None, :]
51             c_mask = (offs_cm[:, None] < M) & (offs_cn[None, :] < N)
52             c = accumulator.to(c_ptr.dtype.element_ty)
53             tl.store(c_ptrs, c, mask=c_mask)
54             __syncthreads()
55
56             thread_idx = tid(0)
57             gemm_barrier_idx = pid_m * num_pid_n + pid_n
58             if thread_idx == 0:
59                 st(gemm_barrier_ptr + gemm_barrier_idx, 1, scope="gpu", semantic="release")

```

The consumer kernel performs the all-reduce operation. It waits for GEMM tiles to be ready, then performs reduction across all ranks.

```

1 @triton.jit(do_not_specialize=[])
2 def consumer_all_reduce_kernel(
3     symm_input_ptr,
4     symm_ar_out_ptr,
5     ar_out_ptr, #
6     gemm_barrier_ptr,
7     multi_st_barrier_ptr, #
8     M,
9     N, #
10    BLOCK_SIZE_M: tl.constexpr, #
11    BLOCK_SIZE_N: tl.constexpr, #
12    NUM_COMM_SMS: tl.constexpr,
13    USE_MULTIMEM_ST: tl.constexpr,
14 ):
15     rank = dl.rank()
16     world_size = dl.num_ranks()
17     pid = tl.program_id(0)
18     num_pid_m = tl.cdiv(M, BLOCK_SIZE_M)
19     num_pid_n = tl.cdiv(N, BLOCK_SIZE_N)
20     num_tiles = num_pid_m * num_pid_n
21     thread_idx = tid(0)
22     block_dim = num_warps() * 32
23     VEC_SIZE: tl.constexpr = 128 // tl.constexpr(symm_input_ptr.dtype.element_ty.primitive_bitwidth)
24
25     tl.static_assert(BLOCK_SIZE_N % VEC_SIZE == 0)
26     VEC_PER_ROW = BLOCK_SIZE_N // VEC_SIZE
27     src_data_mc_ptr = libshmem_device.remote_mc_ptr(libshmem_device.NVSHMEMX_TEAM_NODE,
28                                     symm_input_ptr)
29     if not USE_MULTIMEM_ST:
30         for tile_id in range(pid, num_tiles, NUM_COMM_SMS):
31             pid_m = tile_id // num_pid_n
32             pid_n = tile_id % num_pid_n
33             if thread_idx < world_size:
34                 peer_gemm_barrier_ptr = dl.symm_at(gemm_barrier_ptr, thread_idx)
35                 while ld(peer_gemm_barrier_ptr + tile_id, scope="sys", semantic="acquire") != 1:

```

```

35         pass
36         __syncthreads()
37         tile_m = min(M - pid_m * BLOCK_SIZE_M, BLOCK_SIZE_M)
38         cur_tile_nelem = tile_m * BLOCK_SIZE_N
39         for idx in range(thread_idx, cur_tile_nelem // VEC_SIZE, block_dim):
40             row_id = idx // VEC_PER_ROW
41             col_id = idx % VEC_PER_ROW
42             offset = (row_id + pid_m * BLOCK_SIZE_M) * N + col_id * VEC_SIZE + pid_n *
43                 BLOCK_SIZE_N
44             val0, val1, val2, val3 = multimem_ld_reduce_v4(src_data_mc_ptr + offset)
45             st_v4_b32(ar_out_ptr + offset, val0, val1, val2, val3)
46     else:
47         symm_out_mc_ptr = libshmem_device.remote_mc_ptr(libshmem_device.NVSHMEMX_TEAM_NODE,
48             symm_ar_out_ptr)
49         for tile_id in range(pid + rank * NUM_COMM_SMS, num_tiles, NUM_COMM_SMS * world_size):
50             pid_m = tile_id // num_pid_n
51             pid_n = tile_id % num_pid_n
52             if thread_idx < world_size:
53                 peer_gemm_barrier_ptr = dl.symm_at(gemm_barrier_ptr, thread_idx)
54                 while ld(peer_gemm_barrier_ptr + tile_id, scope="sys", semantic="acquire") != 1:
55                     pass
56                 __syncthreads()
57                 tile_m = min(M - pid_m * BLOCK_SIZE_M, BLOCK_SIZE_M)
58                 cur_tile_nelem = tile_m * BLOCK_SIZE_N
59                 for idx in range(thread_idx, cur_tile_nelem // VEC_SIZE, block_dim):
60                     row_id = idx // VEC_PER_ROW
61                     col_id = idx % VEC_PER_ROW
62                     offset = (row_id + pid_m * BLOCK_SIZE_M) * N + col_id * VEC_SIZE + pid_n *
63                         BLOCK_SIZE_N
64                     val0, val1, val2, val3 = multimem_ld_reduce_v4(src_data_mc_ptr + offset)
65                     multimem_st_v4(symm_out_mc_ptr + offset, val0, val1, val2, val3)
66                 __syncthreads()
67             # barrier on all blocks with same pid
68             # 0. set barrier to all blocks with same pid on all peer ranks
69             if thread_idx < world_size:
70                 peer_ptr = dl.symm_at(multi_st_barrier_ptr, thread_idx)
71                 st(peer_ptr + rank * NUM_COMM_SMS + pid, 1, scope="sys", semantic="release")
72             # 1. wait barrier
73             if thread_idx < world_size:
74                 multi_st_barrier_idx = thread_idx * NUM_COMM_SMS + pid
75                 while ld(multi_st_barrier_ptr + multi_st_barrier_idx, scope="sys", semantic="acquire") !=
76                     1:
77                     pass
78                 st(multi_st_barrier_ptr + multi_st_barrier_idx, 0)
79             # Each block can safely reset the part of the barriers it is waiting for.
80             # In low latency kernel, we ensure that the gemm barriers used for two consecutive iteration are
81             # different,
82             # it can be reset without any sync.
83             for tile_id in range(pid + rank * NUM_COMM_SMS, num_tiles, NUM_COMM_SMS * world_size):
84                 peer_gemm_barrier_ptr = dl.symm_at(gemm_barrier_ptr, thread_idx)
85                 if thread_idx < world_size:
86                     st(peer_gemm_barrier_ptr + tile_id, 0, scope="sys", semantic="relaxed")

```

E.4 REGISTER MEGAKERNEL TASK

In this section, we use linear (GEMM) as an example to show how to register kernels as tasks. The following code contains a tile-level GEMM kernel, a chunk-level (multiple tiles) GEMM kernel, and an interface kernel to keep unified function interface for all various tasks. These code are written by users and only minor modification is required to change users' original kernels.

```

1 @triton.jit
2 def tile_wise_matmul_compute(tile_id, a_ptr, b_ptr, c_ptr, M, N, K, BLOCK_SIZE_M, BLOCK_SIZE_N,
3     BLOCK_SIZE_K,
4     NUM_STAGES):
5     # linear: a (M, K) x b (N, K) -> c (M, N)
6     num_pid_n = tl.cdiv(N, BLOCK_SIZE_N)
7     k_tiles = tl.cdiv(K, BLOCK_SIZE_K)
8     offs_k_for_mask = tl.arange(0, BLOCK_SIZE_K)
9
10    pid_m = tile_id // num_pid_n
11    pid_n = tile_id % num_pid_n
12    start_m = pid_m * BLOCK_SIZE_M
13    start_n = pid_n * BLOCK_SIZE_N
14    offs_am = start_m + tl.arange(0, BLOCK_SIZE_M)
15    offs_bn = start_n + tl.arange(0, BLOCK_SIZE_N)
16    offs_am = tl.where(offs_am < M, offs_am, 0)
17    offs_bn = tl.where(offs_bn < N, offs_bn, 0)
18    offs_am = tl.max_contiguous(tl.multiple_of(offs_am, BLOCK_SIZE_M), BLOCK_SIZE_M)
19    offs_bn = tl.max_contiguous(tl.multiple_of(offs_bn, BLOCK_SIZE_N), BLOCK_SIZE_N)
20    accumulator = tl.zeros((BLOCK_SIZE_M, BLOCK_SIZE_N), dtype=tl.float32)
21    for ki in tl.range(0, k_tiles, num_stages=NUM_STAGES):

```

```

22     offs_k = ki * BLOCK_SIZE_K + tl.arange(0, BLOCK_SIZE_K)
23     a_ptrs = a_ptr + (offs_am[:, None] * K + offs_k[None, :])
24     b_ptrs = b_ptr + (offs_bn[:, None] * K + offs_k[None, :])
25
26     a = tl.load(a_ptrs, mask=offs_k_for_mask[None, :] < K - ki * BLOCK_SIZE_K, other=0.0)
27     b = tl.load(b_ptrs, mask=offs_k_for_mask[None, :] < K - ki * BLOCK_SIZE_K, other=0.0)
28     accumulator = tl.dot(a, b.T, accumulator)
29
30     offs_cm = pid_m * BLOCK_SIZE_M + tl.arange(0, BLOCK_SIZE_M)
31     offs_cn = pid_n * BLOCK_SIZE_N + tl.arange(0, BLOCK_SIZE_N)
32     c_ptrs = c_ptr + N * offs_cm[:, None] + offs_cn[None, :]
33     c_mask = (offs_cm[:, None] < M) & (offs_cn[None, :] < N)
34     c = accumulator.to(c_ptr.dtype.element_ty)
35     tl.store(c_ptrs, c, mask=c_mask)
36
37
38 @triton.jit
39 def tile_range_matmul_compute_and_notify(tile_start, sb_base_ptr, a_ptr, b_ptr, c_ptr, M, N, K,
40     ↪ BLOCK_SIZE_M,
41                                     BLOCK_SIZE_N, BLOCK_SIZE_K, NUM_STAGES, TILE_READY_SIGNAL,
42     ↪ NUM_SMS):
43
44     num_pid_m = tl.cdiv(M, BLOCK_SIZE_M)
45     num_pid_n = tl.cdiv(N, BLOCK_SIZE_N)
46     num_tiles = num_pid_m * num_pid_n
47     for tile_id in tl.range(tile_start, num_tiles, NUM_SMS, flatten=True, warp_specialize=True):
48         tile_wise_matmul_compute(tile_id, a_ptr, b_ptr, c_ptr, M, N, K, BLOCK_SIZE_M, BLOCK_SIZE_N,
49     ↪ BLOCK_SIZE_K,
50                                     NUM_STAGES)
51         st(sb_base_ptr + tile_id, TILE_READY_SIGNAL, "gpu", "release")
52
53 @triton.jit
54 def linear_task_compute(task_base_info: TaskBaseInfo, scoreboard: Scoreboard, BLOCK_SIZE_M:
55     ↪ tl.constexpr,
56                                     BLOCK_SIZE_N: tl.constexpr, BLOCK_SIZE_K: tl.constexpr, NUM_STAGES:
57     ↪ tl.constexpr,
58                                     ALIGNMENT_K: tl.constexpr):
59
60     input: TensorDesc = task_base_info.get_tensor(0)
61     weight: TensorDesc = task_base_info.get_tensor(1)
62     output: TensorDesc = task_base_info.get_tensor(2)
63
64     M = input.size(0)
65     K = input.size(1, ALIGNMENT_K)
66     N = weight.size(0)
67     a_ptr = input.data_ptr(tl.bfloat16)
68     b_ptr = weight.data_ptr(tl.bfloat16)
69     c_ptr = output.data_ptr(tl.bfloat16)
70
71     tile_id = task_base_info.tile_id_or_start
72     tile_wise_matmul_compute(tile_id, a_ptr, b_ptr, c_ptr, M, N, K, BLOCK_SIZE_M, BLOCK_SIZE_N,
73     ↪ BLOCK_SIZE_K,
74                                     NUM_STAGES)
75     scoreboard.release_tile(task_base_info, tile_id)

```

Then, we can register the kernels as tasks using the following interface. Users are required to describe the problem size, data dependency, and the tiling configurations.

```

1  @registry.register_task(op_type="linear", task_cls=LinearTask, config_factory=linear_config_factory,
2     codegen_func=codegen_linear)
3  class LinearTaskBuilder(TaskBuilderBase):
4
5      @classmethod
6      def get_problem_size(cls, io_tensors: List[List['torch.Tensor']], extra_params: Dict[str, Any]):
7          a, b = io_tensors[0]
8          M, K = a.shape
9          N, K = b.shape
10         return (M, N, K)
11
12     @classmethod
13     def _build_tasks_impl(cls, device_prop, layer_id: int, dependency: TaskDependency, io_tensors,
14     ↪ extra_params,
15                             tile_wise=True, config_args={}) -> List[TaskBase]:
16         assert tile_wise == True # noga: E712
17         kernel_config = cls.create_config(**config_args)
18         task_id = cls.get_task_id(layer_id)
19         BLOCK_SIZE_M = kernel_config.BLOCK_SIZE_M
20         BLOCK_SIZE_N = kernel_config.BLOCK_SIZE_N
21         M, N, K = cls.get_problem_size(io_tensors, extra_params)
22         num_tiles_m = cdiv(M, BLOCK_SIZE_M)
23         num_tiles_n = cdiv(N, BLOCK_SIZE_N)
24         num_tiles = num_tiles_m * num_tiles_n
25         x, w = io_tensors[0]
26         y = io_tensors[1][0]
27
28         num_sm = device_prop.NUM_SMS
29         tasks = []
30         cls.log(

```

```

30         f"Linear Task: M = {M}, N = {N}, K = {K}, num_tiles = {num_tiles}, num_sm = {num_sm},
    ↪ tile_wise = {tile_wise}, dependency = {dependency}, BLOCK_SIZE_M = {BLOCK_SIZE_M},
    ↪ BLOCK_SIZE_N = {BLOCK_SIZE_N}"
31     )
32     for tm in range(num_tiles_m):
33         for tn in range(num_tiles_n):
34             tile_id = tm * num_tiles_n + tn
35             bm = min(BLOCK_SIZE_M, M - tm * BLOCK_SIZE_M)
36             bn = min(BLOCK_SIZE_N, N - tn * BLOCK_SIZE_N)
37             x_desc = InputDependencyDesc(x, require_full=False, start_indices=(tm * BLOCK_SIZE_M,
    ↪ 0),
38                                     data_sizes=(bm, K))
39             w_desc = InputDependencyDesc(w, require_full=False, start_indices=(tn * BLOCK_SIZE_N,
    ↪ 0),
40                                     data_sizes=(bn, K))
41             y_desc = OutputTilingDesc(tile_sizes=(BLOCK_SIZE_M, BLOCK_SIZE_N),
42                                     start_indices=(tm * BLOCK_SIZE_M, tn * BLOCK_SIZE_N))
43             inputs_dep = {x: x_desc, w: w_desc}
44             outs_tile_mapping = {y: y_desc}
45             tasks.append(
46                 cls._create_task(layer_id, task_id, tile_id, num_tiles, kernel_config,
    ↪ dependency, io_tensors,
47                                     extra_params, inputs_dep, outs_tile_mapping))
48     return tasks
49
50 @classmethod
51 def build_tasks(cls, device_prop: 'DeviceProp', layer_id: int, dependency: TaskDependency,
52                 io_tensors: List[List['torch.Tensor']], extra_params: Dict[str, Any]) ->
    ↪ List[TaskBase]:
53     return cls._build_tasks_impl(device_prop, layer_id, dependency, io_tensors, extra_params)

```

Using the registration interface, multiple tasks can be added. And DITRON will **automatically** generate MegaKernel. Here we show the generated code. Note that the code is auto-generated, not written by users. The task-level primitives such as *fetch_task* and *scoreboard.wait_deps* are used here.

```

1  @triton.jit
2  def FETCH_TASK(work_queues, idx, INT_PER_TASK, NUM_SMS, MAX_NUM_TENSOR_DIMS,
    ↪ ENABLE_RUNTIME_SCHEDULER=False):
3      sm_id = tl.program_id(axis=0)
4      TASK_TYPE_OFFSET = 0
5      LAYER_ID_OFFSET = 1
6      TASK_ID_OFFSET = 2
7      TILE_ID_OR_START_OFFSET = 3
8      DEPEND_ENTRY_START_OFFSET = 4
9      DEPEND_ENTRY_END_OFFSET = 5
10     IO_TENSORS_OFFSET = 6
11
12     if not ENABLE_RUNTIME_SCHEDULER:
13         offset = INT_PER_TASK * NUM_SMS
14
15         task_type = tl.load(work_queues + idx * offset + sm_id * INT_PER_TASK +
    ↪ TASK_TYPE_OFFSET).to(tl.int32)
16         layer_id = tl.load(work_queues + idx * offset + sm_id * INT_PER_TASK +
    ↪ LAYER_ID_OFFSET).to(tl.int32)
17         task_id = tl.load(work_queues + idx * offset + sm_id * INT_PER_TASK +
    ↪ TASK_ID_OFFSET).to(tl.int32)
18         tile_id_or_start = tl.load(work_queues + idx * offset + sm_id * INT_PER_TASK +
    ↪ TILE_ID_OR_START_OFFSET).to(tl.int32)
19         depend_entry_start = tl.load(work_queues + idx * offset + sm_id * INT_PER_TASK +
    ↪ DEPEND_ENTRY_START_OFFSET).to(tl.int32)
20         depend_entry_end = tl.load(work_queues + idx * offset + sm_id * INT_PER_TASK +
    ↪ DEPEND_ENTRY_END_OFFSET).to(tl.int32)
21         io_tensors_ptr = work_queues + idx * offset + sm_id * INT_PER_TASK + IO_TENSORS_OFFSET
22     else:
23         task_type = tl.load(work_queues + idx * INT_PER_TASK + TASK_TYPE_OFFSET).to(tl.int32)
24         layer_id = tl.load(work_queues + idx * INT_PER_TASK + LAYER_ID_OFFSET).to(tl.int32)
25         task_id = tl.load(work_queues + idx * INT_PER_TASK + TASK_ID_OFFSET).to(tl.int32)
26         tile_id_or_start = tl.load(work_queues + idx * INT_PER_TASK +
    ↪ TILE_ID_OR_START_OFFSET).to(tl.int32)
27         depend_entry_start = tl.load(work_queues + idx * INT_PER_TASK +
    ↪ DEPEND_ENTRY_START_OFFSET).to(tl.int32)
28         depend_entry_end = tl.load(work_queues + idx * INT_PER_TASK +
    ↪ DEPEND_ENTRY_END_OFFSET).to(tl.int32)
29         io_tensors_ptr = work_queues + idx * INT_PER_TASK + IO_TENSORS_OFFSET
30
31     task_base_info = TaskBaseInfo(io_tensors_ptr, task_type, layer_id, task_id, tile_id_or_start,
    ↪ depend_entry_start, depend_entry_end, MAX_NUM_TENSOR_DIMS)
32     return task_base_info
33
34
35 @triton.jit
36 def MEGA_TRITON_KERNEL(
37     work_queue_start, # [1, ] int32, init with zero
38     work_queues, # [MAX_INS, NUM_SMS, INS], int32
39     num_tasks_per_wq, # [num_sms,]
40     scoreboard_ptr,

```

```

1728 task_deps_ptr, # [num_deps_entry_of_all_tasks, INT_PER_DEPS]
1729
1730 INT_PER_DEPS: tl.constexpr,
1731 INT_PER_TASK: tl.constexpr,
1732 MAX_TASK_ID: tl.constexpr,
1733 MAX_NUM_TILES_PER_OP: tl.constexpr,
1734 MAX_NUM_TENSOR_DIMS: tl.constexpr,
1735 NUM_SMS: tl.constexpr,
1736 num_warps: tl.constexpr
1737 ):
1738
1739 WARP_SIZE: tl.constexpr = 32
1740 NUM_THREADS: tl.constexpr = num_warps * WARP_SIZE
1741 scoreboard = Scoreboard(task_deps_ptr, INT_PER_DEPS, scoreboard_ptr, MAX_TASK_ID,
1742 ↪ MAX_NUM_TILES_PER_OP, tl.constexpr(1), NUM_THREADS)
1743 sm_id = tl.program_id(axis=0)
1744
1745 num_tasks = tl.load(num_tasks_per_wq + sm_id)
1746 cur_task_idx = 0
1747
1748 # early exit
1749 if cur_task_idx >= num_tasks:
1750     return
1751
1752 cur_task_base_info = FETCH_TASK(work_queues, cur_task_idx, INT_PER_TASK, NUM_SMS,
1753 ↪ MAX_NUM_TENSOR_DIMS, ENABLE_RUNTIME_SCHEDULER=False)
1754 nxt_task_base_info = cur_task_base_info
1755 cur_task_idx = cur_task_idx
1756
1757 while cur_task_idx < num_tasks:
1758
1759     task_type = cur_task_base_info.task_type
1760
1761     # task kernel need to set signal for each tile
1762
1763     scoreboard.wait_deps(cur_task_base_info)
1764
1765     ##### run task #####
1766     task_base_info = cur_task_base_info
1767
1768     if task_type == 0: # RMSNormTask
1769         rmsnorm_task_compute(task_base_info, scoreboard, RMS_EPS=1e-06, BLOCK_SIZE_N = 2048)
1770
1771     elif task_type == 6: # QKVProjTask
1772         fcl_task_compute(task_base_info, scoreboard, BLOCK_SIZE_M=16, BLOCK_SIZE_N=64,
1773             BLOCK_SIZE_K=256, NUM_STAGES=5)
1774
1775     elif task_type == 9: # QKNormRopeUpdateKVCacheTask
1776         rmsnorm_rope_update_kv_cache_task_compute(
1777             task_base_info, scoreboard, NUM_Q_HEADS=4, NUM_KV_HEADS=1, Q_HEAD_DIM=128,
1778             V_HEAD_DIM=128, PAGE_SIZE=1, MAX_NUM_BLOCKS_PER_SEQ=1024,
1779             Q_RMS_EPS=1e-06, K_RMS_EPS=1e-06
1780         )
1781
1782     elif task_type == 2: # AttnSplitTask
1783         attn_gqa_fwd_batch_decode_split_kv_task_compute(
1784             task_base_info, scoreboard, SM_SCALE=0.08838834764831845, SOFT_CAP=0.0,
1785             NUM_Q_HEADS=4, NUM_KV_HEADS=1, Q_HEAD_DIM=128, V_HEAD_DIM=128, PAGE_SIZE=1,
1786             MAX_NUM_BLOCKS_PER_SEQ=1024, BLOCK_N=64, BLOCK_HEAD_DIM=128,
1787             BLOCK_DPE=0, BLOCK_DV=128, BLOCK_H=16, NUM_KV_SPLITS=32
1788         )
1789
1790     elif task_type == 12: # AttnCombineTask
1791         attn_gqa_fwd_batch_decode_combine_task_compute(
1792             task_base_info, scoreboard, NUM_Q_HEADS=4, V_HEAD_DIM=128, BLOCK_DV=128,
1793             ↪ NUM_KV_SPLITS=32
1794         )
1795
1796     elif task_type == 4: # OProjTask
1797         fcl_task_compute(task_base_info, scoreboard, BLOCK_SIZE_M=16, BLOCK_SIZE_N=128,
1798             BLOCK_SIZE_K=64, NUM_STAGES=7)
1799
1800     elif task_type == 11: # BarrierAllIntraNodeTask
1801         barrier_all_intra_node_task_compute(task_base_info, scoreboard)
1802
1803     elif task_type == 7: # AllReduceTask
1804         allreduce_task_compute(task_base_info, scoreboard, BLOCK_SIZE=1024)
1805
1806     elif task_type == 8: # AddTask
1807         add_task_compute(task_base_info, scoreboard, BLOCK_SIZE=512)
1808
1809     elif task_type == 1: # MLPFC1Task
1810         fcl_task_compute(task_base_info, scoreboard, BLOCK_SIZE_M=16, BLOCK_SIZE_N=64,
1811             BLOCK_SIZE_K=128, NUM_STAGES=6)
1812
1813     elif task_type == 3: # SiLUMulUpTask
1814         silu_mul_up_task_compute(task_base_info, scoreboard, BLOCK_SIZE_M=8, BLOCK_SIZE_N=128)

```

```

132
133     elif task_type == 5: # MLPFC2Task
134         fcl_task_compute(task_base_info, scoreboard, BLOCK_SIZE_M=16, BLOCK_SIZE_N=64,
135                         BLOCK_SIZE_K=256, NUM_STAGES=6)
136
137     elif task_type == 10: # LinearTask
138         linear_task_compute(task_base_info, scoreboard, BLOCK_SIZE_M=16, BLOCK_SIZE_N=128,
139                           BLOCK_SIZE_K=128, NUM_STAGES=4, ALIGNMENT_K=16)
140
141
142     # nxt task
143
144     cur_task_idx = cur_task_idx + 1
145     if cur_task_idx < num_tasks:
146         cur_task_base_info = FETCH_TASK(work_queues, cur_task_idx, INT_PER_TASK, NUM_SMS,
147                                         ↳ MAX_NUM_TENSOR_DIMS, ENABLE_RUNTIME_SCHEDULER=False)
147

```