

PR #50383 完整报告

vllm-project/vllm

Shard the K3 Latent-MoE up-projection on large batches

合并时间: 2026-08-03 15:49

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/50383>

执行摘要

- 一句话: K3 大批量 up-proj 改列并行分片, TTFT 降约 4%
- 推荐动作: 值得精读, 特别是 `_select_tail_tier` 的分层设计 (按 token 数选择不同 collectives 与 kernel 策略) 与 `_shard_up_proj_tail` 中把列并行分片 fold 进最终 all-reduce 的技巧。建议在合并后补充针对三档边界的单元测试, 至少覆盖小批量 / 大批量切换与环境变量开关。

功能与动机

PR body 中给出 PERF 数据: 在 input 8K/output 1K 基准下, TTFT 各并发档位相对 main 下降约 4%, 而 TPOT/ITL 基本持平。这源于原实现中 replicated up-proj 在每个 rank 上计算完整 hidden dim, 大批量 (prefill) 时存在冗余计算; 通过列并行分片可将 up-proj FLOPs 降为原来的 $1/tp$ 。PR 同时依赖 #50000 (KimiK3LatentMoETailOp 融合算子), 并希望在其基础上扩展大批量路径。

实现拆解

1. 在 `vllm/model_executor/layers/fused_moe/runner/latent_moe_runner.py` 中新增 `LatentTailTier` 枚举 (`TAIL_FUSION / ALLREDUCE_OVERLAP / COLUMN_PARALLEL`), 并新增 `_select_tail_tier()` 按 `fused_output.shape[0]` (token 数) 选择实现: 小批量走 SM100 专属 CuTeDSL 融合算子; 不超过 `VLLM_SHARED_EXPERTS_STREAM_TOKEN_THRESHOLD` 且未禁用流重叠时走原 all-reduce + aux stream 重叠; 其余大批量走新增的 `_shard_up_proj_tail()`。
2. 将 `enable_k3_latent_moe_tail_fusion` 的平台判定 (`current_platform.is_cuda()` 且 `capability family 100`) 从 `vllm/models/kimi_k3/nvidia/model.py` 的构造参数 (`runner_args`) 移入 `LatentMoERunner.__init__`, 删除了 `model.py` 中对应逻辑和 `runner_args` 传参, 仅保留两处注释说明尺寸。
3. 在 `LatentMoERunner.__init__` 中新增 `_shared_ar_events` 事件对, 并用 `vllm.utils.multi_stream_utils.maybe_execute_in_parallel` 统一 `_overlap_allreduce_tail` 中的共享专家 all-reduce 与 up-proj GEMM 的流重叠, 替代原先手写的 `aux_stream + record_stream + wait_stream`。
4. 新增 `_shard_up_proj_tail()`: 先对 latent 做 norm + all-reduce, 再按 `tp_size` 用 `weight.narrow(0, shard_start, shard_size)` 切分 replicated up-proj 权重, 通过 `hidden_shard.addmm_(latent, up_proj_shard.t())` 把路由分片累加进共享输出 `partial`,

最后以 `output_is_reduced=False` 调用 `_maybe_reduce_final_output`, 让最终 all-reduce 同时完成跨 rank 求和与分片缝合。

5. 测试配套: 本 PR 未新增或修改测试文件, 完全依赖手工基准 (GSM8K、PERF) 验证正确性与性能。

关键文件:

- `vllm/model_executor/layers/fused_moe/runner/latent_moe_runner.py` (模块 MoE 执行器; 类别 source; 类型 core-logic; 符号 `LatentTailTier`, `_select_tail_tier`, `_small_batch_tail`, `_overlap_allreduce_tail`): 核心改动文件: 新增 `LatentTailTier` 三档枚举与 `_select_tail_tier` 分发, 新增 `_shard_up_proj_tail` 实现大批量列并行 up-proj, 重构 `_overlap_allreduce_tail` 的流重叠逻辑, 并将 tail-fusion 平台门控移入 `init`。
- `vllm/models/kimi_k3/nvidia/model.py` (模块 K3 模型; 类别 source; 类型 refactor; 符号 `KimiK3MegaMoEExperts.init`): 删除 `FusedMoEFactory` 调用中的 `runner_args`, 将 tail-fusion 平台门控下沉到 `runner`, 模型构造更简洁。

关键符号: `_select_tail_tier`, `_small_batch_tail`, `_overlap_allreduce_tail`, `_shard_up_proj_tail`, `_fused_forward`

关键源码片段

`vllm/model_executor/layers/fused_moe/runner/latent_moe_runner.py`

核心改动文件: 新增 `LatentTailTier` 三档枚举与 `_select_tail_tier` 分发, 新增 `_shard_up_proj_tail` 实现大批量列并行 up-proj, 重构 `_overlap_allreduce_tail` 的流重叠逻辑, 并将 tail-fusion 平台门控移入 `init`。

```
# 按 token 数选择当前 batch 的 tail 实现
# 三档共享同一个 replicated up-proj 权重, 选择只需看 token 数, 无需权重重排。
def _select_tail_tier(
    self,
    fused_output: torch.Tensor,
    shared_output: torch.Tensor,
) -> LatentTailTier:
    num_tokens = fused_output.shape[0]
    # tier 0: decode 规模的小 batch, 交给 CuTeDSL 融合算子
    # (内部完成 latent reduce、RMSNorm、shared reduce-scatter 与分片 up-proj 的 Lamport 广播),
    # 仅在 SM100 + TP 8/16 + BF16 下可用。
    if self.enable_k3_latent_moe_tail_fusion and (
        0 < num_tokens <= self._k3_latent_moe_tail_op.contract.max_num_tokens
    ):
        return LatentTailTier.TAIL_FUSION

    transform = self.routed_output_transform
    assert transform is not None
    # tier 1: 默认路径, batch 足够小时用 aux stream 把共享专家 all-reduce
    # 隐藏到 up-projection GEMM 后面。
    if (
```

```

        num_tokens <= envs.VLLM_SHARED_EXPERTS_STREAM_TOKEN_THRESHOLD
        and not envs.VLLM_DISABLE_SHARED_EXPERTS_STREAM
    ):
        return LatentTailTier.ALLREDUCE_OVERLAP
    # tier 2: prefill 规模的大 batch，走列并行分片 up-proj。
    return LatentTailTier.COLUMN_PARALLEL

def _shard_up_proj_tail(
    self,
    fused_output: torch.Tensor,
    shared_output: torch.Tensor,
    trunc_size: int | None,
) -> torch.Tensor:
    """Tier 2: 列并行 up-projection，并把路由部分累加进共享专家 partial。

    与 tier 1 相比，每个 rank 只算 1/tp 的 up-proj FLOPs，代价是放弃了
    aux stream 重叠：因为路由结果要先累加进 shared_output，最后的全减
    必须等累加完成。
    """
    transform = self.routed_output_transform
    assert transform is not None

    if transform.norm is not None:
        latent = self.allreduce_norm_latent_out(fused_output, transform.norm)
    else:
        latent = tensor_model_parallel_all_reduce(fused_output)

    weight = transform.up_proj.weight
    shard_size = weight.shape[0] // self.moe_config.tp_size
    shard_start = get_tensor_model_parallel_rank() * shard_size

    # 列并行切分：每个 rank 取自己的一段输出列与对应权重行。
    up_proj_shard = weight.narrow(0, shard_start, shard_size)
    hidden_shard = shared_output.narrow(-1, shard_start, shard_size)

    # 利用 GEMM 的 beta-add epilogue 把共享专家 partial 直接累加进结果，
    # 少一次独立的加法 kernel。
    hidden_shard.addmm_(latent, up_proj_shard.t())

    # 这里输出尚未全部 reduce，交给 _maybe_reduce_final_output 做跨 rank 求和，
    # 顺带把路由分片与共享专家 partial 一起缝合。
    return self._maybe_reduce_final_output(
        shared_output, trunc_size, output_is_reduced=False
    )

```

评论区精华

该 PR 没有实质性的技术 review 讨论。claude[bot] 自动评论提示该仓库配置为手动 review，未触发深度评审；人类评审者 ZJY0516 直接 APPROVED（无评论内容）。唯一 issue 评论来自 mergify[bot]，提示存在冲突需 rebase，最终已解决。因此本报告的主要洞察来自代码结构本身而非讨论过程。

- 自动评审提示与最终批准 (other): 无设计争议，评审通过。

风险与影响

- 风险:

1. 无任何测试文件变更，三档切换的边界条件 (token 数阈值、环境变量 VLLM_SHARED_EXPERTS_STREAM_TOKEN_THRESHOLD/VLLM_DISABLE_SHARED_EXPERTS_STREAM、TP 8/16 限制) 缺少自动化验证，容易回归。
2. LatentMoERunner.__init__ 无条件创建 torch.cuda.Event()，在非 CUDA 平台（如 CPU、ROCm 的非 CUDA 路径）上实例化该 runner 时可能报错，需确认所有调用方都运行在 CUDA 环境。
3. _shard_up_proj_tail 假设 transform.up_proj.weight.shape[0] 能被 moe_config.tp_size 整除，否则 narrow 会错位或越界；当前 K3 配置下成立，但代码未显式校验。
4. _shard_up_proj_tail 依赖 shared_output 是未 reduce 的跨 rank partial，且 addmm_ 在原缓冲上就地累加，若上游改变 reduce_results 契约会静默出错。
5. 本 PR 依赖 #50000，若合并顺序颠倒，tail-fusion 算子不可用时会回退到默认路径，功能正确但性能优化失效。- 影响：影响面集中在 Kimi-K3 模型的 NVIDIA SM100 (TP 8/16) 场景：TTFT 约提升 4%，TPOT 基本持平，GSM8K 准确率不变。对其他硬件（非 SM100、ROCm、CPU）和 TP 配置，由于平台门控回退，行为与之前一致。对代码结构的影响是让 LatentMoERunner 的 tail 策略更清晰、可扩展，后续其他 latent MoE 模型可复用同一套分层框架；团队协作上需要关注 #50000 的合并顺序。- 风险标记：核心路径变更，缺少测试覆盖，平台相关 CUDA 事件，依赖未合并 PR

关联脉络

- PR #50000 Dependency: K3 latent-MoE tail fusion (per PR body): PR body 明确声明本 PR 依赖 #50000，后者提供 KimiK3LatentMoETailOp 融合算子，本 PR 在其之上扩展大批量分片路径。
- PR #50761 [ROCm][Bugfix][Kimi-K3] Preserve MoE correction bias in FP32: 同为 Kimi-K3 MoE 路径的修复，虽侧重点在 AMD 侧精度，但与本次 NVIDIA 端性能优化共同构成 K3 MoE 的完整演进。