

PR #33021 完整报告

sgl-project/sglang

[AMD] Drop redundant FP8 bpresuffle scale transpose via fused AR kernel

合并时间: 2026-08-25 17:49

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/33021>

执行摘要

- 一句话: 融合 AR 内核直接产出转置 scale, 省去逐层额外拷贝 kernel。
- 推荐动作: 该 PR 值得精读, 尤其关注如何在融合内核中利用编译期模板参数来消除冗余内存操作。设计决策 (在 SGLang 侧透传参数而不是维护后处理) 体现了对底层 kernel 能力的充分利用。但考虑到 CI 验证不足和 aiter 依赖, 建议在合并后补充相应测试, 并确保 aiter 版本正确捆绑。

功能与动机

PR body 指出, 在 gfx95 上 bpresuffle FP8 block-scale GEMM 消费的是列主序 (column-major) 的 per-group activation scale, 而融合 AR+RMSNorm+per-group-FP8-quant kernel 输出的是行主序 (row-major)。为桥接两者, 原实现必须在每个符合条件的 layer 执行显式 post-kernel 转置 (`materialize_bpresuffle_fp8_scale`), 其中 `.contiguous()` 会启动一次额外拷贝 kernel, 这在以大量小 kernel 启动主导的 decode 延迟敏感阶段构成纯开销。aiter 的融合 kernel 已有编译期 `TRANSPOSE_SCALE` 模板 (ROCm/aiter#3652), 可以零额外成本直接写出 GEMM 所需的列主序 scale。

实现拆解

实现步骤:

1. 通信 API 透传 `transpose_scale` 参数: 在 `communication_op.py` 的 `tensor_model_parallel_fused_allreduce_rmsnorm_quant_per_group` 中新增 `transpose_scale: bool = False` 参数, 并转发给 `GroupCoordinator.fused_allreduce_rmsnorm_quant_per_group`。
2. `GroupCoordinator` 转发到 aiter 调用: 在 `parallel_state.py` 的 `fused_allreduce_rmsnorm_quant_per_group` 新增同名参数, 并在调用 `ca_comm.custom_fused_ar_rms_per_group_quant` 时转发 `transpose_scale=transpose_scale`。
3. 调用处启用参数并移除冗余转置: 在 `layernorm.py` 的 `_forward_with_allreduce_fusion_quant_per_group` 中, 对两条融合成功路径 (`keep_bf16=False` 的注意力和 `keep_bf16=True` 的 GDN `emit_bf16`) 调用时传入 `transpose_scale=use_bpresuffle`, 并移除对 `materialize_bpresuffle_fp8_scale` 的调用。

4. 保留 fallback 路径的转置：两条独立 per-group quant 的 fallback 路径（per_1x128_quant）继续使用 materialize_bpreshuffle_fp8_scale，因为 per_1x128_quant 的 transpose_scale 会产生不同的物理排列。
5. FP8 线性层无需改动：aiter_w8a8_block_fp8_linear 消费的 scale 布局不变，其自身的 materialize 调用对已是列主序的张量执行 .contiguous() 会短路，变成幂等 no-op。测试与配套：PR 未新增单元测试，但提供了详尽的准确性和性能基准数据。

关键文件：

- python/sglang/srt/layers/layernorm.py（模块 层实现；类别 source；类型 core-logic）：核心逻辑变更：在融合 AR+RMSNorm+per-group-quant 调用中传入 transpose_scale=use_bpreshuffle，并移除两个融合成功路径上的显式 scale 转置。
- python/sglang/srt/distributed/communication_op.py（模块 通信层；类别 source；类型 core-logic）：公共 API 新增 transpose_scale 参数并透传，是链接 layernorm 与 parallel_state 的桥梁。
- python/sglang/srt/distributed/parallel_state.py（模块 并行状态；类别 source；类型 core-logic）：GroupCoordinator 方法新增 transpose_scale 并最终转发给 aiter 的 custom_fused_ar_rms_per_group_quant 调用。

关键符号：tensor_model_parallel_fused_allreduce_rmsnorm_quant_per_group, fused_allreduce_rmsnorm_quant_per_group, _forward_with_allreduce_fusion_quant_per_group

关键源码片段

python/sglang/srt/layers/layernorm.py

核心逻辑变更：在融合 AR+RMSNorm+per-group-quant 调用中传入 `transpose_scale=use_bpreshuffle`，并移除两个融合成功路径上的显式 scale 转置。

```
# python/sglang/srt/layers/layernorm.py
# 以下为修改后的 `_forward_with_allreduce_fusion_quant_per_group` 核心分支（节选）

# `transpose_scale=use_bpreshuffle` 让融合内核直接以 gfx95 bpreshuffle GEMM
# 所需的列主序 (column-major) 布局写出 per-group scale，与 `materialize_bpreshuffle_fp8_scale`
# 结果一致，因此融合成功路径无需再执行 post-kernel 转置。
# 注意：独立 per-group-quant 的 fallback 路径仍需 `materialize_bpreshuffle_fp8_scale`，
# 因为 `per_1x128_quant` 的 `transpose_scale` 是字节置换，排列不同。
if not keep_bf16:
    result = tensor_model_parallel_fused_allreduce_rmsnorm_quant_per_group(
        x, residual, weight, norm_module.variance_epsilon, group_size,
        transpose_scale=use_bpreshuffle,
    )
    if result is not None:
        fp8_out, residual_out, scale_out = result
        return (fp8_out, scale_out), residual_out

# fallback: 融合 AR+RMSNorm，再单独做 per-group quant
```

```

fused_result = tensor_model_parallel_fused_allreduce_rmsnorm(
    x, residual, weight, norm_module.variance_epsilon
)
if fused_result is None:
    return None
bf16_out, residual_out = fused_result
per_1x128_quant, fp8_dtype = _get_aiter_per_group_quant()
fp8_out, scale_out = per_1x128_quant(
    bf16_out, quant_dtype=fp8_dtype, transpose_scale=False
)
if use_bpreshuffle:
    scale_out = materialize_bpreshuffle_fp8_scale(scale_out)
return (fp8_out, scale_out), residual_out

# keep_bf16=True 的 GDN 路径: 需要同时输出 bf16 和 fp8+scale
result = tensor_model_parallel_fused_allreduce_rmsnorm_quant_per_group(
    x, residual, weight, norm_module.variance_epsilon, group_size,
    emit_bf16=True, transpose_scale=use_bpreshuffle,
)
if result is not None and len(result) == 4:
    fp8_out, residual_out, scale_out, bf16_out = result
    return (bf16_out, fp8_out, scale_out), residual_out

# 同样, fallback 路径保留 materialize

```

python/sglang/srt/distributed/communication_op.py

公共 API 新增 **transpose_scale** 参数并透传, 是链接 layernorm 与 parallel_state 的桥梁。

```

# python/sglang/srt/distributed/communication_op.py
# 新增参数并透传, 保持公共 API 与底层实现一致。
def tensor_model_parallel_fused_allreduce_rmsnorm_quant_per_group(
    input_: torch.Tensor,
    residual_inp_: torch.Tensor,
    weight_: torch.Tensor,
    eps: float,
    group_size: int = 128,
    emit_bf16: bool = False,
    transpose_scale: bool = False,
) -> Optional[Tuple[torch.Tensor, ...]]:
    """Fused TP all-reduce + RMSNorm + per-group FP8 quant (ROCm/aiter)。

    `transpose_scale=True` 时内核直接输出 gfx95 bpreshuffle GEMM 所需的
    列主序 (column-major) scale, 调用方无需再做 post-kernel 转置。
    """
    return get_tp_group().fused_allreduce_rmsnorm_quant_per_group(
        input_, residual_inp_, weight_, eps, group_size,
        emit_bf16=emit_bf16, transpose_scale=transpose_scale,
    )

```

python/sglang/srt/distributed/parallel_state.py

GroupCoordinator 方法新增 `transpose_scale` 并最终转发给 `aiter` 的 `custom_fused_ar_rms_per_group_quant` 调用。

```
# python/sglang/srt/distributed/parallel_state.py
# GroupCoordinator 方法，负责在 gfx95 平台调度融合 AR+RMSNorm+per-group-quant。
def fused_allreduce_rmsnorm_quant_per_group(
    self,
    input_: torch.Tensor,
    residual_inp_: torch.Tensor,
    weight_: torch.Tensor,
    eps: float,
    group_size: int = 128,
    emit_bf16: bool = False,
    transpose_scale: bool = False,
) -> Optional[Tuple[torch.Tensor, ...]]:
    """gfx95 专用入口，非 HIP 或非 gfx95 返回 `None` 以触发 fallback。

    `transpose_scale=True` 时由 aiter 内核直接输出列主序 scale。
    """
    if not (is_hip() and is_gfx95_supported()):
        return None
    ca_comm = self.ca_comm
    if ca_comm is None or getattr(ca_comm, "disabled", True):
        return None
    if not hasattr(ca_comm, "custom_fused_ar_rms_per_group_quant"):
        return None
    # ... 形状和大小合法性检查 ...
    try:
        return ca_comm.custom_fused_ar_rms_per_group_quant(
            input_, residual_inp_, weight_, eps, group_size, use_1stage_ar,
            emit_bf16=emit_bf16, transpose_scale=transpose_scale,
        )
    except Exception:
        return None # 任何异常回退到独立 per-group quant 路径
```

评论区精华

该 PR 的 review 讨论主要集中在 CI 验证状态上。AMD bot 指出，PR 修改的代码路径未被任何 PR CI 测试覆盖，因为该路径在 MI300 运行器上被硬件门控关闭 (`is_gfx95_supported()` 返回 `False`)，且唯一执行的 MI35x 任务测试的是 **Qwen3-8B** (BF16) 且 TP=1，与修改路径无关。NVIDIA 及其它后端流水线因缺少 `run-ci` 标签而被门控关闭。作者 yichiche 在最终评论中澄清，NVIDIA B200 测试失败是无关的 **HFValidationError** (已在 PR #35607 修复)，AMD CI 失败也与本 PR 无关，因此认为 PR 已就绪可合并。hubertlu-tw 和 HaiShaw 均批准 (LGTM)。

- CI 验证不足：修改路径未被任何测试覆盖 (testing): 作者 yichiche 确认所有 CI 失败均与本 PR 无关，且提供了充分的本地基准与正确性数据，评审者认可后合并。

风险与影响

- 风险:

1. 数值一致性风险: PR 声称产出布局与 `materialize_bpreshuffle_fp8_scale` 字节级一致, 但该声明依赖 `aiter` 端 `TRANPOSE_SCALE` 模板的正确实现。若 `aiter` 版本与 PR 不匹配, 可能导致数值偏差或错误, PR 通过异常捕获回退但无法彻底规避。
2. 回归风险: 修改仅影响 `gfx95` 且启用 `aiter` 融合的 FP8 路径 (`use_bpreshuffle` 为 `True` 时), 其它路径 (如 `MI300`、`MFXP4`、未启用融合) 不受影响。但 `use_bpreshuffle` 的判定逻辑在 `fp8_utils.py` 中, 若该逻辑有误, 可能影响路径选择。
3. 兼容性风险: 依赖 `aiter` 的更新 (`ROCm/aiter#4478` 和 `#3652`), 若 `aiter` 版本未同步, 新参数会被忽略或抛出异常。PR 中已通过 `except Exception` 捕获异常并返回 `None`, 触发 `fallback`, 但若 `aiter` 未实现该参数, 则 `transpose_scale` 参数无法转发, 可能导致静默忽略, 无法达到预期优化。
4. 测试覆盖缺口: PR 未新增针对该路径的单元测试, 且 CI 未能实际验证 `gfx95` 上的行为, 存在潜在未覆盖风险。
 - 影响: 影响范围: 对 AMD `gfx95` (`MI35x`) 平台上启用 `--enable-aiter-allreduce-fusion` 的 FP8 模型, 在 `decode` 路径上每个符合条件的 `layer` 减少一次 `kernel` 启动 (每次融合调用一次 `.contiguous()`)。基准显示吞吐提升约 `0.25%~1.25%` (`geomean +0.84%`), `TPOT` 降低 `0.09%~1.67%`。对其它平台 (`NVIDIA`、`MI300`、`MFXP4`) 无影响, 因为改动被硬件门控 (`is_gfx95_supported()`) 限制。用户无需额外配置, 功能自动生效。团队方面, 为后续 `aiter` 内核参数透传提供了模式参考。
 - 风险标记: 缺少测试覆盖, 依赖外部内核改动, 数值一致性敏感

关联脉络

- PR `#3652` `make allreduce_rmsnorm_quant` support `transpose_scale` for `bpreshuffle_gemm`: 该 PR 在 `ROCm/aiter` 中实现了 `transpose_scale` 编译期模板支持, 直接支撑本 PR 的功能, 是本 PR 的底层依赖。
- PR `#35607` `Fix local cache path passed as HuggingFace repo id`: 本 PR 作者在 CI 评论中提及该 PR 修复了 `B200` 测试中的 `HFValidationError`, 与本 PR 的 CI 失败归类相关。