

PR #32248 完整报告

sgl-project/sglang

Migrate CompressedTensorsW4A4Nvfp4MoE TRT-LLM path onto MoeRunner

合并时间: 2026-07-25 11:57

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/32248>

执行摘要

- 一句话: 将 W4A4Nvfp4MoE TRT-LLM 路径迁移至 MoeRunner, 移除内联调度, 净减 62 行。
- 推荐动作: 值得精读, 特别是理解如何将硬编码的内核调度收敛到 runner 抽象模式。设计决策关注点: runner 的注册机制 (@register_fused_func)、FlashInferTrtllmFp4MoeQuantInfo 数据类的作用、以及 shape 切片的兼容性处理。

功能与动机

此前 `use_flashinfer_trtllm` 分支创建了一个从未使用的 `MoeRunner(TRITON, ...)`, 并在 `apply_weights` 中硬编码了 `trtllm_fp4_block_scale_moe` 内核调用及 FP4 量化、对称内存分配等逻辑, 绕过了 runner 抽象 (见 issue #20719)。这与 `ModelOptNvFp4FusedMoEMethod` 的实现不一致, 后者已使用共享的 `MoeRunner(FLASHINFERRTLLM)` 路径。为了统一架构、减少重复代码并便于后续优化, 需要将自定义调度迁移到 runner。

实现拆解

1. 精简导入与清理: 移除不再需要的 `get_tp_group`、`use_symmetric_memory`、`is_allocation_symmetric`、`RoutingMethodType`、`next_power_of_2` 等导入, 仅保留必要依赖。
2. 修改 `create_moe_runner`: 将 `use_flashinfer_trtllm` 分支的 runner 从 `MoeRunner(MoeRunnerBackend.TRITON)` 改为 `MoeRunner(MoeRunnerBackend.FLASHINFERRTLLM)`, 并导入 `moe_runner.flashinfer_trtllm` 以触发 `@register_fused_func` 装饰器, 将 FP4 fused 函数注册到 runner。
3. 重写 `apply_weights` 的 TRT-LLM 分支: 放弃内联的 `fp4_quantize` 和 `trtllm_fp4_block_scale_moe` 调用, 改为构造 `FlashInferTrtllmFp4MoeQuantInfo` 对象并调用 `self.runner.run(dispatch_output, quant_info)`。runner 内部承担 FP4 隐藏状态量化、对称内存输出分配和内核调度。
4. 修复 `global_scale shape`: `w13_input_scale_quant` 原形状为 `[num_local_experts]`, 但 `cute-dsl` 后端要求 `[1]`, 因此传入 `[:1]` 切片, 与 `ModelOptNvFp4FusedMoEMethod` 一致。同时移除了 `correction_bias` 的显式处理 (已由 runner 管理)。

5. 删除无用代码：移除约 80 行内联调度代码和 5 个无用导入，净减 62 行。

关键文件：

- python/sglang/srt/layers/quantization/compressed_tensors/schemes/compressed_tensors_w4a4_nvfp4_moe.py（模块 量化层；类别 source；类型 refactor；符号 create_moe_runner, apply_weights, init）：唯一变更的文件，实现了将硬编码的 TRT-LLM 内核调度迁移到 MoeRunner 的核心重构，包含 create_moe_runner 和 apply_weights 的修改。

关键符号：create_moe_runner, apply_weights

关键源码片段

python/sglang/srt/layers/quantization/compressed_tensors/schemes/compressed_tensors_w4a4_nvfp4_moe.py

唯一变更的文件，实现了将硬编码的 TRT-LLM 内核调度迁移到 MoeRunner 的核心重构，包含 create_moe_runner 和 apply_weights 的修改。

```
def create_moe_runner(self, layer, moe_runner_config):
    self.moe_runner_config = moe_runner_config
    if self.use_flashinfer_trtllm:
        # 导入会触发 @register_fused_func 装饰器，注册 FP4 fused 函数到 runner
        import sglang.srt.layers.moe.moe_runner.flashinfer_trtllm # noqa: F401
        self.runner = MoeRunner(
            MoeRunnerBackend.FLASHINFER_TRTLLM, moe_runner_config
        )
    else:
        import sglang.srt.layers.moe.moe_runner.flashinfer_cutlass # noqa: F401
        self.runner = MoeRunner(
            MoeRunnerBackend.FLASHINFER_CUTLASS, moe_runner_config
        )

def apply_weights(self, layer, dispatch_output):
    x = dispatch_output.hidden_states
    if self.use_flashinfer_trtllm:
        from sglang.srt.layers.moe.moe_runner.flashinfer_trtllm import (
            FlashInferTrtllmFp4MoeQuantInfo,
        )
        # global_scale 形状必须为 [1]，从 [num_local_experts] 切片
        quant_info = FlashInferTrtllmFp4MoeQuantInfo(
            w13_weight=layer.w13_weight,
            w2_weight=layer.w2_weight,
            w13_weight_scale=layer.w13_weight_scale,
            w2_weight_scale=layer.w2_weight_scale,
            g1_scale_c=layer.g1_scale_c,
            g1_alphas=layer.g1_alphas,
            g2_alphas=layer.g2_alphas,
            w13_input_scale_quant=layer.w13_input_scale_quant[:1], # 关键修复
```

```
        global_num_experts=layer.num_experts,
        local_expert_offset=layer.moe_ep_rank * layer.num_local_experts,
        local_num_experts=layer.num_local_experts,
        intermediate_size_per_partition=layer.intermediate_size_per_partition,
    )
    # 委托给 runner 执行 MoE 前向
    return self.runner.run(dispatch_output, quant_info)
else:
    # FLASHINFER_CUTLASS 路径保持不变
    ...
```

评论区精华

审核者 b8zhong 批准了变更，未提出异议。作者在合并前评论说 'nvfp4 model works and the tp8 + mtp failure is unrelated', 表明已验证 B200 比特一致，并指出 CI 中一个不相关测试失败不影响合并。此外，通过 rerun-test 请求触发 Mistral Large 3 测试，B200 上失败（不相关），H200 通过。

- 验证正确性与 CI 结果 (testing): 确认重构后输出比特一致，CI 中不相关测试失败不影响合并。

风险与影响

- 风险：重构目标为比特一致，已在 B200 验证，回归风险极低。但存在潜在风险：1) 如果未来 MoeRunner(FLASHINFER_TRTLLM) 的 run 方法行为有变，可能影响此路径；2) 移除了 correction_bias 的显式处理，需确认 runner 内部正确处理；3) 删除的导入（如 use_symmetric_memory）在其他地方是否仍被依赖？已确认仅在 apply_weights 中使用，故移除安全。整体风险低。
- 影响：对用户：输出与之前完全一致（比特一致），无行为变化。对系统：减少了代码体积和复杂性，统一了 MoE 路由架构，便于后续维护和优化。对团队：所有 TRT-LLM FP4 MoE 量化方案均通过统一的 MoeRunner 调度，降低了认知负担。
- 风险标记：暂无

关联脉络

- 暂无明显关联 PR