

PR #34651 完整报告

sgl-project/sglang

[DCP] Share one pack kernel between both a2a backends

合并时间: 2026-08-14 06:06

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/34651>

执行摘要

- 一句话: 统一 DCP 两条 a2a 后端的 pack kernel, fi_a2a 省 4 次拷贝
- 推荐动作: 值得精读。亮点不在改动规模, 而在两个设计决策: 其一, 用 stride 参数而不是特化 kernel 统一两种互斥布局, 避免两套实现各修各的; 其二, 对 empty() 的论证——以 "字段是否为不透明字节" 判断初始化是否必要, 是减少冗余 kernel 工作的可复用思路。commit 历史还展示了一次完整的方案反转 (持久缓冲区缓存 → 每次调用 empty()), 对理解 CUDA graph capture 约束有参考价值。

功能与动机

PR body 明确指出: #34614 fused the pack/unpack copies on the pynccla2apath but left fi_a2a(FlashInfer MNNVL) untouched, so it still paid four materializing copies plus a zero-fill per MLA layer per decode step. #34614 造成了不对称——在此之前两条路径都做约 4 次拷贝; 本 PR 的目标是 "finishes the job rather than duplicating the fix", 避免同一修复在两条后端上各写一份。

实现拆解

变更入口是 `python/sglang/srt/layers/dcp/comm.py` 中的 `dcp_a2a_lse_reduce` (pynccl 路径) 与 `_dcp_fi_a2a_lse_reduce` (FlashInfer MNNVL 路径), 核心 kernel 位于 `python/sglang/kernels/ops/attention/dcp_kernels.py`。

1. 改造 pack kernel 签名: `dcp_pack_a2a_send` 从接收单个 `send_combined` 交错 buffer 改为接收 `dst_o` 与 `dst_lse` 两个独立目标, 各自携带 N/B/H 三个维度的 stride; `_dcp_pack_a2a_send_kernel` 内部按 `peer = h // H_PER_RANK`、`h_local = h % H_PER_RANK` 计算落点。peer 轴在外层 (pynccl) 还是内层 (FlashInfer)、LSE 是否交错在行尾, 都只是 stride 差异, 因此单个 kernel 覆盖两种传输且都不再需要布局拷贝。
2. 更新 pynccl 路径: `dcp_a2a_lse_reduce` 仍使用预分配的 `send_combined [N, B, H_per_rank, D + lpd]`, 但把 payload 视图 `send_combined[:, :, :, :D]` 与 LSE 视图 `send_words[:, :, :, D // lpd]` 分别传给 kernel, 保持 "单次 all_to_all_single 按字节偏移切分扁平 buffer" 的交错布局不变; 接收侧逻辑完全未动。
3. 重写 fi_a2a 路径: `_dcp_fi_a2a_lse_reduce` 删掉 `partial_o/softmax_stats` 的两次构造拷贝和接收侧两次 `contiguous()`, 改为把 FlashInfer 需要的 peer-inside 视图 (`partial_o.permute(2, 0, 1, 3)` 与 `softmax_stats[..., 0].permute(2, 0, 1)`) 直接交给共享

kernel scatter。softmax_stats 由 zeros() 改为 empty(), 论证是 slot 1 从不被读取、交换只按不透明字节搬运该字段。接收侧 o_out/stats_out 的非连续视图直接送 dcp_lse_combine_triton (它支持任意步幅)。

4. 配套测试: test/registered/kernels/test_dcp_lse_combine.py 更新旧用例适配新签名, 并新增 test_pack_serves_the_split_peer_inside_layout, 用 uint8 视图逐 bit 断言 FlashInfer 的 split/peer-inside 布局与手工构造一致, 同时验证 stats slot 1 保持零。fi_a2a 依赖 MNNVL 硬件、CI 无法覆盖, 此测试在 host 上固定住布局契约。

关键文件:

- python/sglang/srt/layers/dcp/comm.py (模块 通信层; 类别 source; 类型 core-logic; 符号 dcp_a2a_lse_reduce, _dcp_fi_a2a_lse_reduce): DCP a2a 主控逻辑所在: pynccl 组合路径与 fi_a2a 路径都在此分派; 改动让两条路径共用 pack kernel 并按各自 stride 直写目标, 同时删除 fi_a2a 的 4 次物化拷贝与零填充。
- python/sglang/kernels/ops/attention/dcp_kernels.py (模块 内核模块; 类别 source; 类型 core-logic; 符号 _dcp_pack_a2a_send_kernel, dcp_pack_a2a_send): 共享 pack kernel 本体: 从单一交错 send buffer 改为带独立步幅的 dst_o / dst_lse, 用步幅参数统一 pynccl 与 FlashInfer 两种相反布局。
- test/registered/kernels/test_dcp_lse_combine.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 test_pack_serves_the_split_peer_inside_layout): 更新 pack 测试适配新签名, 并新增 peer-inside 布局的逐 bit 断言, 为无法在 CI 覆盖的 fi_a2a 路径提供正确性护栏。

关键符号: dcp_a2a_lse_reduce, _dcp_fi_a2a_lse_reduce, dcp_pack_a2a_send, _dcp_pack_a2a_send_kernel, test_pack_serves_the_split_peer_inside_layout

关键源码片段

python/sglang/srt/layers/dcp/comm.py

DCP a2a 主控逻辑所在: pynccl 组合路径与 fi_a2a 路径都在此分派; 改动让两条路径共用 pack kernel 并按各自 stride 直写目标, 同时删除 fi_a2a 的 4 次物化拷贝与零填充。

```
def _dcp_fi_a2a_lse_reduce(
    cp_attn_out: torch.Tensor,
    cp_attn_lse: torch.Tensor,
    cp_group: "GroupCoordinator",
    is_lse_base_on_e: bool = True,
) -> torch.Tensor:
    """fi_a2a: 把跨 rank 交换交给 FlashInfer 的 MNNVL kernel, 组合仍复用本地 Triton。
```

FlashInfer 需要独立的 partial_o [B, H_per_rank, cp_size, D] (peer 轴在 heads 内层) 和 softmax_stats [B, H_per_rank, cp_size, 2] fp32 (S 补齐到 2)。
这里用共享的 dcp_pack_a2a_send 直接把输入散列到这两个布局, 不做物化拷贝。

```
"""
```

```
from flashinfer.comm.dcp_alltoall import decode_cp_a2a_alltoall
```

```
state = _FI_A2A_STATE
```

```

assert state is not None, (
    "fi_a2a workspace not initialized — call init_fi_a2a_workspace(dcp_group) "
    "at model-runner init (before CUDA graph capture)."
)

N = cp_group.world_size
B, H, D = cp_attn_out.shape
assert H % N == 0, f"num_heads ({H}) must be divisible by dcp_size ({N})"
H_per_rank = H // N

# 用 empty() 而非 zeros(): pack 会填满 partial_o 和 stats 的 slot 0,
# 而 slot 1 没有任何代码读取——交换把 stats 当作不透明字节搬运,
# 我们只读回 stats_out[..., 0]。零填充从未被需要过。
partial_o = torch.empty(
    B, H_per_rank, N, D, dtype=cp_attn_out.dtype, device=cp_attn_out.device
)
softmax_stats = torch.empty(
    B, H_per_rank, N, 2, dtype=torch.float32, device=cp_attn_out.device
)

# peer 轴在 heads 内侧: 把 [N, B, H_pr, D] 的非连续视图交给 kernel,
# 步幅参数会处理非连续写入, 接收端也不需要 contiguous()。
dcp_pack_a2a_send(
    cp_attn_out,
    cp_attn_lse,
    partial_o.permute(2, 0, 1, 3),
    softmax_stats[..., 0].permute(2, 0, 1),
)

o_out, stats_out = decode_cp_a2a_alltoall(
    partial_o,
    softmax_stats,
    state["workspace"],
    state["cp_rank"],
    N,
)

# dcp_lse_combine_triton 接受任意步幅, 非连续视图直接就地读取。
recv_output = o_out.permute(2, 0, 1, 3)
recv_lse = stats_out[..., 0].permute(2, 0, 1)

combined, _ = dcp_lse_combine_triton(
    recv_output, recv_lse, is_lse_base_on_e=is_lse_base_on_e
)
return combined

```

[test/registered/kernels/test_dcp_lse_combine.py](#)

更新 pack 测试适配新签名，并新增 peer-inside 布局的逐 bit 断言，为无法在 CI 覆盖的 fi_a2a 路径提供正确性护栏。

```
def test_pack_serves_the_split_peer_inside_layout(self):
    # FlashInfer MNNVL 要求 peer 轴在 heads 内侧且 stats 独立存放；
    # 本用例以手工构造的 [B, H_pr, N, ...] 布局为 oracle，逐 bit 校验共享 kernel。
    from sglang.kernels.ops.attention.dcp_kernels import dcp_pack_a2a_send

    for N, B, H_per_rank, D in ((2, 4, 8, 128), (4, 1, 16, 512)):
        with self.subTest(N=N, B=B, H_per_rank=H_per_rank, D=D):
            H = H_per_rank * N
            out = torch.randn(B, H, D, device=self.device, dtype=torch.bfloat16)
            lse = torch.randn(B, H, device=self.device, dtype=torch.float32)

            partial_o = torch.empty(
                B, H_per_rank, N, D, dtype=torch.bfloat16, device=self.device
            )
            stats = torch.zeros(
                B, H_per_rank, N, 2, dtype=torch.float32, device=self.device
            )
            dcp_pack_a2a_send(
                out,
                lse,
                partial_o.permute(2, 0, 1, 3), # [N, B, H_pr, D] 非连续视图
                stats[..., 0].permute(2, 0, 1), # [N, B, H_pr] 非连续视图
            )

            want_o = out.view(B, N, H_per_rank, D).permute(0, 2, 1, 3)
            want_lse = lse.view(B, N, H_per_rank).permute(0, 2, 1)
            # 用 uint8 视图比较保证字节级一致（与 fp8 场景同口径）。
            self.assertTrue(
                torch.equal(
                    partial_o.view(torch.uint8),
                    want_o.contiguous().view(torch.uint8),
                )
            )
            self.assertTrue(torch.equal(stats[..., 0], want_lse))
            # slot 1 从未被写入也没有人读取，应保持初始零值。
            self.assertTrue(
                torch.equal(stats[..., 1], torch.zeros_like(stats[..., 1]))
            )
```

评论区精华

该 PR 的 4 条 review 评论全部来自作者本人，属于典型的自审迭代：

- 关于 _alloc-fi_a2a_send 辅助函数，作者先写 "Actually I want to note this part more", 随后又自我质疑 "This is maybe not that useful?". 早期提交确实在 workspace init 时预分配 fi_a2a 发送缓冲区以避免 CUDA graph capture 期间分配，但最终提交（Drop the

fi_a2a send-buffer cache; allocate per call with empty()) 推翻了该方案：持久缓冲区只为摊销 softmax_stats 的零填充，而零填充从来不需要，于是每次调用直接 torch.empty() 反而更简单。

- 在 base_runner.py 的 _pre_initialize-fi_a2a_workspace docstring 上作者直接批示 "This is obvious. Remove", 随后多次提交裁剪注释 (Trim the comments、Fold the two notes into one) , 最终只保留两条承重注释: empty() 的原因、交换发生在 decode 图内。
- 对 "dtype is None when the caller could not derive the shape (non-MLA)" 的注释留下疑问 "?", 后续合并精简。

没有外部 reviewer 提出异议，合并者是作者本人。讨论的价值在于展示了 " 用步幅统一布局 " 和 " 删掉从未被读的初始化 " 这两个判断如何被反复验证和收敛。

- fi_a2a 发送缓冲区的分配与生命周期 (design): 最终提交 bb4cb4f 删除了 helper 与缓冲区缓存，改为每次调用 torch.empty() 内联分配：持久缓冲区只为摊销 softmax_stats 的零填充，而零填充从来不需要 (slot 1 无人读) 。
- 注释与 docstring 裁剪 (documentation): 采纳：大量自解释注释被删除，只剩 empty() 原因与 CUDA graph 交换位置两条关键说明。
- 非 MLA 场景下的形状推导注释 (question): 相关注释被整合或删除，未引入额外逻辑变更。

风险与影响

- 风险：
 - fi_a2a 无 CI 覆盖：fi_a2a 需要真实 MNNVL 交换网络，CI 只能跑 pyncccl 路径与布局单测；一旦 FlashInfer 的 decode_cp_a2a_alltoall 布局契约改变，回归只能靠布局断言的间接暴露。
 - empty() 依赖隐含不变量：softmax_stats[..., 1] 永不被读取、a2a 把 stats 当作不透明字节搬运。若未来 FlashInfer 版本读写 slot 1，或接入非 FlashInfer 的 fi 后端，未初始化值会沿下游传播。
 - pyncccl 路径按 fp32 word 重解释写入：依赖 $D \% \text{lpd} == 0$ 与 word 对齐，逻辑已有校验，但未来引入其他 dtype (如 fp8) 时需重新走读。
 - 精度验证单点：最大误差 $7.8e-3$ 相对 fp32 参考，属 bf16 舍入而非漂移，但只覆盖了 B=8、H_per_rank=12、D=512 的单一配置 (PR body 自述) 。
- 影响：
 - 性能影响：fi_a2a 下每个 MLA 层 decode 步从 4 次物化拷贝 + 1 次零填充降为 0 次额外物化；对 DeepSeek-V3.1 这类 MLA + DCP 模型，decode 步的 a2a 交换是热点路径之一，PR body 中的 profile 前后截图对比明显。pyncccl 路径行为不变，仅内部改用共享 kernel，无回退预期。
 - 用户影响：主要惠及使用 MNNVL 交换网络 + fi_a2a 后端 + DCP 的用户；pyncccl 用户几乎无感。
 - 工程影响：消除了两条路径的布局拷贝差异，未来新增传输后端只需提供 stride 约定即可复用同一 pack kernel；新增单测成为布局契约的护栏。

- 风险标记: fi_a2a 路径无 CI 覆盖, 依赖 stats slot 1 未被读取的不变量, 核心 decode 路径变更, 精度验证为单机单配置

关联脉络

- PR #34614 前序 PR: pynccl a2a 路径融合 pack/unpack (标题以仓库记录为准): PR body 明确引用: 本 PR 是 #34614 的续作, 后者在 pynccl a2a 路径融合了 pack/unpack 拷贝但留下 fi_a2a 未处理, 本 PR 通过共享 pack kernel 补齐对称性。