

PR #30700 完整报告

sgl-project/sglang

[NVIDIA] Add flashinfer MNNVL backend for allreduce only

合并时间: 2026-08-12 07:46

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/30700>

执行摘要

- 一句话: FlashInfer 纯 allreduce 后端, 复用融合 workspace 加速 DSV3/V4
- 推荐动作: 值得精读, 尤其是以下三个设计决策: 1) `can_use_flashinfer_allreduce` 与 `flashinfer_allreduce` 的拆分——如何在 Dynamo 符号执行下保持判定 rank-invariant 且不引入 Dynamo guard 重编译; 2) workspace 与进程组精确绑定的必要性——hybrid EP+TP 下同 world size 不同 rank 集合会静默产生错误归约, 这是分布式正确性中易被忽视的细节; 3) custom op 包装对 piecewise CUDA graph 的支持方式。建议结合 PR#32330 的 SM120 修复与 PR#32339 的性能对比一起阅读, 可形成对多节点 allreduce 路径选型的完整认识。

功能与动机

PR body 指出: SGLang 已支持通过 FlashInfer 的 `kARResidualRMSNorm` 做 fused allreduce (allreduce + residual + RMSNorm 单 kernel), 但并非每个 allreduce 点都能走融合路径——`RowParallelLinear`、attention 输出 fallback、MoE TP fallback 都是纯 (非融合) allreduce, 目前即使 FlashInfer mnnvl/trtllm workspace 已初始化可用也总是回退到 NCCL/custom-allreduce。本 PR 将这类纯 allreduce 路由到 FlashInfer `AllReduceFusionPattern.kAllReduce`, 复用现有 workspace, 无额外内存开销。nvpohanh 在评论中补充了关键背景: DSV4 因 mHC 无法做 AR+Norm 融合, 此前 SGLang 只对 DSV3 支持 FlashInfer MNNVL AR, 本 PR 补齐 DSV4 支持, 可显著提升 DSV4 在 GB200 TP8 小并发场景的性能。

实现拆解

实现按以下 5 步拆解:

1. 新增纯 allreduce 内核封装: 在 `python/sglang/srt/layers/flashinfer_comm_fusion.py` 中新增 `can_use_flashinfer_allreduce()` 与 `flashinfer_allreduce()` 两个函数。前者做全部可用性判定 (FlashInfer 可用性、2D 连续形状、workspace 初始化、world size 与 group 精确匹配、buffer 容量), 后者只负责调用 `allreduce_fusion(pattern=kAllReduce)`, 且刻意不捕获 kernel 异常——吞异常会让本 rank 回退 NCCL 而其他 rank 仍在 kernel 内, 导致集体通信失配挂起。判定与执行分离是为了适配 Dynamo: 判定在普通 Python 中于 trace 期完成, 执行注册为 opaque custom op 可进入 piecewise CUDA graph。
2. 并行状态层 dispatch 与 group 打标: 在 `python/sglang/srt/distributed/parallel_state.py` 中注册 `flashinfer_allreduce(Tensor, group_name)` custom op, 并在

- `GroupCoordinator.all_reduce()` 的编译与非编译两条路径上、在既有 NCCL/custom-allreduce 链之前插入 dispatch。新增 `_tag_groups_for_flashinfer_allreduce_only()`：只在真正拥有 fusion workspace 的组上打 `_fi_workspace_hint` 标记（attention-TP 一个、MoE 一个；MoE workspace 在 `moe_ep_size > 1` 时 rendezvous 于 EP 组，否则于 MoE-TP 组，二者只取其一），避免 hybrid EP+TP 下归约错 peer。
3. 启动流程接线与参数简化：`python/sglang/srt/distributed/bootstrap.py` 的 `_set_all_reduce_flags()` 调用 `set_flashinfer_allreduce_only(server_args.flashinfer_allreduce_fusion_backend is not None)`，`_init_parallel_groups()` 在 `initialize_model_parallel()` 之后调用 `_tag_groups_for_flashinfer_allreduce_only()`。按 review 意见删除了独立的 `--enable-flashinfer-pure-allreduce` flag 和架构白名单，`kAllReduce` 直接复用融合后端 flag，随融合路径默认开启。
 4. DSV3/V4 自动启用：`python/sglang/srt/arg_groups/overrides.py` 将 `DeepseekV4ForCausalLM` 加入 `_FLASHINFER_ALLREDUCE_FUSION_ARCHS`，使 DSV4 用户无需手动传任何参数即可获得该加速路径。
 5. hybrid EP+TP 配套修复与测试：`python/sglang/srt/layers/communicator.py` 的 `should_fuse_mlp_allreduce_with_next_layer()` 新增守卫：当 `moe_ep_size > 1` 且 `moe_tp_size > 1` 时禁止融合 post-experts all-reduce（融合会让下一层 residual+LN 只在一个组上归约，导致激活只归约一半 peer 的静默错误）。测试侧在 `test/registered/unit/layers/test_flashinfer_comm_fusion.py` 扩展了 `TestFlashInferAllReduceOnly`（12 个用例，覆盖正确性、形状守卫、容量 /dtype 守卫、group 匹配守卫）与 `TestTagGroupsForFlashInferAllReduceOnly`，并新增 `test/registered/unit/layers/test_layer_communicator_fusion_gate.py` 覆盖融合门控的 hybrid/pure 场景。

关键文件：

- `python/sglang/srt/distributed/parallel_state.py`（模块 并行组；类别 source；类型 core-logic；符号 `flashinfer_allreduce`, `_can_use_flashinfer_allreduce`, `_flashinfer_allreduce`, `set_flashinfer_allreduce_only`）：核心 dispatch 所在：注册 `flashinfer_allreduce` custom op，在 `GroupCoordinator.all_reduce()` 热路径插入 `FlashInfer` 分支，并通过 `_tag_groups_for_flashinfer_allreduce_only()` 精确打标避免 hybrid EP+TP 下归约错 peer。
- `python/sglang/srt/layers/flashinfer_comm_fusion.py`（模块 通信融合；类别 source；类型 core-logic；符号 `can_use_flashinfer_allreduce`, `flashinfer_allreduce`）：新增 `can_use_flashinfer_allreduce` 与 `flashinfer_allreduce` 两个核心函数：前者集中所有可用性判定（rank-invariant），后者执行 `AllReduceFusionPattern.kAllReduce` kernel 调用且不吞异常，是整条路径的正确性基石。
- `python/sglang/srt/layers/communicator.py`（模块 通信层；类别 source；类型 bugfix；符号 `should_fuse_mlp_allreduce_with_next_layer`）：修复 hybrid EP+TP 下融合 post-experts all-reduce 会静默产生未完全归约激活的严重正确性问题（Qwen3-30B-A3B tp=4 ep=2 实测为垃圾输出）。
- `python/sglang/srt/distributed/bootstrap.py`（模块 启动流程；类别 source；类型 dependency-wiring；符号 `_set_all_reduce_flags`, `_init_parallel_groups`,

set_flashinfer_allreduce_only, _tag_groups_for_flashinfer_allreduce_only) : 启动接线 : 在 _set_all_reduce_flags 中根据融合后端 flag 设置 set_flashinfer_allreduce_only, 在组初始化后调用 _tag_groups_for_flashinfer_allreduce_only, 是功能开关生效的入口。

- python/sglang/srt/arg_groups/overrides.py (模块 参数覆盖; 类别 source; 类型 configuration; 符号 _FLASHINFER_ALLREDUCE_FUSION_ARCHS) : 将 DeepseekV4ForCausalLM 加入 FlashInfer allreduce fusion 自动启用架构集, 使 DSV4 用户开箱即用; 同时清理了按 review 删除的独立纯 allreduce 架构白名单。
- test/registered/unit/layers/test_flashinfer_comm_fusion.py (模块 单元测试; 类别 test ; 类型 test-coverage; 符号 TestFlashInferAllReduceOnly, TestTagGroupsForFlashInferAllReduceOnly, _make_manager, _patched_attn_workspace) : 新增 TestFlashInferAllReduceOnly (12 个用例) 与 TestTagGroupsForFlashInferAllReduceOnly, 覆盖输出正确性、2D/ 连续形状守卫、容量 /dtype 守卫、group 精确匹配守卫及打标拓扑选择, 是判定逻辑的主要回归保障。
- test/registered/unit/layers/test_layer_communicator_fusion_gate.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 TestFuseMlpAllReduceGate, test_hybrid_ep_tp_does_not_fuse, test_pure_tp_still_fuses, test_pure_ep_still_fuses) : 新增测试验证 hybrid EP+TP 下禁止融合、纯 TP/ 纯 EP 下仍允许融合, 直接守护 communicator.py 的正确性修复。

关键符号: flashinfer_allreduce, can_use_flashinfer_allreduce, _can_use_flashinfer_allreduce, _flashinfer_allreduce, _tag_groups_for_flashinfer_allreduce_only, set_flashinfer_allreduce_only, should_fuse_mlp_allreduce_with_next_layer

关键源码片段

python/sglang/srt/distributed/parallel_state.py

核心 dispatch 所在: 注册 `flashinfer_allreduce` custom op, 在 `GroupCoordinator.all_reduce()` 热路径插入 FlashInfer 分支, 并通过 `_tag_groups_for_flashinfer_allreduce_only()` 精确打标避免 hybrid EP+TP 下归约错 peer。

```
# python/sglang/srt/distributed/parallel_state.py
# 注册为 custom op, 使 Dynamo 将其视为不透明调用, 可进入 piecewise CUDA graph;
# 判定逻辑在 trace 期由 _can_use_flashinfer_allreduce 完成, op 自身无回退。
@register_custom_op(out_shape="tensor")
def flashinfer_allreduce(tensor: torch.Tensor, group_name: str) -> torch.Tensor:
    assert group_name in _groups, f"Group {group_name} is not found."
    group = _groups[group_name]()
    if group is None:
        raise ValueError(f"Group {group_name} is destroyed.")
    return group._flashinfer_allreduce(tensor)
```

```
class GroupCoordinator:
    def _can_use_flashinfer_allreduce(self, input_: torch.Tensor) -> bool:
        # 只有被 _tag_groups_for_flashinfer_allreduce_only 标记过的组
```

```

# （即真正拥有某个 fusion workspace 的组）才允许走该路径。
if self._fi_workspace_hint is None:
    return False
from sglang.srt.layers.flashinfer_comm_fusion import (
    can_use_flashinfer_allreduce,
)

return can_use_flashinfer_allreduce(
    input_,
    use_attn_tp_group=(self._fi_workspace_hint == "attn_tp"),
    expected_world_size=self.world_size,
    expected_group_key=(self.device_group, self.cpu_group),
)

def _flashinfer_allreduce(self, input_: torch.Tensor) -> torch.Tensor:
    from sglang.srt.layers.flashinfer_comm_fusion import (
        flashinfer_allreduce as _impl,
    )

    return _impl(input_, use_attn_tp_group=(self._fi_workspace_hint == "attn_tp"))

def _tag_groups_for_flashinfer_allreduce_only():
    """给真正拥有 fusion workspace 的组打上 _fi_workspace_hint 标记。

    只有两个 workspace: attention-TP 一个、MoE 一个；MoE workspace 在
    moe_ep_size > 1 时于 EP 组上 rendezvous，否则于 MoE-TP 组上。
    因此 _MOE_EP 与 _MOE_TP 中只有一个是合法的——两个都打标会让
    MoE-TP 的 allreduce 在 hybrid EP+TP（如 tp=4, ep=2）下错误归约 EP peer。
    """
    if not _ENABLE_FLASHINFER_ALLREDUCE_ONLY:
        return

    moe_group = _MOE_EP if (_MOE_EP is not None and _MOE_EP.world_size > 1) else _MOE_TP
    for group, hint in ((moe_group, "moe"), (_ATTN_TP, "attn_tp")):
        if group is not None:
            group._fi_workspace_hint = hint

```

python/sglang/srt/layers/flashinfer_comm_fusion.py

新增 `can_use_flashinfer_allreduce` 与 `flashinfer_allreduce` 两个核心函数：前者集中所有可用性判定（rank-invariant），后者执行 `AllReduceFusionPattern.kAllReduce` kernel 调用且不吞异常，是整条路径的正确性基石。

```

# python/sglang/srt/layers/flashinfer_comm_fusion.py
# 纯 allreduce 的判定与执行分离：判定在 trace 期普通 Python 中完成，
# 执行被包裹为不透明 custom op；这样 Dynamo 不会拆开 kernel 调用，
# 而数据相关的守卫都留在 trace 期，避免重编译。

```

```

def can_use_flashinfer_allreduce(
    input_: torch.Tensor,
    *,
    use_attn_tp_group: bool,
    expected_world_size: int,
    expected_group_key: Tuple[Optional[ProcessGroup], Optional[ProcessGroup]],
) -> bool:
    """判定本次 all-reduce 能否走 FlashInfer kAllReduce 路径。

    所有检查必须保持 rank 无关 (rank-invariant) : 如果某个 rank 悄悄
    回退到 NCCL 而其余 rank 进入 kernel, 集体通信会失配并挂起。
    """

    if _flashinfer_allreduce_unavailable or _flashinfer_comm is None:
        return False

    # kAllReduce 只支持 2D 连续张量 (如 3D 的 vocab embedding 输出则回退) 。
    if input_.ndim != 2 or not input_.is_contiguous():
        return False

    workspace_manager = _get_workspace_manager(use_attn_tp_group)
    if not workspace_manager.initialized or workspace_manager.workspace is None:
        return False

    # workspace 必须是在完全相同的 peer 集合上 rendezvous 的:
    # hybrid EP+TP 下 MoE workspace 可能对应 EP 或 MoE-TP 组, 两者 world size
    # 相同但 rank 集合不同, 混用会静默归约错 peer (错误输出而非崩溃) 。
    if (
        workspace_manager.world_size != expected_world_size
        or workspace_manager.group != expected_group_key
    ):
        return False

    # 尺寸检查放在最后: token 维在 Dynamo 下是符号值, 静态不可用的配置
    # 必须先短路, 避免 trace 期访问 workspace 对象 (与融合路径同规则) 。
    token_num, hidden_dim = input_.shape
    if torch.compiler.is_compiling():
        return (
            workspace_manager.max_token_num is not None
            and workspace_manager.hidden_dim is not None
            and workspace_manager.dtype is not None
            and token_num <= workspace_manager.max_token_num
            and hidden_dim <= workspace_manager.hidden_dim
            and workspace_manager.dtype == input_.dtype
        )

    return workspace_manager.is_buffer_size_sufficient(
        token_num=token_num,
        hidden_dim=hidden_dim,
        dtype=input_.dtype,
    )

```

```

)

def flashinfer_allreduce(
    input_: torch.Tensor,
    *,
    use_attn_tp_group: bool,
) -> torch.Tensor:
    """执行 FlashInfer kAllReduce (仅 allreduce, 无融合) 。

    假定调用前 can_use_flashinfer_allreduce 已返回 True; 这里刻意不捕获
    kernel 异常——吞掉异常会让本 rank 走 NCCL 而其他 rank 留在 kernel 内,
    失配挂起而不是报错。
    """
    workspace_manager = _get_workspace_manager(use_attn_tp_group)

    output = torch.empty_like(input_)
    kwargs = dict(
        input=input_,
        workspace=workspace_manager.workspace,
        pattern=_flashinfer_comm.AllReduceFusionPattern.kAllReduce,
        launch_with_pdl=True,
        fp32_acc=False, # 与融合路径保持一致的数值语义
        output=output,
    )
    if _flashinfer_allreduce_supports_trigger_completion:
        kwargs["trigger_completion_at_end"] = False # 与融合路径保持一致
    _flashinfer_comm.allreduce_fusion(**kwargs)
    return output

```

python/sglang/srt/layers/communicator.py

修复 hybrid EP+TP 下融合 post-experts all-reduce 会静默产生未完全归约激活的严重正确性问题 (Qwen3-30B-A3B tp=4 ep=2 实测为垃圾输出) 。

```

# python/sglang/srt/layers/communicator.py
# 决定 MLP 的 all-reduce 是否可以与下一层的 residual+LN 融合。
def should_fuse_mlp_allreduce_with_next_layer(
    self, forward_batch: ForwardBatch
) -> bool:
    # MOE_FULLL 时无法融合: 融合路径会跳过 postprocess_layer (含 moe_cp scatter) ,
    # 导致 hidden_states 与 residual 形状不匹配。
    if is_enable_moe_cp_allgather():
        return False

    # 关键修复: hybrid EP+TP (moe_ep_size > 1 且 moe_tp_size > 1) 下,
    # post-experts 的归约要跨两个不相交的组 (_MOE_EP 再 _MOE_TP) 。
    # 融合让下一层 residual+LN 只在一个组上归约, 而
    # should_skip_post_experts_all_reduce() 会把两个都跳过——
    # 结果激活只归约了一半 peer, 静默错误而非崩溃。

```

```

parallel = get_parallel()
if parallel.moe_ep_size > 1 and parallel.moe_tp_size > 1:
    return False

# 其余原有守卫（EAGLE、input_scattered、SCATTERED 模式等）保持不变。
if (
    is_dp_attention_enabled()
    and self._speculative_algo is not None
    and self._speculative_algo.is_eagle()
):
    return False

if get_attn_tp_context().input_scattered:
    return False

batch_size = (
    forward_batch.input_ids.shape[0]
    if hasattr(forward_batch, "input_ids")
    else 0
)
if self.layer_scatter_modes.mlp_mode == ScatterMode.SCATTERED:
    return False

return apply_flashinfer_allreduce_fusion(batch_size) or ...

```

评论区精华

Review 中围绕「如何安全地把 NCCL 热路径换成 FlashInfer」展开了高质量讨论：

- b8zhong 认为独立 flag 与架构白名单冗余，建议直接复用 `--flashinfer-allreduce-fusion-backend` 控制，作者采纳并删除 `_FLASHINFER_ALLREDUCE_ONLY_ARCHS`。
- mmangkad 连续指出三个正确性问题：`flashinfer_allreduce()` 的 try-except 回退不是跨 rank 同步的，divergent rank 会 NCCL/FlashInfer 失配挂起；`fp32_acc=True` 与融合路径默认 `False` 不一致；`trigger_completion_at_end` 未显式设置而继承默认 `True`，与融合路径 `False` 不一致。作者逐一修正：移除 try-except、改为 `fp32_acc=False`、显式传 `trigger_completion_at_end=False`（按 capability 门控）。
- mmangkad 指出 `_MOE_TP` 与 `_MOE_EP` 同时打标会在 hybrid EP+TP 下归约错 peer，并建议用 custom op 包装以支持 piecewise CUDA graph；作者据此实现 `can_use_/execute` 分离与精确打标。
- b8zhong 与 shyeh25 将本 PR 与 PR#32339（custom-allreduce 多节点升级）对比实测，二者性能相当（1P1D 约 51.4 TPS，1P4D 约 929 TPS）；结论是本 PR 对 TP16/TP32 仍有独立价值。
- ormandj 在 SM120 上实测发现本 PR fails closed（arch 校验拒绝 + 非 multicast TRT-LLM workspace 被 preflight 拒绝），需叠加两个 commit 才能启用（+5.2% median），并开 PR#32330 跟进。

- 是否保留独立的纯 allreduce 开关与架构白名单 (design): 作者删除独立 flag 与白名单, kAllReduce 随融合后端 flag 默认开启, 仅保留 DSV3/V4 自动启用逻辑 (DSV4 加入 `_FLASHINFER_ALLREDUCE_FUSION_ARCHS`) 。
- fallback 路径的跨 rank 一致性 (correctness): 作者移除 try-except, 改为「判定 (rank-invariant) 与执行 (不吞异常) 分离」, 并强调所有判定必须保持 rank 无关。
- fp32_acc 与 trigger_completion_at_end 与融合路径不一致 (correctness): 作者修正为 fp32_acc=False, 并显式传 trigger_completion_at_end=False (按 `_flashinfer_allreduce_supports_trigger_completion` 门控), 与融合路径对齐。
- workspace 与进程组的精确绑定 (hybrid EP+TP) (correctness): 作者改为只打标真正持有 workspace 的组 (EP 或 MoE-TP 二选一), `_fi_workspace_hint` 在 `__init__` 初始化; 同时将 kernel 调用注册为 custom op, 并在 `is_compiling()` 路径也接入 dispatch。
- 与 PR#32339 (custom-allreduce 多节点升级) 的性能对比 (performance): 两者性能 on par; 本 PR 作为不依赖 NCCL ring 的替代路径保留, 对超大 TP 场景有价值。
- SM120 上 fails closed (performance): 合入前未在 SM120 上启用, 由 #32330 以叠加 commit 方式跟进; 作者未在本 PR 内处理。
- 测试类与注释的规范性 (style): 作者已将新增测试类改为 CustomTestCase, 并精简注释。

风险与影响

- 风险:
 1. 核心路径变更回归风险: `GroupCoordinator.all_reduce()` 是所有 TP 通信的热路径, dispatch 插入位置在既有 `pynccl/custom-allreduce` 之前, 任何守卫的 rank 不一致都会导致集体通信失配挂起而非报错。代码通过「所有判定 rank-invariant + kernel 异常不捕获」缓解, 但覆盖多拓扑 (hybrid EP+TP、DP attention、PP) 的端到端验证仍有限。
 2. CUDA graph 兼容性: 作者在评论中承认 PCG capture 未完全端到端验证 ("PCG capture not yet validated end to end"), 虽然 custom op 设计上可进入 piecewise graph, 但真实图捕获场景仍需回归。
 3. SM120 兼容性缺口: ormandj 实测在 SM120 (RTX PRO 6000 Blackwell) 上该 PR fails closed, 需额外 commit; 这意味着一部分 Blackwell 用户 (非 GB200/GB300 MNNVL 环境) 在合入后无法立即受益。
 4. 数值语义变化: fp32_acc=False 与融合路径对齐后, 纯 allreduce 的数值行为与 NCCL sum 存在潜在差异; GSM8K 精度持平 (0.958), 但其他模型 / 量化组合未覆盖。
 5. hybrid EP+TP 禁用融合的影响: `communicator.py` 的守卫会让 hybrid EP+TP 场景失去 `fuse_mlp_allreduce` 优化, 这是为避免静默错误的有意取舍, 但可能影响该拓扑下的性能预期。
 - 影响: 用户影响: DeepSeek-V3/V4 用户在没有额外参数的情况下自动获得 FlashInfer kAllReduce 加速路径; Blackwell (SM100 系) 小并发场景收益显著 (shyeh25 实测 GB200 TP8 单并发 TPS per GPU 从 78.77 提升到 97.12, 约 +23%; BS 1-16 的 decode tok/s 提升 +2.6%~+6.9%), 大并发基本持平 (BS=64 约 -1.4%); SM90 单机仍可用 TRT-LLM 后端, 但不能用于多节点。系统影响: 零额外显存开销, 复用现有融合 workspace; 与 PR#32339 的 custom-allreduce v2 多节点路径性能相当, 但本 PR 不依赖 NCCL ring, 在 TP16/TP32 等大 TP 场景有独立价值。团队影响: 确立了「可用性判定 (trace 期) 与 kernel 执行 (图内) 分离 + 精确 group 绑定」的模式

，为后续其他 FlashInfer 融合原语接入提供了可复用的设计范式。 - 风险标记：核心路径变更，跨 rank 一致性风险，CUDA graph 兼容性待验证，SM120 兼容性缺口，hybrid EP+TP 边界

关联脉络

- PR #32339 Upgrade custom-allreduce kernels for multi-node GB200/300（评论中提及）：b8zhong 提出其可能替代本 PR；shyeh25 实测两者性能相当，本 PR 对 TP16/TP32 仍有独立价值。
- PR #32963 Fix hybrid EP+TP post-experts all-reduce fusion（评论中提及，作者声明本 PR 基于它）：作者在 review 回复中说明该 PR 修复的 hybrid EP+TP 归约问题被 #32963 正式修复，本 PR rebase 在其上。
- PR #32330 Enable flashinfer allreduce-only on SM120（ormandj 所开，叠加本 PR 的附加 commit）：ormandj 实测 SM120 fails closed 后提供两个修复 commit（arch 选择器与 multicast preflight 移除），并建议折回本 PR。