

PR #46807 完整报告

vllm-project/vllm

PD disagg with Mooncake Connector: GDN support (Qwen3.5) and MLA support (Deepseek-V4-Flash)

合并时间: 2026-06-30 14:29

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/46807>

执行摘要

- 一句话: Mooncake 连接器支持 GDN 和 MLA 混合 KV cache P/D 分离
- 推荐动作: 值得精读, 尤其是 `_expand_transfer_regions` 的重构和 MambaSpec N-1 边界设计。展示了如何将同质传输路径演化为缓存感知传输路径, 对于理解 vLLM 的 KV cache 抽象很有帮助。

功能与动机

Mooncake P/D 分离在混合 KV cache 模型 (如 Qwen3.5) 中失效, 原因包括: Full Attention 和 GDN 具有不同块语义; HMA 返回每个 KV cache 组的块 ID, 但 Mooncake 之前将其扁平化复用; 注意力逻辑块需要扩展为物理核块, 而 GDN 状态块应保持状态槽空间; 块优先的 FA 区域需要 K/V 分割, 而 GDN 状态区域需保持整体。来自 PR #46807 的描述。

实现拆解

1. 传输区域扩展重构: 修改 `_expand_transfer_regions` 函数, 新增 `kv_block_lens`、`group_indices` 和 `split_kv_regions` 参数, 使每个区域可指定自己的 `kv_block_len` 和所属组索引。TransferRegion 数据类新增 `group_index` 字段。
2. 内存注册区分区域类型: `register_kv_caches` 保留所有逻辑传输区域, 即使共享同一个 backing 分配, 也同时发出 FA 和 GDN 区域。对 FA 区域应用 K/V 分割 (`kv_block_len = block_len // 2`), GDN/Mamba 状态区域保持 `kv_block_len = block_len`。
3. 传输规划保留组身份: 请求块 ID 按组消费: FA 组使用 FA 组块 ID, GDN 组使用自己的组块 ID; 来自对齐 MambaSpec 状态的空占位块被跳过, 不发起传输。
4. 逻辑到物理块映射: 仅对 FA 区域应用 `_logical_to_kernel_block_ids` 将逻辑块 ID 扩展为物理块 ID; GDN 状态区域不做扩展, 因其块 ID 代表状态槽而非 per-token KV 页面。
5. MambaSpec N-1 边界: 在调度器端实现。Prefiller 传输状态时只计算到 token N-2 (状态 `h(N-1)`), decoder 本地重算最后一个 token N-1。通过 `_get_remote_prefill_token_count` 和 `_truncate_mamba_request_for_prefill` 在 Mooncake 和 NIXL 两侧统一实现。

关键文件:

- `vllm/distributed/kv_transfer/kv_connector/v1/mooncake/mooncake_connector.py` (模块连接器; 类别 source; 类型 dependency-wiring; 符号

`_get_remote_prefill_token_count`, `_truncate_mamba_request_for_prefill`, `logical_to_kernel_block_ids`) : 核心文件, 实现 Mooncake 连接器混合 KV cache 传输的全部逻辑。包括传输区域扩展、组感知块规划、逻辑到物理映射、MambaSpec N-1 边界处理。

- `tests/v1/kv_connector/unit/test_mooncake_connector_hybrid_mamba.py` (模块 混合测试; 类别 test; 类型 test-coverage; 符号 `noop_shutdown`, `make_hybrid_gdn_kv_cache_config`, `make_hybrid_gdn_scheduler`, `test_hybrid_gdn_remote_prefill_uses_mamba_n_minus_one`) : 新增专用测试文件, 覆盖 Mooncake 混合 FA+GDN 场景, 验证 N-1 边界、区域注册、传输参数保留组身份等关键行为。
- `vllm/distributed/kv_transfer/kv_connector/v1/nixl/base_scheduler.py` (模块 调度器; 类别 source; 类型 core-logic; 符号 `_mamba_prefill_token_count`, `_get_remote_prefill_token_count`) : NIXL 基础调度器, 重命名 `_mamba_prefill_token_count` 为 `_get_remote_prefill_token_count` 以对齐 Mooncake 修改, 保持一致性。
- `vllm/distributed/kv_transfer/kv_connector/v1/nixl/pull_scheduler.py` (模块 调度器; 类别 source; 类型 core-logic) : 更新对重命名方法的调用, 从 `_mamba_prefill_token_count` 改为 `_get_remote_prefill_token_count`。
- `vllm/distributed/kv_transfer/kv_connector/v1/nixl/push_scheduler.py` (模块 调度器; 类别 source; 类型 core-logic) : 更新对重命名方法的调用, 从 `_mamba_prefill_token_count` 改为 `_get_remote_prefill_token_count`。
- `tests/v1/kv_connector/unit/test_mooncake_connector.py` (模块 连接器测试; 类别 test; 类型 test-coverage) : 适配性修改, 为 worker mock 添加 `kv_block_len_per_layer` 等字段, 支持新传输参数。

关键符号: `_expand_transfer_regions`, `_logical_to_kernel_block_ids`, `_get_remote_prefill_token_count`, `_truncate_mamba_request_for_prefill`, `get_num_new_matched_tokens`, `register_kv_caches`, `build_xfer_params`

关键源码片段

`vllm/distributed/kv_transfer/kv_connector/v1/mooncake/mooncake_connector.py`

核心文件, 实现 Mooncake 连接器混合 KV cache 传输的全部逻辑。包括传输区域扩展、组感知块规划、逻辑到物理映射、MambaSpec N-1 边界处理。

```
@dataclass(frozen=True)
class TransferRegion:
    """单个可传输区域的元数据, 每个区域属于特定 KV cache 组。"""
    layer_name: str
    layer_index: int
    base_addr: int
    block_len: int # 底层张量中一个块的字节长度
    kv_block_len: int # 实际传输中每个 KV 块的字节长度。
        # 对于 FA blocks-first 布局, kv_block_len = block_len // 2 (K/V 分割);
        # 对于 GDN/Mamba 状态, kv_block_len = block_len (整体传输)。
```

```
group_index: int = 0 # 新增: 所属 KV cache 组索引
```

```
def _expand_transfer_regions(
    base_addrs: list[int],
    block_lens: list[int],
    kv_block_lens: list[int], # 新增: 每个区域自己的 kv_block_len
    layer_names: list[str],
    layer_indices: list[int],
    is_kv_layout_blocks_first: bool,
    group_indices: list[int] | None = None, # 新增: 组索引列表
    split_kv_regions: list[bool] | None = None, # 新增: 是否需要 K/V 分割
) -> list[TransferRegion]:
    # ... 参数校验 ...
    for (base_addr, block_len, kv_block_len, layer_name,
        layer_index, group_index, split_kv_region) in zip(
        base_addrs, block_lens, kv_block_lens,
        layer_names, layer_indices,
        group_indices, split_kv_regions):
        if is_kv_layout_blocks_first and split_kv_region:
            kv_block_len = block_len // 2
        regions.append(TransferRegion(
            layer_name=layer_name,
            layer_index=layer_index,
            group_index=group_index,
            base_addr=base_addr,
            block_len=block_len,
            kv_block_len=kv_block_len,
        ))
    return regions
```

tests/v1/kv_connector/unit/test_mooncake_connector_hybrid_mamba.py

新增专用测试文件，覆盖 Mooncake 混合 FA+GDN 场景，验证 N-1 边界、区域注册、传输参数保留组身份等关键行为。

```
def test_register_kv_caches_emits_fa_and_gdn_regions(monkeypatch):
    monkeypatch.setenv("VLLM_MOONCAKE_ABORT_REQUEST_TIMEOUT", "5")
    vllm_config = create_vllm_config(
        kv_connector="MooncakeConnector",
        kv_role="kv_consumer",
    )
    kv_cache_config = make_hybrid_gdn_kv_cache_config(
        vllm_config.cache_config.block_size
    )
    with set_current_vllm_config(vllm_config), patch_worker_dependencies():
        connector = MooncakeConnector(
            vllm_config,
            KVConnectorRole.WORKER,
            kv_cache_config,
```

```

)
worker = connector.connector_worker

# 注册 FA 和 GDN 张量 (共享同一 backing 分配以验证去重)
fa_cache = torch.empty((2, 2, 11), dtype=torch.float16)
gdn_conv_state = torch.empty((2, 22), dtype=torch.float16)
gdn_ssm_state = torch.empty((2, 4), dtype=torch.float16)

worker.register_kv_caches({
    "model.layers.0.self_attn": fa_cache,
    "model.layers.1.linear_attn.conv_state": gdn_conv_state,
    "model.layers.1.linear_attn.ssm_state": gdn_ssm_state,
})

regions = worker.transfer_regions
# 应包含三个区域: 1 个 FA, 2 个 GDN (conv_state 和 ssm_state 分别注册)
assert len(regions) == 3
# FA 区域: group_index=0, kv_block_len=block_len//2
fa_region = [r for r in regions if r.layer_name == "model.layers.0.self_attn"][0]
assert fa_region.group_index == 0
assert fa_region.kv_block_len == fa_region.block_len // 2
# GDN 区域: group_index=1, kv_block_len=block_len (不分割)
gdn_regions = [r for r in regions if "linear_attn" in r.layer_name]
for r in gdn_regions:
    assert r.group_index == 1
    assert r.kv_block_len == r.block_len

```

评论区精华

Review 中 reviewer zhewenl 建议将 Mooncake 调度器中的 `_mamba_prefill_token_count` 重命名为 `_get_remote_prefill_token_count` 以避免混淆。作者 andakai 同意并同步修改了 NIXL 调度器的同一方法名，保持一致性。

- 方法重命名: `_mamba_prefill_token_count` 改为 `_get_remote_prefill_token_count` (design): 方法重命名为 `_get_remote_prefill_token_count`, 同时更新 NIXL 调度器以保持一致。

风险与影响

- 风险:
 1. DeepSeek V4 Flash 支持验证不足: 作者在 issue 评论中指出 DeepSeek V4 Flash 在 main 分支上已有精度问题, 本 PR 未解决, MLA 传输可能不完整, 需要后续修复。
 2. 架构调整范围大: 核心传输逻辑从同质 KV cache 假设改为感知缓存组类型, 可能引入纯 attention 模型的回归。
 3. 测试覆盖: 新增混合场景单元测试, 但缺少端到端多 GPU 集成测试。
 4. 性能: 组身份保留和逻辑到物理映射增加额外计算开销, 但预期影响有限。- 影响: 影响范围集中在使用 Mooncake 连接器的 P/D 分离部署, 特别是 Qwen3.5 用户将获得正

确性提升（表格显示精度接近独立模型，GSM8K 1p1d 准确度 0.323730 vs 独立 0.330553）。NIXL 调度器用户无功能变化，仅方法重命名。团队需关注 DeepSeek V4Flash 的后续修复。 - 风险标记：DeepSeek V4 Flash 精度问题未解决，核心传输路径重构，新增混合测试，可能影响纯 attention 模型

关联脉络

- PR #39482 Initial Mooncake hybrid attention support for Qwen3.5: 本 PR 建立在此 PR 的基础上，继续完善混合注意力支持。
- PR #41869 NIXL P/D disaggregation support for Qwen3.5 GDN: 本 PR 将 NIXL 已验证的 GDN P/D 分离行为迁移到 Mooncake 连接器路径。
- PR #44456 Broader KV-cache layout refactor around Mamba cache standardization: 相关但不同，关注 Mamba 缓存标准化与本 PR 的请求级 Mooncake P/D 传输修复。
- PR #46334 Mooncake logical/physical block-ratio handling: 涉及物理块映射，本 PR 利用此机制但处理混合 FA+GDN 组语义。
- PR #46004 Mooncake shared KV metadata for DeepSeek V4 PP/PD: 相关 Mooncake 区域元数据工作，但针对不同模型 / 拓扑。