

# PR #44455 完整报告

vllm-project/vllm

[2/N][KV-Cache Layout Refactor] Pack K/V into the content dim across attention backends

合并时间: 2026-07-11 23:11

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/44455>

## 执行摘要

- 一句话: 将 K/V 打包到内容维度, 统一 KV 缓存布局为 4D
- 推荐动作: 建议所有 v1 用户关注此 PR, 特别是在使用 PD disagg 的场景下。设计决策值得学习: 通过统一布局后端接口, 减少连接器耦合。合并前需确保 AMD 测试修复通过。

## 功能与动机

参照 RFC #42082, 标准化 KV 缓存布局以消除后端间的差异, 减少 KV-connector 中 is\_mamba/is\_mla 等标志的泛滥。当前布局导致代码紧耦合, 且不利于异构 TP 和跨层传输。本 PR 集中解决 K/V 打包问题, 是整体布局标准化系列的一部分。

## 实现拆解

按以下步骤实现:

1. 修改注意力后端形状定义: 在所有注意力后端 (FlashAttention、FlashInfer、Triton、ROCm 等) 的 get\_kv\_cache\_shape 方法中, 将返回的形状从 (num\_blocks, 2, block\_size, num\_kv\_heads, head\_size) 改为 (num\_blocks, num\_kv\_heads, block\_size, 2\*head\_size)。量化缓存类似, 但使用 padded\_hs。
2. 更新 stride order: 对应的 get\_kv\_cache\_stride\_order 方法适配 4D 形状, 在 NHD 和 HND 布局下分别返回合适的 permutation (例如 NHD 下为 (0,2,1,3))。
3. 改造 K/V 提取逻辑: 在各个后端的 forward/do\_kv\_cache\_update 中, 用 kv\_cache.transpose(1,2).split(head\_size, dim=-1) 替换旧的 kv\_cache.unbind(1) 和 PagedAttention.split\_kv\_cache 调用, 获得零成本 view。
4. 简化 NVFP4 量化缓存处理: 移除 nvfp4\_kv\_cache\_split\_views 函数, 改为调用 nvfp4\_split\_data\_scale 对一侧处理。量化缓存使用 separate head groups 布局时, 直接在 dim=1 上 split。
5. 调整 KV-connector 传输拓扑: 在 TransferTopology 中判断形状是否为 4D blocks-first 并断言, 删除 is\_kv\_layout\_blocks\_first、split\_k\_and\_v 等属性, 简化虚拟 split 逻辑。NIXL 和 MoRPIO 连接器均适配新的 4D 布局, K/V 不再拆分为独立区域。
6. 更新测试和平台代码: 测试文件 (test\_nixl\_connector、test\_mooncake\_connector、test\_minimax\_m3 等) 更新以匹配新形状; ROCm 平台代码调整以适应包装布局。

关键文件:

- vllm/utils/torch\_utils.py (模块 工具函数; 类别 source; 类型 core-logic; 符号 \_nvfp4\_split\_data\_scale, nvfp4\_split\_data\_scale, nvfp4\_kv\_cache\_split\_views) : 修改 NVFP4 KV 缓存的拆分函数, 移除旧的分拆视图函数, 精简为单侧处理, 反映布局变化对量化缓存的影响。
- vllm/distributed/kv\_transfer/kv\_connector/utils.py (模块 KV 连接器; 类别 source; 类型 core-logic; 符号 is\_kv\_layout\_blocks\_first, split\_k\_and\_v) : 核心修改: 在 TransferTopology 中识别 4D packed 形状, 简化 blocks-first 判断, 移除 split\_k\_and\_v 属性, 直接使用统一布局。
- vllm/v1/attention/backends/triton\_attn.py (模块 注意力后端; 类别 source; 类型 core-logic; 符号 to\_f32\_units) : 展示了核心的 get\_kv\_cache\_shape 和 get\_kv\_cache\_stride\_order 变化, 反映 4D 打包布局的实现。
- vllm/v1/attention/backends/flashinfer.py (模块 注意力后端; 类别 source; 类型 core-logic) : 修改 TRTLLM 去量化 kernel 的 stride 计算以支持打包布局, 调整量化缓存处理路径。
- vllm/distributed/kv\_transfer/kv\_connector/v1/nixl/base\_worker.py (模块 NIXL 连接器; 类别 source; 类型 core-logic) : NIXL 连接器移除旧 K/V 拆分传输逻辑, 适配单一区域传输; 调整 host buffer 分配以匹配 4D 布局。
- vllm/distributed/kv\_transfer/kv\_connector/v1/moriio/moriio\_layout.py (模块 MoRIIO 连接器; 类别 source; 类型 core-logic; 符号 \_content\_packed\_dim) : 添加对 4D packed 形状的传输几何计算支持。
- tests/v1/kv\_connector/unit/test\_nixl\_connector.py (模块 NIXL 测试; 类别 test; 类型 test-coverage; 符号 test\_hybrid\_mamba\_attention\_remote\_descs\_use\_packed\_head\_slices) : 更新测试以匹配新形状和传输区域计数; 添加新的测试用例验证打包头部切片。

关键符号: nvfp4\_split\_data\_scale, is\_kv\_layout\_blocks\_first, split\_k\_and\_v, \_content\_packed\_dim, TritonAttentionBackend.get\_kv\_cache\_shape, TritonAttentionBackend.get\_kv\_cache\_stride\_order, FlashInferBackend.get\_kv\_cache\_shape, TransferTopology.post\_init

## 关键源码片段

### vllm/utils/torch\_utils.py

修改 NVFP4 KV 缓存的拆分函数, 移除旧的分拆视图函数, 精简为单侧处理, 反映布局变化对量化缓存的影响。

```
def nvfp4_split_data_scale(kv_side: torch.Tensor) -> tuple[torch.Tensor, torch.Tensor]:
    """
    将NVFP4缓存的一侧 (K或V) 拆分为数据和缩放因子。
    输入形状为 (B, H, N, full_dim), 其中 full_dim = data_dim + scale_dim。
    布局为 [K_data | K_scale | V_data | V_scale] 每页连续存储。
    调用方需先通过 split 或切片获得单侧。
    """
    num_pages, dim_1, dim_2, full_dim = kv_side.shape
    data_dim = full_dim * 8 // 9
    scale_dim = full_dim - data_dim
```

```

data_per_kv = dim_1 * dim_2 * data_dim
page_bytes = kv_side.stride(0)

# 从原始 stride 推导 data 和 scale 的 stride, 保持物理布局 (NHD 或 HND)
s1 = kv_side.stride(1) * data_dim // full_dim
s2 = kv_side.stride(2) * data_dim // full_dim
data_shape = (num_pages, dim_1, dim_2, data_dim)
data_strides = (page_bytes, s1, s2, 1)

s1_s = kv_side.stride(1) * scale_dim // full_dim
s2_s = kv_side.stride(2) * scale_dim // full_dim
scale_shape = (num_pages, dim_1, dim_2, scale_dim)
scale_strides = (page_bytes, s1_s, s2_s, 1)

base = kv_side.storage_offset()
data = torch.as_strided(kv_side, data_shape, data_strides, storage_offset=base)
scale = torch.as_strided(kv_side, scale_shape, scale_strides, storage_offset=base + data_per_kv).view(torch.float8_e4m3fn)
return data, scale

```

### vllm/distributed/kv\_transfer/kv\_connector/utils.py

核心修改: 在 TransferTopology 中识别 4D packed 形状, 简化 blocks-first 判断, 移除 split\_k\_and\_v 属性, 直接使用统一布局。

```

# 在 TransferTopology.__post_init__ 中, 识别标准 4D packed 布局
head_size = 1
# Mock 一个块形状来检查 layout
kv_cache_shape = attn_backend.get_kv_cache_shape(
    num_blocks=1, block_size=_MOCK_BLOCK_SIZE, num_kv_heads=1, head_size=head_size,
)
logger.debug("Test kv_cache_shape: %s", kv_cache_shape)

# 新的注意力缓存形状应为 4D: [num_blocks, num_kv_heads, block_size, content_size]
assert kv_cache_shape[0] == 1, (
    "KV cache layout must be blocks-first; expected mocked "
    f"num_blocks=1 in leading dim, got shape {kv_cache_shape}."
)
if not self.is_mla:
    assert len(kv_cache_shape) == 4, (
        "Attention KV cache layout must be standardized as "
        "[num_blocks, num_kv_heads, block_size, content_size], "
        f"got shape {kv_cache_shape}."
    )

# 旧属性 is_kv_layout_blocks_first 和 split_k_and_v 被移除,
# blocks-first 已成为唯一布局, K/V 不再需要拆分传输。

```

### vllm/v1/attention/backends/triton\_attn.py

展示了核心的 `get_kv_cache_shape` 和 `get_kv_cache_stride_order` 变化，反映 4D 打包布局的实现。

```
@staticmethod
def get_kv_cache_shape(num_blocks, block_size, num_kv_heads, head_size, cache_dtype_str=
"auto"):
    # K 和 V 打包到最后二维：逻辑形状为 (B, H, N, 2*C)
    if block_size % 16 != 0:
        raise ValueError("Block size must be a multiple of 16.")
    if kv_cache_uses_per_token_head_scales(cache_dtype_str):
        # 量化情况下额外 padding 用于存放缩放因子
        from vllm.utils.torch_utils import STR_DTYPE_TO_TORCH_DTYPE, get_dtype_size
        cache_dtype = STR_DTYPE_TO_TORCH_DTYPE[cache_dtype_str]
        scale_pad = get_dtype_size(torch.float32) // get_dtype_size(cache_dtype)
        if get_kv_quant_mode(cache_dtype_str) == KVQuantMode.INT4_PER_TOKEN_HEAD:
            data_head_size = head_size // 2
        else:
            data_head_size = head_size
        padded_hs = data_head_size + scale_pad
        return (num_blocks, num_kv_heads, block_size, 2 * padded_hs)
    return (num_blocks, num_kv_heads, block_size, 2 * head_size)

@staticmethod
def get_kv_cache_stride_order(cache_layout=None):
    """根据cache_layout返回从逻辑形状到物理内存的permutation。
    逻辑形状 (B, H, N, 2*C)，物理排列由stride_order控制。
    NHD: 物理为 (B, N, H, 2*C)
    HND: 物理为 (B, H, N, 2*C) (恒等)
    """
    if cache_layout == "NHD":
        return (0, 2, 1, 3)
    elif cache_layout == "HND":
        return (0, 1, 2, 3)
    else:
        raise ValueError(f"Unknown cache layout: {cache_layout}")
```

## 评论区精华

核心讨论包括：

- AMD 测试失败：NickLucche 发现三个 AMD 侧测试失败（MLA backend、MoRIIO connector 等），需要修复。
- prefix\_prefill stride 映射错误：depthfirst-app 指出在 `prefix_prefill.py` 中，新旧布局下 stride 索引交换，可能导致 KV 数据读写出界（高严重性）。LucasWilkinson 后续修复。
- 量化缓存兼容性：MatthewBonanni 担心 NVFP4 + 异构 TP NIXL 传输，LucasWilkinson 回应 NVFP4 + PD 在主线上当前未测试 / 不工作，因此暂时 break 是允许的。
- Intel 异构硬件：NickLucche 担忧去掉 `permute` 逻辑会影响 Intel 等异构部署，需要确保 backend 支持 HND。

- ROCm\_ATTN 支持: NickLucche 疑问 ROCm\_ATTN 是否支持 blocks-first, LucasWilkinson 确认在 MI300x 上测试无误。
- 文档删除: NickLucche 对 Claude 删除 NIXL 文档表示不满。
- AMD 测试失败 (testing): 需要在合并前修复, 后续提交似乎已解决。
- prefix\_prefill stride 映射错误 (correctness): LucasWilkinson 在后续提交中修复。
- NVFP4 量化缓存兼容性 (question): LucasWilkinson 回应 NVFP4 + PD 在主线上当前未测试 / 不工作, 暂时 break 可以接受, 后续再完善。
- Intel 异构硬件支持 (design): 暂无明确结论, 需 Intel 方确认。
- DOCS 文档被删除 (question): 属于自动化 review 工具的副作用。

## 风险与影响

- 风险:
  1. stride 计算错误 (高): 在 prefix\_prefill 和 scale cache 中, 深度优先分析发现 stride 映射错误, 可能导致 KV 数据损坏或越界读取。已在后续提交中修复, 但需警惕未覆盖的 kernel 路径。
  2. 量化缓存回归 (中): NVFP4 和 FP8 量化缓存布局特殊, 新方案使用 head-group 布局, 若未彻底测试可能引入精度或性能问题。当前 NVFP4+PD 未被测试, 可能隐藏回归。
  3. 异构硬件兼容性 (中): Intel GPU、AMD 等平台需要额外验证新布局, 特别是 NIXL 连接器移除 permute 逻辑后对非默认布局的影响。
  4. 性能退化 (低): 新增的 transpose 和 split 操作理论上为零成本 view, 但某些路径可能因 stride 不连续导致后续 kernel 变慢。复审评论中提到 perf 影响应极小。- 影响: 影响范围广: 所有使用 v1 attention backends 的模型都会受到影响, 因为 KV 缓存形状改变。但经过 E2E lm-eval PD disagg 测试验证, 覆盖了 Qwen、DeepSeek、Nemotron 等模型在多种 backend 和拓扑配置下, 结果通过。用户升级到该 PR 后需确保模型兼容。团队需注意与外部 PR (如 Intel GPU、AMD ROCm) 的冲突。测试配套更新完整, 包括 connector 单元测试。- 风险标记: 核心路径变更, 多个 backend 同时修改, 量化缓存路径未充分测试, AMD 测试失败待修复, prefix\_prefill stride 隐患, 异构硬件兼容性待确认

## 关联脉络

- PR #42082 [RFC]: Standardize KV-cache Layouts: 本 PR 是 RFC #42082 中标准化 KV 缓存布局系列的第 2 部分, RFC 提出了全局布局规范。
- PR #42374 first part of RFC #42082: 本 PR 基于 #42374 拆分为 4 个 PR, #44455 是其中的 [2/N], 负责打包 K/V 到内容维度。
- PR #44454 [1/N][KV-Cache Layout Refactor] Refactor DSV4 KV cache config: 本 PR 依赖 #44454 作为基础, 必须先合并。
- PR #44456 [3/N][KV-Cache Layout Refactor] Standardize Mamba cache; drop get\_transfer\_cache\_regions: 与本 PR 同系列, 后续步骤。

- PR #44458 [4/N][KV-Cache Layout Refactor] Standardize KV cache layout: 系列的最后部分。
- PR #41657 related bug PR: 本 PR 动机之一提到该 PR 导致的 bug。