

PR #51738 完整报告

vllm-project/vllm

[Perf] Avoid more GPU<->CPU syncs on the model execution path

合并时间: 2026-08-12 10:30

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51738>

执行摘要

- 一句话: 模型执行路径再消除多组 GPU-CPU 同步
- 推荐动作: 值得精读。这是 vLLM 官方维护者在 GPU 同步消除主线上的第二批落地, 展示了一类可复用的性能优化范式: 用主机侧已有精确数据推导设备端元数据、用 `async` 拷贝与 `fill_` 消除隐式同步、用 `keep_on_cpu` 避免无意义 staging。建议重点阅读 `vllm/v1/attention/backends/mamba_attn.py` 的 `_prefill_cpu_metadata` 抽取与 `mamba2_attn.py` 的推导替换、`gemma4_mm.py` 的 per-device 索引缓存模式、`diffusion_gemma.py` 的 `init_canvas` 契约收紧。若你维护多模态模型或自定义 attention backend, 需要对照检查自己的 `MultiModalFieldConfig` 与元数据构建是否也存在同类同步。

功能与动机

模型执行主路径上每个 step 都可能出现 GPU 与 CPU 之间的阻塞同步 (D2H 或 H2D), 这会打断流水线、拉高单步延迟。PR body 明确说明这些改动的每一条都是消除主机往返而不是抑制它: 'Each of these removes a host roundtrip rather than suppressing it'。具体动机包括: kv-sharing fast prefill 中每步的 `.max().item() / .sum().item()` 对设备张量产生 D2H 同步; mamba2 的 `torch.any(has_initial_states_p).item()` 同样需要设备回读; EPLB 的 `_pad_out_tensor` 漏了 `non_blocking=True` 阻塞了异步 worker 线程; 多个多模态 processor 把 token-id 和 size 字段搬到设备上只是为了立刻读回来。

实现拆解

按 PR body 拆解的六个方向, 实现过程可归纳为以下几步:

1. kv-sharing fast prefill: 元数据改由主机侧推导。在 `vllm/v1/attention/backends/utils.py` 的 `make_kv_sharing_fast_prefill_common_attn_metadata` 中, `total_num_decode_tokens` 直接用 `num_logits_indices` (本来就是 Python int), `decode_max_query_len` 改为读取新增的 `common_attn_metadata.max_logits_per_req`, 去掉每步对设备张量的 `.max().item()` 与 `.sum().item()`。配套在 `vllm/v1/attention/backend.py` 的 `CommonAttentionMetadata` 中新增 `max_logits_per_req` 字段, 并在 `vllm/v1/worker/gpu_model_runner.py` 中把 `_prepare_inputs` 的返回值扩展为三元组, 将主机上已有的 `num_sampled_tokens` 最大值一路下传到 `_build_attention_metadata`。
2. mamba2: prefill 初始状态判定改由 CPU 元数据推导。在 `vllm/v1/attention/backends/mamba_attn.py` 中把原来 `_build_chunk_metadata_tensors`

内基于 CPU 数据的 prefill 上下文长度推导抽成可复用的 `_prefill_cpu_metadata`；随后在 `vllm/v1/attention/backends/mamba2_attn.py` 的 `build` 中，用 `(num_computed_tokens_p_cpu > 0).any()` 替代 `torch.any(common.has_initial_states_p).item()`，避免 D2H 同步。关键前提是 `seq_lens_cpu_upper_bound` 对 prefill 行在所有模式下（含 `async spec decode`）都是精确值。

3. EPLB：补齐非阻塞拷贝。`_pad_out_tensor` 缺少 `non_blocking=True`，导致逻辑映射的提交阻塞异步 worker 线程，本次补上后提交不再同步阻塞。
4. gemma4 / diffusion_gemma：将一次性索引搬运改为设备端操作。`gemma4_mm.py` 的 `compute_logits` 中，把每步用 Python 列表做 `logits[:, self._suppress_token_ids] = -inf` 改为：首次按设备缓存 `async_tensor_h2d` 构造的 `suppressed-token` 索引张量，之后每步用 `logits.index_fill_(1, suppress_idx, -inf)`。`diffusion_gemma.py` 的 `init_canvas` 签名从 `np.ndarray` 改为已就位的 `torch.Tensor`，`_finish_prefills` 中把 slot 索引一次性 `async_copy_to_gpu` 到 `ps_gpu`，后续 `init_canvas(ps_gpu)`、`draft_tokens[ps_gpu] = states.canvas[ps_gpu]`、`is_encoder_phase.index_fill_(0, ps_gpu, False)` 全部复用；`add_request / remove_request` 里的标量赋值改为 `.fill_()`，新 slot 索引用 `async_tensor_h2d` 上传。
5. glm4_1v / phi4mm：替换 H2D 与类型转换同步。以 `async_tensor_h2d` 替代 `torch.tensor(..., device=)` 的同步 H2D 拷贝，并用 `.bool()` 替代 `.type(torch.BoolTensor)` 的主机往返。
6. 多模态 processor：标记 `keep_on_cpu=True`。`internvl`、`interns1`、`phi3v`、`step3_vl`、`keye`、`audioflamingo3` 的 `MultiModalFieldConfig` 中，对 token-id 与 size 类字段（如 `image_token_id`、`video_token_id`、`image_sizes`、`num_patches`、`image_grid_thw`、`video_grid_thw`）加上 `keep_on_cpu=True`，避免被 staging 到设备后又立即读回 CPU；消费侧循环改为先 `.tolist()` 再迭代，例如 `keye.py` 的 `_process_image_input`、`step3_vl.py` 的 `_process_image_input`、`phi3v.py` 的 `hd_feature_transform`。

测试与部署配套：本次没有新增或修改测试文件，也未涉及配置、schema 或部署脚本；验证主要依赖 CI 三轮回跑（Buildkite #83228/#83293/#83425）。

关键文件：

- `vllm/v1/attention/backends/mamba_attn.py`（模块 注意力后端；类别 source；类型 core-logic；符号 `_build_chunk_metadata_tensors`，`_prefill_cpu_metadata`）：将 CPU 侧 prefill 上下文长度推导从 `_build_chunk_metadata_tensors` 抽成独立 `_prefill_cpu_metadata`，成为 `mamba/mamba2` 两个 backend 共享的无同步推导基座
- `vllm/v1/attention/backends/mamba2_attn.py`（模块 注意力后端；类别 source；类型 core-logic；符号 `build`）：`prep_initial_states` 判定从设备张量的 `torch.any().item()` 改为 CPU 推导，消除每步 D2H 同步
- `vllm/v1/worker/gpu_model_runner.py`（模块 模型执行器；类别 source；类型 data-contract；符号 `_prepare_inputs`，`_build_attention_metadata`）：将 `max_num_sampled_tokens` 从主机侧下传到 `CommonAttentionMetadata.max_logits_per_req`，是 kv-sharing fast prefill 消除同步的数据契约改动

关键符号: _prefill_cpu_metadata, _build_chunk_metadata_tensors, build, _prepare_inputs, _build_attention_metadata, init_canvas, compute_logits, _finish_prefills, make_kv_sharing_fast_prefill_common_attn_metadata, _process_image_input, hd_feature_transform, _get_mm_fields_config, _keye_field_config

关键源码片段

vllm/v1/attention/backends/mamba_attn.py

将 CPU 侧 prefill 上下文长度推导从 _build_chunk_metadata_tensors 抽成独立 _prefill_cpu_metadata, 成为 mamba/mamba2 两个 backend 共享的无同步推导基座

```
# vllm/v1/attention/backends/mamba_attn.py
# 将 prefill 上下文长度推导从设备侧计算中剥离, 全部基于 CPU 已有数据。
# 关键前提: seq_lens_cpu_upper_bound 对 prefill 行在所有模式下
# (包括 async spec decode) 都是精确值, 因此无需 D2H 同步。
def _prefill_cpu_metadata(
    self,
    common: M,
    common_attn_metadata: CommonAttentionMetadata,
) -> tuple[torch.Tensor, torch.Tensor]:
    """从 CPU 数据推导 prefill 上下文长度与 query 偏移。

    Returns (num_computed_tokens_p_cpu, query_start_loc_p_cpu)。
    """
    seq_lens_cpu = common_attn_metadata.seq_lens_cpu_upper_bound
    assert seq_lens_cpu is not None
    num_reqs = common.num_reqs
    num_prefills = common.num_prefills
    # 取 query_start_loc_cpu 尾部 prefill 段, 减去 decode token 数即 prefill 起点
    query_start_loc_p_cpu = (
        common_attn_metadata.query_start_loc_cpu[-num_prefills - 1 :]
        - common.num_decode_tokens
    )
    # prefill 各请求的 query 长度 = 相邻偏移差
    prefill_query_lens_cpu = query_start_loc_p_cpu[1:] - query_start_loc_p_cpu[:-1]
    # 已计算 token 数 = 序列长度 - 本次 query 长度
    num_computed_tokens_p_cpu = (
        seq_lens_cpu[num_reqs - num_prefills : num_reqs] - prefill_query_lens_cpu
    )
    return num_computed_tokens_p_cpu, query_start_loc_p_cpu

def _build_chunk_metadata_tensors(
    self,
    chunk_size: int,
    common: M,
    common_attn_metadata: CommonAttentionMetadata,
) -> tuple[torch.Tensor, torch.Tensor, torch.Tensor]:
    """计算 chunk 元数据并返回设备张量。
```

```

Returns (cu_chunk_seqlen_p, seq_idx_p, last_chunk_indices_p)。
"""
num_prefills = common.num_prefills
# 复用在 CPU 上完成的所有推导，不再触碰设备张量
num_computed_tokens_p_cpu, query_start_loc_p_cpu = self._prefill_cpu_metadata(
    common, common_attn_metadata
)
cu_chunk_seqlen, seq_idx, last_chunk_indices = self._compute_chunk_metadata(
    chunk_size,
    num_prefills,
    num_computed_tokens_p_cpu,
    query_start_loc_p_cpu,
)
device = common_attn_metadata.query_start_loc.device
# 在 pinned CPU 内存上构建后非阻塞上传，避免 torch.as_tensor(list, device=cuda)
# 触发的同步 H2D 拷贝
cu_chunk_seqlen_p = async_tensor_h2d(
    cu_chunk_seqlen, dtype=torch.int32, device=device
)
seq_idx_p = async_tensor_h2d(seq_idx, dtype=torch.int32, device=device)
last_chunk_indices_p = async_tensor_h2d(
    last_chunk_indices, dtype=torch.int32, device=device
)
return cu_chunk_seqlen_p, seq_idx_p, last_chunk_indices_p

```

vllm/v1/attention/backends/mamba2_attn.py

prep_initial_states 判定从设备张量的 torch.any().item() 改为 CPU 推导，消除每步 D2H 同步

```

# vllm/v1/attention/backends/mamba2_attn.py
# build 中 prefill 初始状态判定：原来需要 torch.any(has_initial_states_p).item()
# 触发 D2H 同步，现在完全从 CPU 数据推导。
if common.num_prefills > 0:
    prep_initial_states = False
    if common.has_initial_states_p is not None:
        # 与 has_initial_states_p 等价的 CPU 侧条件：
        # 任一 prefill 请求有已计算 token 即需要准备初始状态。
        # seq_lens_cpu_upper_bound 对 prefill 行精确，因此无需 D2H。
        num_computed_tokens_p_cpu, _ = self._prefill_cpu_metadata(
            common, common_attn_metadata
        )
        prep_initial_states = bool((num_computed_tokens_p_cpu > 0).any())

# chunk 元数据同样走 CPU 推导 + 非阻塞上传路径
cu_chunk_seqlen_p, seq_idx_p, last_chunk_indices_p = (
    self._build_chunk_metadata_tensors(
        self.chunk_size,
        common,
        common_attn_metadata,

```

```
)  
)
```

vllm/v1/worker/gpu_model_runner.py

将 `max_num_sampled_tokens` 从主机侧下传到 `CommonAttentionMetadata.max_logits_per_req`，是 kv-sharing fast prefill 消除同步的数据契约改动

```
# vllm/v1/worker/gpu_model_runner.py  
# 关键数据流（示意片段，省略 _prepare_inputs 中部逻辑）：  
# 主机侧已知的 max_num_sampled_tokens 一路下传，  
# 替代原来对设备张量执行 .max().item() 的 D2H 同步。  
  
# 1. _prepare_inputs 的返回点：num_sampled_tokens 是 host 数组，  
# int(num_sampled_tokens.max()) 在主机侧即可得到，零同步。  
return (  
    logits_indices,  
    spec_decode_metadata,  
    int(num_sampled_tokens.max()),  
)  
  
# 2. execute_model 中接收三元组后传给 _build_attention_metadata  
logits_indices, spec_decode_metadata, max_num_sampled_tokens = (  
    self._prepare_inputs(scheduler_output, num_scheduled_tokens_np)  
)  
  
# 3. _build_attention_metadata 内写入 CommonAttentionMetadata，  
# utils 侧直接读取，免去对 decode 张量的 .max()/sum() 回读。  
if logits_indices is not None and self.cache_config.kv_sharing_fast_prefill:  
    cm_base.num_logits_indices = logits_indices.size(0)  
    cm_base.max_logits_per_req = max_num_sampled_tokens  
    cm_base.logits_indices_padded = self._prepare_kv_sharing_fast_prefill(  
        logits_indices  
    )
```

评论区精华

Review 与讨论整体非常收敛：WoosukKwon 直接 APPROVED，claude[bot] 因 fork 来源自动跳过 review（提示 maintainer 可手动触发）。PR 内没有实质性的技术争议线程，review 评论为空。PR body 中作者自己给出了最重要的技术约束：对 mamba2 的 CPU 推导，明确依赖 `seq_lens_cpu_upper_bound` 在 prefill 行上精确这一前提（'including async spec decode'）；对 kv-sharing fast prefill，明确指出 `num_decode_tokens` 是 `logits_indices` 的直方图，因此总和就是 `num_logits_indices` 这个已知 Python int。这些是本次改动的正确性论据，而非外部讨论残留。

- 暂无高价值评论线程

风险与影响

- 风险：

1. 语义等价性风险（最高）：mamba2 的 `prep_initial_states` 判定从 `has_initial_states_p` 改为 `(num_computed_tokens_p_cpu > 0).any()`，依赖 '`seq_lens_cpu_upper_bound` 对 `prefill` 行精确' 这一前提。若未来出现 `prefill` 行不精确的场景（例如新的异步投机解码路径），可能产生与设备端真实状态不一致的判定，导致初始状态漏加载。同样，`kv-sharing fast prefill` 的 `decode_max_query_len` 改用 `num_sampled_tokens.max()` 下传，若调用链上某个分支未更新 `max_num_sampled_tokens`，会退化为断言失败或错误元数据。
2. 数据契约变更风险：CommonAttentionMetadata 新增 `max_logits_per_req`，`_prepare_inputs` 返回类型从二元组变三元组、`init_canvas` 参数从 `np.ndarray` 改为 `torch.Tensor`，MultiModalFieldConfig 新增 `keep_on_cpu` 语义。这些是跨模块签名变更，任何遗漏调用点（如其他 runner 或自定义 processor）都会在运行时暴露，且本次没有测试覆盖来兜底。
3. 多模态字段语义风险：`keep_on_cpu=True` 意味着这些字段不再出现在设备端 staging 中，若某条处理路径仍期望它们在 device 上（例如直接交给视觉塔的 `image_grid_thw`），会出现 device 不匹配。`keye` 的 `_process_image_input` 改用 `.tolist()` 后，`image_grid_thw` 只用于 CPU 循环，风险较低，但 `internvl` 等文件改动较小，需确认没有下游设备消费。
4. 缓存一致性风险：gemma4 的 `_suppress_token_ids_cache` 按 `logits.device` 缓存索引张量，若 `_suppress_token_ids` 内容可变（当前是静态集合），缓存会过期；该风险当前不存在，但属于隐藏约束。
5. 回归风险：无直接测试文件配套，CI 只有三轮回跑记录，对多模态 processor 的修改（`internvl`、`phi3v` 等）缺少针对性回归验证。
 - 影响：影响范围是 vLLM V1 引擎的模型执行主路径：每个 `decode/prefill step` 上减少若干次 GPU-CPU 阻塞同步，直接降低单步延迟与 CPU 等待，对 `kv-sharing fast prefill`、mamba2/SSM 模型、gemma4/diffusion 系列、GLM-4.1V/Phi-4-mm 以及 InternVL/Phi-3-V/Step3-VL/Keye/AudioFlamingo3 等多模态模型的在线服务吞吐与首 token 延迟有正面作用。团队侧，CommonAttentionMetadata 与多个 model processor 的数据契约被收紧，后续开发需要遵循 '主机侧可推导的数据不要搬设备' 这一约定；这也是 VLLM_GPU_SYNC_CHECK 系列工作的组成部分，为未来更系统的同步审计铺路。影响程度中等偏高，但扩散在多个模型文件上、单点风险可控。
 - 风险标记：核心路径变更，元数据推导依赖精确性假设，数据契约跨模块改动，缺少测试覆盖

关联脉络

- PR #43107 VLLM_GPU_SYNC_CHECK（上游 sync-check 分支）：PR body 明确说明这些修复是从 <https://github.com/vllm-project/vllm/pull/43107> 提取出来的，属于同一 GPU 同步消除主线
- PR #46747 [Bugfix][V1][Multimodal] Recover from P0/P1 processor cache drift (#46747)：同为 vllm/v1 多模态执行链路相关，一个管 processor 缓存一致性，一个管多模态字段搬运同步，属于同一多模态执行路径的持续加固
- PR #51840 [Bugfix][TieredOffloading]：Return HIT_PENDING when KV promotion is triggered：同为 vllm/v1 执行路径的性能 / 同步优化，KV 提升触发时避免多余前缀扫描，

与本 PR 都在减少每 step 的不必要工作量

- PR #51865 [Bugfix][MRV2] Require all requests to be decoding for uniform-decode dispatch: 同为 vllm/v1/worker/gpu 路径的批处理与元数据构建改动, 与本 PR 的 _prepare_inputs/execute_model 改动重叠在 V1 worker 执行路径上