

PR #49731 完整报告

vllm-project/vllm

[Spec Decode][Perf] Replicate DSpark Markov head across TP ranks

合并时间: 2026-07-29 23:43

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49731>

执行摘要

- 一句话: 复制 Markov head 权重消除 TP 通信开销
- 推荐动作: 值得精读。本 PR 展示了如何通过复制轻量级 head 权重来消除频繁 TP 通信的典型模式, 同时保持了与现有 LogitsProcessor 后处理的兼容。VocabParallelEmbedding 的 `disable_tp` 参数设计可作为类似优化 (如 MoE 路由网络) 的参考。Review 中关于 Gemma4 兼容性的及时发现也说明充分回归测试的重要性。

功能与动机

DSpark Markov head 在解码每个 draft token 时需顺序执行 embedding 和投影, 若使用 TP 分片则每次都需要 all-reduce 和全词汇 gather, 成为性能瓶颈。PR body 明确说明 “Replicate the DSpark Markov embedding and projection across TP ranks, removing an all-reduce and full-vocabulary gather per draft position.”

实现拆解

1. VocabParallelEmbedding 新增 `disable_tp` 能力 (vllm/model_executor/layers/vocab_parallel_embedding.py) : 在 `__init__` 中增加 `disable_tp` 关键字参数。启用时强制 `tp_rank=0`, `tp_size=1`, 并调用新方法 `update_param_tp_status()` 更新所有 vLLM 参数的 `tp_rank` / `tp_size`, 确保权重重载机制兼容。forward 中根据 `self.tp_size > 1` 决定是否执行 all-reduce。
2. DSparkMarkovHead 全面复制 (vllm/model_executor/models/qwen3_dspark.py) : 将 `markov_w1` 从 VocabParallelEmbedding 替换为普通 `nn.Embedding` (天然复制); `markov_w2` (ParallelLMHead) 传入 `disable_tp=True`, 避免分片。`bias` 方法保留 `logits_processor` 参数, 确保 `soft_cap` / `scale` 等后处理被正确应用。移除旧的 TODO 注释。
3. LogitsProcessor 条件化 gather (vllm/model_executor/layers/logits_processor.py) : `_get_logits` 中原先无条件调用 `_gather_logits`, 改为仅在 `lm_head.tp_size > 1` 时调用; `get_top_tokens` 中原先通过全局函数 `get_tensor_model_parallel_world_size()` 获取 `tp_size`, 改为直接使用 `lm_head.tp_size`。减少对全局状态的依赖。
4. 配套测试 (tests/v1/sample/test_head_dtype.py) : 新增 `test_replicated_lm_head_skips_tp_communication_and_preserves_processing`, 构造 `ParallelLMHead(disable_tp=True)` 并 mock `tp_size=2`, 验证 `_gather_logits` 和 `tensor_model_parallel_all_gather` 均未被调用, 且输出与手动计算一致。同时调整已有测

试以适配 `_FakeLmHead` 新增的 `tp_size` 属性。

关键文件：

- `vllm/model_executor/layers/vocab_parallel_embedding.py`（模块 嵌入层；类别 `source`；类型 `data-contract`；符号 `update_param_tp_status`）：新增 `disable_tp` 参数和 `update_param_tp_status` 方法，为 `ParallelLMHead` 提供全复制模式，是本次优化的基础设施。
- `vllm/model_executor/models/qwen3_dspark.py`（模块 `DSpark` 模型；类别 `source`；类型 `data-contract`；符号 `DSparkMarkovHead.bias`）：`DSparkMarkovHead` 改用全复制权重：`markov_w1` 替换为 `nn.Embedding`，`markov_w2` 启用 `disable_tp=True`，彻底消除 TP 通信；`bias` 方法继续使用 `LogitsProcessor` 确保正确处理 `soft_cap` 等。
- `tests/v1/sample/test_head_dtype.py`（模块 测试；类别 `test`；类型 `test-coverage`；符号 `test_replicated_lm_head_skips_tp_communication_and_preserves_processing`）：新增测试验证复制后 LM head 跳过 TP 通信且正确应用后处理，并调整已有测试适配 `FakeLmHead` 新增的 `tp_size` 属性。
- `vllm/model_executor/layers/logits_processor.py`（模块 日志处理器；类别 `source`；类型 `data-contract`）：`LogitsProcessor` 的 `_get_logits` 和 `get_top_tokens` 改为根据 `lm_head.tp_size` 条件执行 `gather`，减少对全局函数的依赖。

关键符号：`DSparkMarkovHead.init`, `DSparkMarkovHead.bias`,
`VocabParallelEmbedding.init`, `VocabParallelEmbedding.update_param_tp_status`,
`ParallelLMHead.init`, `LogitsProcessor._get_logits`, `LogitsProcessor.get_top_tokens`,
`test_replicated_lm_head_skips_tp_communication_and_preserves_processing`

关键源码片段

`vllm/model_executor/layers/vocab_parallel_embedding.py`

新增 `disable_tp` 参数和 `update_param_tp_status` 方法，为 `ParallelLMHead` 提供全复制模式，是本次优化的基础设施。

```
class VocabParallelEmbedding(PluggableLayer):
    def __init__(
        self,
        ...,
        *, disable_tp: bool = False, # 新增：完全复制模式
    ):
        super().__init__()
        self.disable_tp = disable_tp
        if disable_tp:
            # 完全复制模式下，强制 tp_rank=0, tp_size=1
            tp_rank, self.tp_size = 0, 1
        else:
            # 正常 TP 分片模式
            tp_rank = get_tensor_model_parallel_rank()
            self.tp_size = get_tensor_model_parallel_world_size()
        self.tp_rank = tp_rank
```

```

# ... 其他初始化 ...
# 初始化后更新所有参数的 tp 状态
self.update_param_tp_status()

def update_param_tp_status(self):
    """确保所有 vLLM 参数的 tp_rank 和 tp_size 与层一致。
    权重重载需要靠此信息正确分片或复制。"""
    for param in self.parameters():
        if isinstance(param, BasevLLMParameter):
            param.tp_rank = self.tp_rank
            param.tp_size = self.tp_size

def forward(self, input_):
    # ... 原有逻辑 ...
    if self.tp_size > 1:
        # 仅当启用 TP 时才执行 all-reduce
        output_parallel.masked_fill_(input_mask.unsqueeze(-1), 0)
        return tensor_model_parallel_all_reduce(output_parallel)
    return output_parallel

```

vllm/model_executor/models/qwen3_dspark.py

DSparkMarkovHead 改用全复制权重：markov_w1 替换为 nn.Embedding，markov_w2 启用 disable_tp=True，彻底消除 TP 通信；bias 方法继续使用 LogitsProcessor 确保正确处理 soft_cap 等。

```

class DSparkMarkovHead(nn.Module):
    """Sequential transition-bias head (low-rank  $V \times r, r \times V$ ).

    Both weights are replicated because the head runs sequentially for every
    draft position. Sharding them would add an all-reduce and a full-vocab
    gather to each position.
    """
    def __init__(
        self,
        vocab_size: int,
        draft_vocab_size: int,
        markov_rank: int,
        prefix: str,
    ) -> None:
        super().__init__()
        # 使用 nn.Embedding 实现全复制，避免 TP 分片
        self.markov_w1 = nn.Embedding(vocab_size, markov_rank)
        # ParallelLMHead 启用 disable_tp=True，保持完整权重
        self.markov_w2 = ParallelLMHead(
            draft_vocab_size,
            markov_rank,
            bias=False,
            prefix=maybe_prefix(prefix, "markov_w2"),
            disable_tp=True,

```

)

```
def embed(self, token_ids: torch.Tensor) -> torch.Tensor:
    """r-dim Markov embedding of token_ids ([B] -> [B, r])."""
    return self.markov_w1(token_ids)

def bias(
    self,
    markov_embed: torch.Tensor,
    logits_processor: LogitsProcessor,
) -> torch.Tensor:
    """Vocab-size transition bias from a Markov embedding ([B, r] -> [B, V]).
    经过 logits_processor 保证 soft_cap 和 scale 被正确处理。"""
    return logits_processor(self.markov_w2, markov_embed)
```

评论区精华

Review 中主要围绕三个关键问题展开：

- `_freeze` 参数的必要性：benchislett 指出 vLLM 运行在推理模式，无需梯度，询问为何需要 `_freeze`。mgoin 承认这是仿照其他位置的写法，属于过度设计，已在后续提交中移除。
- Bias 方法是否应继续使用 LogitsProcessor：benchislett 质疑直接使用 `markov_w2` 而不经过程 LogitsProcessor 可能会遗漏 `soft_cap` / `scale` 等处理。mgoin 确认这会破坏 Gemma4 DSpark。最终保留了 `logits_processor` 参数，`bias` 方法保持原接口。
- `update_param_tp_status` 的用途：benchislett 询问新增方法的作用。mgoin 解释这是为了兼容权重重载机制，参照 PR#48025 中 LinearBase 的做法，确保参数 `tp_rank` / `tp_size` 与层一致，无争议。
 - `_freeze` 参数的必要性 (design)：mgoin 承认这只是仿照其他位置的写法，属于过度设计，已在后续提交中移除。
 - Bias 方法是否应继续使用 LogitsProcessor (correctness)：最终保留 `logits_processor` 的调用，`bias` 方法保持原接口。
 - `update_param_tp_status` 的作用 (design)：被接受，无争议。

风险与影响

- 风险：
 - 跨模型兼容性：LogitsProcessor 的行为改为依赖 `lm_head.tp_size` 而非全局函数，若其他自定义模型未正确设置 `tp_size` 属性（例如通过非标准方式创建 VocabParallelEmbedding），可能导致 gather 逻辑错误。已在 `_FakeLmHead` 测试类中新增 `tp_size=1` 默认值，降低风险。
 - 权重加载与重载：新增 `update_param_tp_status` 方法在 `__init__` 末尾调用，若下游模型在构造后手动修改参数 `tp_rank` / `tp_size`，可能被覆盖。但 PR 作者声明这是为了与 PR#48025 的权重重载兼容，当前实现与 LinearBase 一致，风险可控。
 - 性能回归：仅对 `disable_tp=True` 的 ParallelLMHead 生效，默认行为不变。但若错误地在普通 LM head 上启用 `disable_tp`，将失去 TP 内存节省，需通过文档或断言保护。

- 影响：
 - 用户影响：仅影响使用 DSpark 模型（如 Qwen3-14B DSpark）的用户，获得 3–4% 端到端吞吐提升，无需额外配置。Gemma4 DSpark 用户需确认后续提交已修复兼容性。
 - 系统影响：对非 DSpark 模型，因 `disable_tp` 默认为 `False`，行为完全不变。
 - 团队影响：为后续类似“复制小头”优化提供了可复用的 `disable_tp` 模式，LogitsProcessor 的 `tp_size` 属性简化了未来重构。
 - 风险标记：核心路径变更，跨模型兼容性，需要回归测试

关联脉络

- 暂无明显关联 PR