

PR #47070 完整报告

vllm-project/vllm

[Feature] Support sequence parallel without the need for DP, 1.9%~5.0% E2E Throughput Improvement

合并时间: 2026-07-05 20:41

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47070>

执行摘要

- 一句话: 移除 SP 必须依赖 DP 的限制, 提升 MoE 吞吐 1.9%-5%
- 推荐动作: 建议关注并行配置的相关变更, 特别是 DP=1 时启用 SP 的新能力。对于使用 MoE 且未开启 DP 的用户, 推荐尝试本 PR 以获取性能提升。同时注意 LoRA 与 SP 的互斥约束。

功能与动机

原设计中序列并行必须依赖数据并行, 但该约束并非必需。去除后可在无 DP 时启用 SP, 进一步提高 TP 下 MoE 模型的性能。来自 PR body: 'Originally we must have DP to enable sequence parallel, but this constraint is not needed, this PR fixes the issue.'

实现拆解

1. 并行配置解耦: vllm/config/parallel.py 的 use_sequence_parallel_moe 属性中去掉 data_parallel_size > 1 条件, 使 SP 可独立于 DP 启用。
2. 前向上下文适配: vllm/forward_context.py 放宽 DPMetadata.make 断言和 set_forward_context 条件, 使 SP 场景也能创建 DPMetadata; 当 DP 大小为 1 时构造单元素 num_tokens_across_dp。
3. 通信器初始化扩展: base_device_communicator.py 中 EP all2all 管理器的启用条件增加 use_sequence_parallel_moe, 确保 SP 下正确初始化。
4. 模型层调整: deepseek_v2.py 中 SP padding 计算简化为负数取模, 消除条件分支; gpt_oss.py 的 MLPBlock.__init__ 中增加 and vllm_config.lora_config is None 以禁用 LoRA 与 SP 同时使用。
5. MoE 运行时分发优化: moe_runner.py 的 do_naive_dispatch_combine 条件从 dp_size > 1 扩展为 dp_size > 1 or is_sequence_parallel, 使 SP 下 naive dispatch 正确工作。
6. 无直接测试文件: 变更未包含独立测试, 依赖现有集成测试和手动基准验证。

关键文件:

- vllm/config/parallel.py (模块 并行配置; 类别 source; 类型 core-logic; 符号 use_sequence_parallel_moe): 核心配置属性 use_sequence_parallel_moe 移除 DP 强制要求, 是解耦的关键。

- `vllm/forward_context.py` (模块 前向上下文; 类别 source; 类型 core-logic; 符号 `DPMetadata.make`, `set_forward_context`) : 修改了 `DPMetadata` 创建条件和 `set_forward_context` 逻辑, 使 SP 独立于 DP 初始化元数据。
- `vllm/model_executor/models/deepseek_v2.py` (模块 Deepseek 模型; 类别 source; 类型 data-contract; 符号 forward) : 简化 SP 下的 padding 计算, 去除条件分支, 提升性能。
- `vllm/model_executor/models/gpt_oss.py` (模块 GPT-OSS 模型; 类别 source; 类型 data-contract; 符号 `MLPBlock.init`) : 确保 LoRA 与 MoE SP 互斥, 避免运行时错误。
- `vllm/distributed/device_communicators/base_device_communicator.py` (模块 通信器; 类别 source; 类型 core-logic; 符号 init) : 扩展 EP all2all 的触发条件到序列并行 MoE, 确保通信管理器正确初始化。
- `vllm/model_executor/layers/fused_moe/runner/moe_runner.py` (模块 MoE 运行时; 类别 source; 类型 data-contract; 符号 `do_naive_dispatch_combine`) : 使 naive dispatch combine 在 SP 下也能正确工作。
- `vllm/model_executor/layers/fused_moe/config.py` (模块 MoE 配置; 类别 source; 类型 data-contract) : 配套配置属性调整。

关键符号: `use_sequence_parallel_moe`, `DPMetadata.make`, `set_forward_context`, `forward`, `MLPBlock.init`, `do_naive_dispatch_combine`, `init(base_device_communicator)`

关键源码片段

`vllm/forward_context.py`

修改了 `DPMetadata` 创建条件和 `set_forward_context` 逻辑, 使 SP 独立于 DP 初始化元数据。

```
# From vllm/forward_context.py, set_forward_context() context manager
```

```
dp_metadata: DPMetadata | None = None
if (
    (
        vllm_config.parallel_config.data_parallel_size > 1
        or vllm_config.parallel_config.use_sequence_parallel_moe
        # NOTE: now sequence parallel MoE also initializes DPMetadata
        # because it needs to track per-rank token counts
    )
    and vllm_config.parallel_config.is_moe_model is not False
    and (attn_metadata is not None or num_tokens is not None)
):
    if (
        num_tokens_across_dp is None
        and vllm_config.parallel_config.data_parallel_size > 1
    ):
        # If DP > 1, coordinate tokens across DP ranks via all-reduce
        assert ubatch_slices is None
        assert num_tokens is not None
        _, num_tokens_across_dp, _ = coordinate_batch_across_dp(
            num_tokens_unpadded=num_tokens,
```

```

        parallel_config=vllm_config.parallel_config,
        allow_microbatching=False,
    )
    assert num_tokens_across_dp is not None
elif num_tokens_across_dp is None:
    # For SP without DP, create a single-entry tensor with all tokens
    assert num_tokens is not None
    num_tokens_across_dp = torch.tensor([num_tokens], dtype=torch.int32)
dp_metadata = DPMetadata.make(
    vllm_config.parallel_config, num_tokens or 0, num_tokens_across_dp
)

```

vllm/model_executor/models/deepseek_v2.py

简化 SP 下的 padding 计算，去除条件分支，提升性能。

```
# From DeepseekV2DecoderLayer.forward, sequence parallel MoE section
```

```

if self.use_sequence_parallel_moe:
    tp_world_size = get_tensor_model_parallel_world_size()
    # small trick using minus, eg. -17 % 8 = 7
    sp_pad = (-hidden_states.shape[0]) % tp_world_size
    # pad if not divisible by world size
    hidden_states = torch.nn.functional.pad(hidden_states, (0, 0, 0, sp_pad))
    hidden_states = tensor_model_parallel_reduce_scatter(hidden_states, 0)
    if not input_is_sequence_parallel:
        residual = sequence_parallel_chunk(residual)

```

评论区精华

在后续修复 PR #47290 中，gcanlin 发现序列并行与 LoRA 不兼容，因而在 `gpt_oss.py` 的 MLPBlock 初始化中加入了 `vllm_config.lora_config is None` 检查，确保启用 SP 时禁用 LoRA，避免运行时错误。

- LoRA 与 MoE 序列并行冲突 (bugfix): 在 `gpt_oss.py` 中添加 `vllm_config.lora_config is None` 检查，启用 SP 时禁用 LoRA。

风险与影响

- 风险：核心路径变更：修改了并行条件，影响所有 MoE 模型的前向。DP=1 但启用 SP 时，之前未初始化 DPMetadata，现在会初始化单元素 tensor，需要确保所有使用 DPMetadata 的地方都正确处理。缺少测试覆盖：本次 PR 未包含直接测试文件，仅依赖集成测试和手动 benchmark。LoRA 与 SP 的互斥修复是后来补充的，可能存在其他未发现的兼容性问题。
- 影响：用户：之前必须同时开启 DP 和 TP 才能使用 SP 的用户现在可以仅用 TP，带来 1.9-5% 吞吐提升，尤其适合未配置 DP 的场景。系统：减少冗余 allreduce 通信，降低开销。团队：对 DeepSeek 和 GPT-OSS 模型用户影响最大，建议更新文档和配置示例。
- 风险标记：核心路径变更，缺少测试覆盖，配置条件修改，LoRA 兼容限制

关联脉络

- PR #47290 [BugFix for #47070] Lora should not be with MoE SP: 修复了本 PR 中 LoRA 与 SP 不兼容的问题，并合并到本分支。