

PR #24955 完整报告

sgl-project/sglang

Support Nemotron DP attention and MTP

合并时间: 2026-06-13 03:21

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/24955>

执行摘要

- 一句话: 为 Nemotron-H 模型启用 DP attention 和 MTP
- 推荐动作: 该 PR 的架构设计 (通过 LayerCommunicator 解耦布局转换) 值得精读。适合关注混合模型 DP 推理、推测解码与数据并行结合的工程师。建议在集成后尽快补充单元测试和集成测试。

功能与动机

Nemotron-H 是混合模型 (Mamba2 + full attention + MoE), 但 DP attention 此前只支持标准 attention+MLP transformer, 无法在 Nemotron-H 上启用数据并行布局, 导致 attention 和 Mamba 只能运行在全局 TP 上 (存在 KV 显存浪费)。本 PR 让 attention 和 Mamba 运行在 attention-TP 子组上, 减少 KV 和 Mamba 状态复制, 同时 MoE 保持全局 TP, 从而提升显存效率。

实现拆解

1. 抽取公共基类 NemotronHMLPLikeDecoderLayer ([nemotron_h.py](#)): 将 dense MLP 和 MoE decoder layer 的 forward 逻辑抽取到公共基类, 加入 DP attention 判断: 启用时通过 LayerCommunicator 进行 hidden_states 布局转换 (TP_ATTN_FULL ↔ FULL), 否则沿用原有 norm + mixer 路径。
2. 新增 nemotron_h_utils.py 工具模块: 包括 [is_attn_layer](#) (判断是否为 attn/Mamba 层)、[get_real_num_tokens](#) (计算非 padding 的真实 token 数)、[pad_to_original_num_tokens](#) (将输出 padding 回原始尺寸)、[make_layer_communicator](#) (构建 per-layer LayerCommunicator, 配置 scatter mode 为 attn 前 FULL、attn 后 TP_ATTN_FULL)。
3. Mamba 层适配 DP attention ([mamba.py](#)): 在 [MambaMixer2](#) 中, 当 DP attention 启用时, 将 tp_size/tp_rank 切换为 attention-TP 组, 将线性层 (conv1d、in_proj、out_proj) 的 tp_rank/tp_size 参数化, 并设置 [reduce_results=False](#)。forward 中处理 padding: 截断 hidden_states_B_C 和 dt 到真实 token 数, 确保 padding 行不参与状态更新。
4. MTP 解码层适配 ([nemotron_h_mtp.py](#)): 对 Attention 和 MoE 解码层的 MTP 变体做类似改造: [eh_proj](#) 的 gather_output 设为 False, 使用 attention-TP 组; forward 后若 DP 启用则进行 all_gather。此外, 增加 [prev_layer_is_attn](#) 判断以在合适时机做 attn_tp_all_reduce。

5. ForwardBatch 扩充 (`forward_batch_info.py`) : 增加 `_original_batch_size` 和 `_original_forward_mode` 字段, 用于 DP 场景下恢复原始 batch 信息; 修改 `prepare_mlp_sync_batch` 以支持 hybrid-SSM (Mamba) 模型在 idle/extend 时的特殊 padding 转换。
6. FlashInfer 后端适配 (`flashinfer_backend.py`) : 新增 `_cuda_graph_capture_max_bs` 函数, 将 cuda graph 捕获时的 max_bs 对齐到 attn_TP/CP 倍数的上取整, 避免因 DP 导致的 batch size 不对齐。

关键文件:

- `python/sglang/srt/models/nemotron_h.py` (模块 模型定义; 类别 source; 类型 core-logic; 符号 `NemotronHMLPDecoderLayer`, `NemotronHMLPLikeDecoderLayer`, `forward`, `NemotronHMoEDecoderLayer`) : 核心文件, 引入了 `NemotronHMLPLikeDecoderLayer` 基类, 将所有 decoder layer 的 forward 抽取统一, 并嵌入 DP attention 布局转换。
- `python/sglang/srt/models/nemotron_h_utils.py` (模块 工具函数; 类别 source; 类型 data-contract; 符号 `is_attn_layer`, `get_real_num_tokens`, `pad_to_original_num_tokens`, `_build_layer_scatter_modes`) : 新增工具模块, 包含 DP 注意力关键辅助函数, 如 `is_attn_layer`、`get_real_num_tokens`、`make_layer_communicator`。
- `python/sglang/srt/model_executor/forward_batch_info.py` (模块 前向批次信息; 类别 source; 类型 data-contract) : DP attention 需要保存原始 batch_size 和 forward_mode, 在 `prepare_mlp_sync_batch` 中增加了 hybrid-SSM 模型特殊的 idle/extend 转换逻辑。
- `python/sglang/srt/layers/attention/mamba/mamba.py` (模块 Mamba 层; 类别 source; 类型 dependency-wiring) : `MambaMixer2` 的 TP 组切换和 padding 行截断, 是关键依赖组件。
- `python/sglang/srt/layers/attention/flashinfer_backend.py` (模块 Attention 后端; 类别 source; 类型 core-logic; 符号 `_cuda_graph_capture_max_bs`) : cuda graph 捕获的 batch size 对齐, 避免 DP 导致索引越界。

关键符号: `NemotronHMLPLikeDecoderLayer.forward`, `NemotronHMLPDecoderLayer.init`, `NemotronHMoEDecoderLayer.init`, `NemotronHMambaDecoderLayer.forward`, `NemotronHAttnLikeDecoderLayer.forward`, `MambaMixer2.init`, `MambaMixer2.forward`, `get_real_num_tokens`, `pad_to_original_num_tokens`, `make_layer_communicator`, `_build_layer_scatter_modes`, `is_attn_layer`, `ForwardBatch.prepare_mlp_sync_batch`, `_cuda_graph_capture_max_bs`

关键源码片段

`python/sglang/srt/models/nemotron_h.py`

核心文件, 引入了 `NemotronHMLPLikeDecoderLayer` 基类, 将所有 decoder layer 的 forward 抽取统一, 并嵌入 DP attention 布局转换。

```
class NemotronHMLPLikeDecoderLayer(nn.Module):
```

```
"""Shared forward for the dense-MLP / MoE decoder layers."""
```

```
def forward(
    self,
    *,
    hidden_states: torch.Tensor,
    residual: Optional[torch.Tensor],
    forward_batch: ForwardBatch,
) -> tuple[torch.Tensor, torch.Tensor]:
    if is_dp_attention_enabled():
        # DP attention 启用时: 通过 LayerCommunicator 准备 MLP 输入 (从 TP_ATTN_FULL
        # 转为 FULL)
        hidden_states, residual = self.layer_communicator.prepare_mlp(
            hidden_states, residual, forward_batch
        )
        hidden_states = self.mixer.forward(hidden_states)
        # 后处理: 从 FULL 转回 TP_ATTN_FULL
        hidden_states, residual = self.layer_communicator.postprocess_layer(
            hidden_states, residual, forward_batch
        )
        return hidden_states, residual

    # 非 DP 路径 (与原来一致)
    if residual is None:
        residual = hidden_states
        hidden_states = self.norm(hidden_states)
    else:
        hidden_states, residual = self.norm(hidden_states, residual)

    hidden_states = self.mixer.forward(hidden_states)
    return hidden_states, residual
```

```
class NemotronHMLPDecoderLayer(NemotronHMLPLikeDecoderLayer):
    def __init__(self, ...):
        super().__init__()
        # ... 原有初始化 ...
        self.layer_communicator = make_layer_communicator(self.norm, for_attn=False)
```

python/sglang/srt/model_executor/forward_batch_info.py

DP attention 需要保存原始 batch_size 和 forward_mode, 在 prepare_mlp_sync_batch 中增加了 hybrid-SSM 模型特殊的 idle/extend 转换逻辑。

```
# 在 ForwardBatch 类中新增字段
_original_batch_size: Optional[int] = None
_original_forward_mode: Optional[ForwardMode] = None

# prepare_mlp_sync_batch 方法中的关键新增逻辑 (针对 hybrid-SSM)
hybrid_ssm = (
```

```

model_runner.mambaish_config is not None
or (
    model_runner.is_draft_worker
    and getattr(model_runner.model_config.hf_config, "mtp_hybrid_override_pattern", None) is
    not None
)
)
if hybrid_ssm and self.spec_info is not None and not self.spec_info.is_draft_input():
    if self.forward_mode.is_idle():
        # 将 idle 模式转换为 TARGET_VERIFY, 使 DP attention 可以生成 fabricate 行
        self._original_forward_mode = self.forward_mode
        self.forward_mode = ForwardMode.TARGET_VERIFY
        bs = self.batch_size = num_tokens // self.spec_info.num_tokens_per_req
    elif self.is_extend_in_batch and dp_padding_mode.is_max_len():
        self._original_forward_mode = self.forward_mode
        self.forward_mode = ForwardMode.EXTEND
    # ... 处理 extend_seq_lens 等字段 ...

```

评论区精华

- all_reduce 错误: gemini-code-assist 指出 get_attention_tp_group().all_reduce(local_sums) 返回 Work 而非 reduced 张量, 应使用 attn_tp_all_reduce。作者后续提交修复。
- helper 函数位置: b8zhong 建议将 _is_attn_layer 等辅助函数移出建模文件, 作者将其抽取到 nemotron_h_utils.py。
- model_executor 侵入: Fridge003 和 ispobock 认为应避免对共享 forward_batch_info.py 的修改, 可通过条件分支保护。作者后续逐步缩小了 model_executor 的改动范围。
- emulate_global_tp_chunks 必要性: Fridge003 问是否真正需要该 chunk 逻辑, 后来作者通过 benchmark 证明去掉后精度不变且性能提升 6%, 因此最终移除了该功能。
- dynamic attr 和 forward mode 状态翻转: ispobock 认为对 forward_mode 的隐性修改较 hacky, 作者回应已修复但未完全移除。
 - all_reduce 错误使用 (correctness): 作者在后续提交中修复为 attn_tp_all_reduce。
 - 辅助函数应外移 (design): 作者创建了 nemotron_h_utils.py 并将这些函数迁移过去。
 - model_executor 侵入风险评估 (design): 作者逐步缩减 model_executor 改动, 最终仅保留 forward_batch_info.py 的必需字段和分支逻辑。
 - emulate_global_tp_chunks 必要性 (performance): 作者通过 benchmark 验证去掉后精度不变且性能提升 6%, 最终在后续提交中移除了该逻辑。
 - dynamic attr 和 forward_mode 隐性修改 (design): 作者回应 "Good catch! fixed", 但未完全移除所有 setattr; 该设计在后续 review 中仍有争议。

风险与影响

- 风险:

- 性能退化：DP attention 本身带来额外的集体通信和 layout 转换，测试显示延迟增加约 30-40% (TP=8, DP=4 时 latency 118.3s vs 89.6s)，但提升了显存利用率（允许更大 batch size）。
- 数值精度：o_proj 的 deferred reduction 和 padding 零化可能引入数值偏差，但 GSM8K 测试结果显示无精度回归。
- Mamba 状态完整性：padding 行截断逻辑复杂：必须确保 padding token 不写入 KV cache、不更新 Mamba 状态、不污染残差。任何遗漏可能导致显存或数值错误。
- 共享基础设施侵入：forward_batch_info.py 和 flashinfer_backend.py 的修改可能影响其他模型，尤其是 _cuda_graph_capture_max_bs 的改变会影响所有使用 DP attention 的模型。
- 影响：
 - 用户影响：Nemotron-H 模型用户可以通过设置 --enable-dp-attention 启用 DP 模式，在相同显存下支持更大的 batch size 和序列长度；但单个请求延迟会有所增加。
 - 系统影响：新增了 LayerCommunicator 布局转换和 DP padding 机制，增加了前向通路的复杂性；测试覆盖率不足（PR 未包含自动化测试）。
 - 团队影响：开发者需要理解 DP attention 与 Mamba 混合模型的交互，后续维护成本中等。
 - 风险标记：共享基础设施变更，Mamba 状态处理复杂，性能下降风险，无自动化测试覆盖

关联脉络

- PR #23862 Fix --mem-fraction-static not accounting for EAGLE draft model KV cache: 同样涉及 Nemotron-H 模型和推测解码的 KV cache 内存管理，与本次的 MTP 支持有重叠。