

PR #1980 完整报告

THUDM/slime

[Fix] Fix FLOPs accounting for non-MLA attention

合并时间: 2026-05-29 20:33

原文链接: <http://prhub.com.cn/THUDM/slime/pull/1980>

执行摘要

- 一句话: 修复非 MLA 模型 FLOPs 计算错误的 bug
- 推荐动作: 值得立即合并。这是明确的计量 bugfix, 改动小、逻辑清晰, 且经过本地验证。建议后续补充单元测试, 覆盖 MLA/ 非 MLA 两种场景的 FLOPs 计算。

功能与动机

PR body 指出 Megatron 注册的 MLA 参数会在全局存在默认值 (如 `kv_lora_rank=32`), 导致非 MLA 模型 (如 GQA 的 Qwen3-4B) 进入 MLA 分支, FLOPs 统计错误。提供的调试日志明确展示了这一现象。

实现拆解

1. 新增辅助函数: 在 `slime/utils/flops_utils.py` 中添加 `_is_multi_latent_attention(args)` 函数, 通过 `getattr(args, "multi_latent_attention", False)` 安全获取布尔标志。
2. 改造 QKV 投影 FLOPs 计算 (`calculate_qkv_projection_flops`):
 - 先调用 `_is_multi_latent_attention` 获得 `is_mla`。
 - MLA 条件由原来的 `q_lora_rank is None / kv_lora_rank is None` 改为 `is_mla and q_lora_rank is not None / is_mla and kv_lora_rank is not None`。
 - 非 MLA 分支中, `q_head_dim` 使用 `args.kv_channels` 而非 MLA 的 `qk_head_dim + qk_pos_emb_head_dim`。
3. 改造 Attention FLOPs 计算 (`calculate_attention_flops`):
 - QK^T 和 $A*V$ 的分支条件从 `args.qk_pos_emb_head_dim / args.v_head_dim` 替换为 `is_mla`, 因为这些维度仅在 MLA 中有效。
4. 无测试文件新增: PR 作者仅做了本地手动验证和 `py_compile` 检查。

关键文件:

- `slime/utils/flops_utils.py` (模块 工具模块; 类别 source; 类型 core-logic; 符号 `_is_multi_latent_attention`): 核心改动文件, 新增 `_is_multi_latent_attention` 函数并修改了 `calculate_qkv_projection_flops` 和 `calculate_attention_flops` 两个关键函数的控制流。

关键符号: `_is_multi_latent_attention`, `calculate_qkv_projection_flops`, `calculate_attention_flops`

关键源码片段

slime/utils/flops_utils.py

核心改动文件，新增 `_is_multi_latent_attention` 函数并修改了 `calculate_qkv_projection_flops` 和 `calculate_attention_flops` 两个关键函数的控制流。

```
# slime/utils/flops_utils.py
```

```
def _is_multi_latent_attention(args):
```

```
    # 通过显式的 multi_latent_attention 标志判断，而非依赖 MLA 参数默认值
```

```
    return bool(getattr(args, "multi_latent_attention", False))
```

```
def calculate_qkv_projection_flops(args, seqlen, hidden_size, num_attention_heads, num_query_groups):
```

```
    is_mla = _is_multi_latent_attention(args)
```

```
    if is_mla and args.q_lora_rank is not None:
```

```
        # MLA 专用的 Q LoRA 公式
```

```
        q_flops = (
```

```
            2
```

```
            * seqlen
```

```
            * args.q_lora_rank
```

```
            * (args.hidden_size + args.num_attention_heads * (args.qk_head_dim + args.qk_pos_emb_head_dim))
```

```
        )
```

```
    else:
```

```
        # 标准 MHA/GQA 公式，q_head_dim 使用 kv_channels (非 MLA 时 qk_head_dim 无意义)
```

```
        q_head_dim = args.qk_head_dim + args.qk_pos_emb_head_dim if is_mla else args.kv_channels
```

```
        q_flops = 2 * seqlen * hidden_size * num_attention_heads * q_head_dim
```

```
    if is_mla and args.kv_lora_rank is not None:
```

```
        # MLA 专用的 KV LoRA 公式
```

```
        kv_flops = (
```

```
            2
```

```
            * seqlen
```

```
            * (
```

```
                args.kv_lora_rank
```

```
                * (args.hidden_size + args.num_attention_heads * (args.qk_head_dim + args.v_head_dim))
```

```
                + args.hidden_size * args.qk_pos_emb_head_dim
```

```
            )
```

```
        )
```

```
    else:
```

```
        # 标准 GQA 的 KV projection 公式 (2 倍是因为 K 和 V 两个投影)
```

```
        kv_flops = 2 * 2 * seqlen * hidden_size * num_query_groups * args.kv_channels
```

```
    return q_flops + kv_flops
```

```
def calculate_attention_flops(args, seqlen, num_attention_heads):
    is_mla = _is_multi_latent_attention(args)
    # QK^T 计算（因果掩码，除 2）
    if is_mla:
        # MLA 的 QK 维度包含 qk_head_dim 和 qk_pos_emb_head_dim
        flops = 2 * num_attention_heads * seqlen * seqlen * (args.qk_head_dim + args.qk_pos_emb_head_dim) / 2
    else:
        # 标准模型使用 kv_channels 作为 QK 维度
        flops = 2 * num_attention_heads * seqlen * seqlen * args.kv_channels / 2
    # A*V 计算
    if is_mla:
        flops += num_attention_heads * seqlen * seqlen * args.v_head_dim
    else:
        flops += num_attention_heads * seqlen * seqlen * args.kv_channels
    return flops
```

评论区精华

无 review 讨论。

- 暂无高价值评论线程

风险与影响

- 风险：风险极低。变更仅限 FLOPs 统计工具，不影响模型前向计算或训练正确性。但缺少单元测试覆盖，未来修改可能引入回归。
- 影响：影响范围：所有使用 calculate_fwd_flops 进行 FLOPs 统计的训练 / 评估流程。对非 MLA 模型（MHA/GQA），FLOPs 统计值将从错误的 MLA 公式修正为正确的标准公式。对 MLA 模型，行为不变。
- 风险标记：缺少测试覆盖

关联脉络

- PR #1947 feat: add FlashQLA backend for Qwen GDN ...: 该 PR 引入了对 Qwen GDN (MLA) 的支持，可能与 FLOPs 统计相关。