

PR #27065 完整报告

sgl-project/sglang

[diffusion] bug: remove boolean arithmetic guard to fix compiling

合并时间: 2026-06-10 13:55

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/27065>

执行摘要

- 一句话: 修复 FLUX.2 compile 模式的布尔运算兼容性 bug
- 推荐动作: 值得合并: 小改动无风险, 但修复了重要模型的 compile 兼容问题, 建议参与审核的工程师确认语义等价后合并。

功能与动机

FLUX.2 模型启用 compile 后, `USPAttention.forward` 中校验 `num_replicated_prefix / suffix / kv_prefix` 的布尔加法导致 `AttributeError: 'BooleanFalse' object has no attribute 'as_coeff_Mul'`, 需要修改为 compile 兼容的写法以恢复编译模式。

实现拆解

修改 `python/sglang/multimodal_gen/runtime/layers/attention/layer.py` 中 `USPAttention` 的 `forward` 方法, 将内部的 replicated-token 模式互斥校验从 `sum(bool(n) for n in ...) > 1` 改为三组显式的 `and` 条件组合, 移除 `bool` 和 `sum` 调用, 确保与 `torch.compile` 的符号追踪兼容。

关键文件:

- `python/sglang/multimodal_gen/runtime/layers/attention/layer.py` (模块 注意力层; 类别 source; 类型 core-logic): 核心文件: `USPAttention` 的 replicated-token 互斥校验由布尔算术改为显式比较, 以兼容 `torch.compile`。

关键符号: 未识别

关键源码片段

`python/sglang/multimodal_gen/runtime/layers/attention/layer.py`

核心文件: `USPAttention` 的 replicated-token 互斥校验由布尔算术改为显式比较, 以兼容 `torch.compile`。

```
# 原始代码: 使用 sum(bool(n) ...) > 1 来检查至少两个 replicated-token 模式被同时启用
# 这在 torch.compile 的符号追踪下会触发 BooleanFalse.as_coeff_Mul 错误
#
# 修改后: 显式列举三个两两组合的条件, 用 or 连接, 避免布尔值算术运算
sp_size = get_ulysses_parallel_world_size()

# 确保最多只有一个 replicated-token 模式是激活的 -> 正确性校验
```

```
if (
    (num_replicated_prefix > 0 and num_replicated_suffix > 0)
    or (num_replicated_prefix > 0 and num_replicated_kv_prefix > 0)
    or (num_replicated_suffix > 0 and num_replicated_kv_prefix > 0)
):
    raise ValueError(
        "USPAttention supports at most one replicated-token mode per call."
    )
```

评论区精华

无 Review 评论，仅 Maintainer mickqian 直接批准。

- 暂无高价值评论线程

风险与影响

- 风险：变更极小（+3/-9），仅修改校验表达式，语义不变，逻辑等价。回归风险极低。
- 影响：用户：修复 FLUX.2 在 compile 模式下崩溃的问题，提升用户体验。系统：无其他影响。兼容性：参数和接口均未改变，向后兼容。
- 风险标记：核心路径变更，低风险

关联脉络

- PR #24994 [BUG] introduced boolean arithmetic guard in attention layer: 本 PR 修复的布尔算术 bug 是由该 PR 引入的。