

PR #33785 完整报告

sgl-project/sglang

Fix Mistral-Large-3 EAGLE draft skipping DeepseekV2Model.__init__

合并时间: 2026-08-07 04:59

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/33785>

执行摘要

- 一句话: 修复 EAGLE 草稿缺失父类初始化, 并抑制预热 NaN 探针
- 推荐动作: 值得精读。推荐关注两点: 一是 commit 2 中“helper 不是真实分类”的论证与唯一初始化路径的设计取舍; 二是作者区分“引入 bug 的 commit”与“暴露 bug 的 commit”的归因方法 (nightly 时间线对比)。合并后建议跟进: 量化 draft 路径的前缀修正复验、warmup NaN 是否需要在探针或草稿 warmup 路径上修复。

功能与动机

nightly-test-general-8-gpu-b200 的 test_mistral_large3.py TP8+MTP 变体持续失败, 报 `AttributeError: 'MistralLarge3EagleModel' object has no attribute 'use_dsa'`。PR body 指出根因: `MistralLarge3EagleModel` 继承 `DeepseekV2Model` 却跳过其 `init`, 手工重建了 `embed_tokens/layers/fc`, 但仍继承 `forward`, 而 #28785 在 `forward` 首行引入 `_dsa_forward_uses_topk()` 后缺失属性成为必然崩溃; 紧接着的 `next_full_attention_layer_id`、`first_k_dense_replace`、`cp_size` 也会依次崩溃。同时 PR 揭示存在第二个被掩盖的问题: EAGLE CUDA 图捕获预热期的 NaN 断言 (#27461 默认开启探针后才开始报告)。

实现拆解

1. 定位根因: `MistralLarge3EagleModel` 继承 `DeepseekV2Model` 但用 `nn.Module.__init__` 手工重建 `embed_tokens / layers / fc` 等状态, 而 `forward` 被继承, 读取 `self.use_dsa` 等父类属性时缺失; #28785 将 `_dsa_forward_uses_topk()` 提到 `forward` 首行后崩溃显式化。涉及文件 `python/sglang/srt/models/mistral_large_3_eagle.py`。
2. 中间方案与回退: 提交 1 曾在 draft 上补 `_init_forward_attrs / _init_next_full_attention_layer_id` 等 helper; 提交 2 认为“forward 恰好读到的属性”不是可靠分类, 任何要求子类记得调用的 helper 都是同一故障模式, 于是回退 `deepseek_v2.py` 到 main, 改为调用 `super().__init__`。
3. 核心变更: 删除约 40 行手工初始化, 只保留 `fc` 与 `forward` 覆写; 父类初始化同时带来 `embed_tokens` 走 `get_embedding_tp_kwargs()` (保证与 target 的 vocab-parallel 布局一致)、修正 `add_prefix` 参数顺序 (权重名从 `layers.0.model / fc.model` 恢复为 `model.layers.0 / model.fc`)、decoder layers 获得 `alt_stream` 与 `offloader hooks`。
4. 测试配套: `test/registered/8-gpu-models/test_mistral_large3.py` 用 `envs.SGLANG_ENABLE_ASYNC_ASSERT.override(0)` 包裹 `run_combined_tests`, 抑制 EAGLE 草稿捕获预热期 NaN 断言, 注释明确这是 suppression 而非修复, 与

Nemotron-3 三个 nightly 测试既有一致。

5. 验证: B200 rerun 从 `AttributeError` 转为通过, H200 因非 Blackwell 跳过; 合并 main 时解决 `PerformanceTestParams.profile_dir -> result_dir` 的冲突。

关键文件:

- `python/sglang/srt/models/mistral_large_3_eagle.py` (模块 草稿模型; 类别 source; 类型 data-contract; 符号 `MistralLarge3EagleModel.init`, `MistralLarge3EagleModel.forward`): 核心修复文件: EAGLE draft 由手工重建状态改为调用 `DeepseekV2Model.init`, 消除 `forward` 属性缺失崩溃, 并修正 `add_prefix` 参数顺序与 `embedding` 布局。
- `test/registered/8-gpu-models/test_mistral_large3.py` (模块 夜间测试; 类别 test; 类型 test-coverage; 符号 `TestMistralLarge3.test_mistral_large3_all_variants`): 夜间测试配套: 用 `SGLANG_ENABLE_ASYNC_ASSERT.override(0)` 抑制 EAGLE 预热期 NaN 断言, 与 Nemotron-3 做法一致, 注释说明了抑制理由。

关键符号: `MistralLarge3EagleModel.init`, `MistralLarge3EagleModel.forward`, `TestMistralLarge3.test_mistral_large3_all_variants`

关键源码片段

`python/sglang/srt/models/mistral_large_3_eagle.py`

核心修复文件: EAGLE draft 由手工重建状态改为调用 `DeepseekV2Model.init`, 消除 `forward` 属性缺失崩溃, 并修正 `add_prefix` 参数顺序与 `embedding` 布局。

```
# MistralLarge3EagleModel 修复后的完整初始化与 forward 路径。
# 关键改动: 不再用 nn.Module.__init__ 手工重建 embed_tokens / layers / norm,
# 而是调用父类 DeepseekV2Model.__init__, 让 forward 依赖的所有状态
# (use_dsa、next_full_attention_layer_id、first_k_dense_replace、cp_size 等)
# 由唯一初始化路径产生, 避免子类状态与父类漂移。
class MistralLarge3EagleModel(DeepseekV2Model):
    def __init__(self, config, quant_config=None, prefix=""):
        # super().__init__ 会构建 embed_tokens (经 get_embedding_tp_kwargs 保持
        # vocab-parallel 布局与 target 一致)、DeepseekV2DecoderLayer 堆栈、
        # start_layer / end_layer、norm, 以及 DSA / CP 相关属性
        super().__init__(config, quant_config, prefix=prefix)
        assert self.pp_group.world_size == 1

        # 草稿模型唯一额外需要的是 fc: 拼接 token embedding 与 target hidden states
        self.fc = RowParallelLinear(
            config.hidden_size * 2,
            config.hidden_size,
            bias=False,
            quant_config=quant_config,
            # 修复 add_prefix 参数顺序: 旧代码写 add_prefix(prefix, "fc"),
            # 权重名变成 fc.model, quant 配置解析与 remapping (目标为 model.fc.*) 对不上
            prefix=add_prefix("fc", prefix),
            input_is_parallel=False,
```

)

```
def forward(self, input_ids, positions, forward_batch,
            input_embeds=None, pp_proxy_tensors=None):
    if input_embeds is None:
        input_embeds = self.embed_tokens(input_ids)
    # 将 target hidden_states 与当前输入 embedding 拼接后过 fc,
    # 得到草稿输入, 再走父类 DeepseekV2Model.forward
    input_embeds, _ = self.fc(
        torch.cat((input_embeds, forward_batch.spec_info.hidden_states), dim=-1)
    )
    output = super().forward(
        input_ids, positions, forward_batch, input_embeds, pp_proxy_tensors
    )
    assert isinstance(output, torch.Tensor)
    return output
```

test/registered/8-gpu-models/test_mistral_large3.py

夜间测试配套：用 `SGLANG_ENABLE_ASYNC_ASSERT.override(0)` 抑制 EAGLE 预热期 NaN 断言，与 Nemotron-3 做法一致，注释说明了抑制理由。

```
# test/registered/8-gpu-models/test_mistral_large3.py
# TP8+MTP 变体在 EAGLE 草稿 CUDA 图捕获阶段的 flashinfer-autotune 预热批次上
# 触发 `NaN detected! draft_forward step 0`; 该 NaN 在 2026-06-06 之前就已潜在
# （那时探针未启用，accept_len 2.45、gsm8k 0.960 均正常），#27461 默认开启
# SGLANG_ENABLE_ASYNC_ASSERT 后才变成启动期 abort。此处按 Nemotron-3 夜间测试
# 既有做法在测试期间关闭探针（suppression，非 root cause 修复），等待探针作者
# 对“预热全零批次是否应被断言”的结论。
with envs.SGLANG_ENABLE_ASYNC_ASSERT.override(0):
    run_combined_tests(
        models=variants,
        test_name="Mistral-Large-3",
        accuracy_params=AccuracyTestParams(dataset="gsm8k", baseline_accuracy=0.85),
        performance_params=PerformanceTestParams(
            result_dir="performance_results_mistral_large3",
        ),
    )
```

评论区精华

核心讨论集中在 PR body 的自我论证与提示（review 评论区无其他评论，Fridge003 直接 APPROVE）：

- 两个 bug 堆叠的证据链：作者用 nightly 历史证明 NaN 自 06-07 探针启用后开始报告，#28785 在 06-21 加入 `use_dsa` 读取后掩盖了它，本 PR 移除掩盖后失败回到与 06-20/21 完全一致的形态，说明 NaN 与 `use_dsa` 无关。
- 继承设计之争：“Attributes forward happens to read is not a real category”，作者放弃 helper 补齐方案，改为唯一初始化路径。

- 量化前缀修正的告警：add_prefix 参数交换导致权重名畸形 (layers.0.model / fc.model) ，修正后 quant 方法解析会变化，作者要求 second opinion。
- 对探针的开放质疑：warmup 全零批次上出现 NaN 是否是真实缺陷，还是探针断言位置不当，应交给 #27461 作者与 spec-decode 负责人裁决。
- 两个 bug 堆叠：NaN 先于 AttributeError 存在 (correctness): AttributeError 是新增崩溃；NaN 是更早存在的独立问题，与本 PR 无关。
- add_prefix 参数交换对量化解析的影响 (design): 作者与 reviewer (Fridge003 APPROVE) 接受修正方向，但量化 draft 路径需额外复验。
- warmup NaN 是缺陷还是探针位置不当 (question): 未解决，应交给 #27461 探针作者与 spec-decode 负责人裁决。
- 抑制探针 vs 修复根因 (design): 先接受抑制恢复 nightly 健康，后续按裁决移除。

风险与影响

- 风险：
 1. 回归风险：权重前缀从 layers.0.model / fc.model 变为 model.layers.0 / model.fc，依赖旧前缀的量化配置解析、离线权重脚本可能出现行为变化；现有覆盖只有 FP8 变体，量化 draft (NVFP4 等) 路径缺少回归测试。
 2. 测试抑制风险：关闭 SGLANG_ENABLE_ASYNC_ASSERT 会让 warmup NaN 不再被报告，如果该 NaN 是真实缺陷（例如稳态 decode 数值问题的前兆），夜间测试会失去报警；作者和 Nemotron-3 先例都依赖后续裁决。
 3. 继承契约风险：draft 完全依赖 DeepseekV2Model.__init__，父类新增状态会隐式影响子类，虽然本次消除了漂移，但未来对父类初始化的改动需要同步评估所有继承者。
 4. 性能风险：低。模块数不变，embed_tokens 在 nightly 配置 (enable_tp=True, use_attn_tp_group=False) 下值不变，仅在 DP attention 或 SGLANG_ENABLE_EMBED_REPLICATION 下行为有差异。- 影响：影响范围集中在 Mistral-Large-3-675B-Instruct-2512 的 EAGLE 草稿配置与对应 nightly-8-gpu-common 套件：修复后 B200 上 TP8+MTP 变体恢复健康（回到 06-06 的 accept_len 2.45 与 gsm8k 0.960 水平），消除 nightly 噪音。对团队而言，确立了 EAGLE draft 必须复用父类初始化的模式，为其他继承 DeepseekV2Model 的草稿模型提供了样板；同时把 CI 探针策略的开放问题摆上台面。影响程度中等偏低：源码改动小 (+6/-46)，但修复的是继承契约类深层 bug。- 风险标记：核心路径变更，量化前缀行为变更，测试探针被抑制，NaN 根因未解决

关联脉络

- PR #28785 （标题未在材料中提供）：该 PR 在 DeepseekV2Model.forward 首行引入 _dsa_forward_uses_topk()，读取 self.use_dsa，把 draft 缺失父类状态的 latent bug 变成显式崩溃；本 PR 的修复正是围绕它展开。
- PR #27461 Enable async-assert invariant probes by default in CI: 该 PR 在 nightly-test-nvidia.yml 与 rerun-test.yml 默认设置 SGLANG_ENABLE_ASYNC_ASSERT，暴露了 EAGLE 草稿捕获预热期的 NaN，本 PR 的测试抑制正是针对它。