

PR #29546 完整报告

sgl-project/sglang

Clean up follow-ups for eagle hidden dim clean up

合并时间: 2026-06-29 12:30

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29546>

执行摘要

- 一句话: 清理 #29464 遗留的陈旧注释和多余 model_config 返回值
- 推荐动作: 建议精读 kv_cache_builder.py 中 get_draft_kv_pool 的返回值简化, 了解如何安全地清理不再使用的契约; inspect 用法值得在其他类似场景推广, 以消除硬编码假设。

功能与动机

Address follow-up review comments from #29464: remove stale comment, use inspect for arg position, simplify init_disaggregation. 这些清理确保代码与当前逻辑保持一致, 消除误导性注释和硬编码。

实现拆解

1. 移除陈旧注释 (prefill_cuda_graph_runner.py): 删除了 __init__ 中关于 multimodal input_embeds slot 注册的多行注释, 该注释内容已过时, 容易误导开发者。
2. 用 inspect 替换硬编码位置参数 (prefill_cuda_graph_runner.py): 在 __init__ 中通过 inspect.signature(self.layer_model.forward).parameters 动态获取 input_embeds 参数索引, 存储为 self._input_embeds_arg_idx; 在 replay_layer_forward 中改用 ie_idx 替代硬编码的 args[3], 使代码更健壮, 不受 forward 签名变化影响。
3. 简化 init_disaggregation (scheduler.py 和 kv_cache_builder.py): get_draft_kv_pool 不再返回 model_config (已无人使用), 改为仅返回 draft_token_to_kv_pool; scheduler.py 中移除对应的 fallback 分支, 减少死代码。
4. 格式调整 (prefill_cuda_graph_runner.py): 执行了代码格式化 (Format prefill CUDA graph runner 提交)。

关键文件:

- python/sglang/srt/model_executor/runner/prefill_cuda_graph_runner.py (模块 调度器; 类别 source; 类型 core-logic; 符号 replay_layer_forward, init): 移除陈旧注释并用 inspect 动态解析 input_embeds 参数位置, 消除硬编码 args[3]。
- python/sglang/srt/managers/scheduler.py (模块 调度器; 类别 source; 类型 core-logic; 符号 init_disaggregation): 简化 init_disaggregation, 移除不再使用的 model_config 变量和 fallback 分支。
- python/sglang/srt/mem_cache/kv_cache_builder.py (模块 缓存层; 类别 source; 类型 data-contract; 符号 get_draft_kv_pool, maybe_register_hicache_draft):

get_draft_kv_pool 返回值从二元组改为单一值，清理内部引用。

关键符号：replay_layer_forward, init_disaggregation, get_draft_kv_pool, maybe_register_hicache_draft

关键源码片段

python/sglang/srt/model_executor/runner/prefill_cuda_graph_runner.py

移除陈旧注释并用 inspect 动态解析 input_embeds 参数位置，消除硬编码 args[3]。

```
# python/sglang/srt/model_executor/runner/prefill_cuda_graph_runner.py
# 以下为 __init__ 和 replay_layer_forward 中的关键变更
import inspect # 新增导入

class BreakableCudaGraphRunner(BaseCudaGraphRunner):
    def __init__(self, model_runner):
        # ...
        # 删除陈旧的多行注释后，buffer_registry 调用保持原样
        self.buffer_registry = build_prefill_registry(
            # ...
            source=self.buffers,
        )
        # ...
        # 解析 layer_model.forward 的参数列表，找到 input_embeds 的索引
        params = list(inspect.signature(self.layer_model.forward).parameters)
        self._input_embeds_arg_idx = (
            params.index("input_embeds") if "input_embeds" in params else None
        )

    def replay_layer_forward(self, *args, **layer_kwargs):
        # 通过关键字参数或动态索引获取 input_embeds，不再硬编码 args[3]
        if self.buffer_registry.has_slot("input_embeds"):
            ie = layer_kwargs.get("input_embeds")
            ie_idx = self._input_embeds_arg_idx
            if ie is None and ie_idx is not None and len(args) > ie_idx:
                ie = args[ie_idx]
            if ie is not None:
                self.buffer_registry.get_slot("input_embeds").slice_for(1, static_n)
```

python/sglang/srt/managers/scheduler.py

简化 init_disaggregation，移除不再使用的 model_config 变量和 fallback 分支。

```
# python/sglang/srt/managers/scheduler.py
def init_disaggregation(self):
    # ...
    # 简化调用：get_draft_kv_pool 不再返回 model_config
    draft_token_to_kv_pool = kv_cache_builder.get_draft_kv_pool(
        draft_worker=self.draft_worker,
        spec_algorithm=self.spec_algorithm,
```

```

        server_args=self.server_args,
    )
    # 删除旧的 fallback 分支: if model_config is None: model_config = self.model_config
    # 因为 model_config 已无人使用

```

python/sglang/srt/mem_cache/kv_cache_builder.py

get_draft_kv_pool 返回值从二元组改为单一值，清理内部引用。

```

# python/sglang/srt/mem_cache/kv_cache_builder.py

def get_draft_kv_pool(
    *,
    draft_worker: BaseTpWorker,
    spec_algorithm: SpeculativeAlgorithm,
    server_args: ServerArgs,
):
    """Return the draft token-to-KV pool for the current draft worker,
    or None when no draft KV pool is available."""
    if draft_worker is None or spec_algorithm.is_ngram():
        return None # 之前返回 (None, None)

    # V2 workers nest the draft runner under `draft_worker`.
    if server_args.enable_multi_layer_eagle:
        draft_runner = draft_worker.draft_worker.draft_runner_list[0]
    else:
        draft_runner = draft_worker.draft_worker.draft_runner
    return draft_runner.token_to_kv_pool # 之前返回 (pool, model_config)

```

评论区精华

无 review 评论。作者 merrymercy 自审后标记 approved 并合并。

- 自审批准 (other): 无需进一步讨论，直接合并。

风险与影响

- 风险：低风险。所有变更均为清理性更改：移除注释不影响逻辑；inspect 方式更安全，可自动适应 forward 签名变化；model_config 的移除确认了所有调用者已改用 get_draft_recurrent_hidden_state_spec。但需注意 inspect.signature 可能会带来极小的性能开销，仅在初始化时执行一次，可忽略。
- 影响：影响范围有限，仅涉及 3 个源码文件。对用户和系统无功能影响，主要是代码可维护性提升。开发者阅读 #29464 相关代码时会更清晰，减少误解。
- 风险标记：低风险，代码清理

关联脉络

- PR #29464 [Spec] Replace shared-infra dflash special-cases with capabilities (WAR barrier + seq_lens_cpu): 本 PR 是对 #29464 review 意见的后续清理，涉及相同文件中的

注释、参数硬编码和返回值简化。