

PR #31145 完整报告

sgl-project/sglang

Clean up ModelRunner by renaming effective-token property and remove dead code

合并时间: 2026-07-14 15:50

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/31145>

执行摘要

- 一句话: 清理 ModelRunner 中的死代码并重命名属性
- 推荐动作: 值得快速合并。作为大型重构链 (ModelRunner 拆分) 的一部分, 保持代码整洁。建议确认是否有外部依赖引用旧属性名。

功能与动机

根据 PR body, `MLA_ATTENTION_BACKENDS` 和 `add_mla_attention_backend` 在 `python/sglang/srt` 下没有任何读取方, 属于死代码。`RankZeroFilter` 未被引用。`alloc_memory_pool` 末尾的六个 `None` 默认值是旧 `disable_cuda_graph` 提前跳过路径的遗留物, 现在 `init_model_worker` 无条件调用 `init_attention_backends` 和 `init_cuda_graphs`, 这些默认值不再需要。属性重命名是为了更准确地表达语义 (考虑混合 SWA 设置)。

实现拆解

1. 移除 `MLA_ATTENTION_BACKENDS` 和 `add_mla_attention_backend`: 在 `model_runner.py` 中删除了整个列表定义和函数定义, 因为没有任何消费者。
2. 删除 `RankZeroFilter` 类: 移除了该日志过滤器类及其 `__init__` 和 `filter` 方法, 因为未被使用。
3. 删除 `alloc_memory_pool` 中的遗留默认值: 移除了六个字段的 `None` 赋值 (`attn_backend`、`decode_attn_backend`、`decode_attn_backend_group`、`decode_cuda_graph_runner`、`graph_mem_usage`、`prefill_cuda_graph_runner`), 因为这些现在由后续调用保证。
4. 重命名属性 `max_token_pool_size` 为 `effective_max_total_num_tokens`: 在 `model_runner.py` 中将 `property` 重命名, 更新注释以反映混合 SWA 场景。然后在 `tp_worker.py` (两处) 和 `disaggregation/prefill.py` (一处) 中更新所有引用。

关键文件:

- `python/sglang/srt/model_executor/model_runner.py` (模块 模型执行器; 类别 `source`; 类型 `data-contract`; 符号 `add_mla_attention_backend`, `RankZeroFilter`, `init`, `filter`): 核心变更文件: 删除 `MLA_ATTENTION_BACKENDS`、`add_mla_attention_backend`、`RankZeroFilter` 和 `alloc_memory_pool` 中的遗留默认值; 重命名属性 `effective_max_total_num_tokens`。
- `python/sglang/srt/managers/tp_worker.py` (模块 TP 工作器; 类别 `source`; 类型 `core-logic`): 两处引用 `max_token_pool_size` 更新为 `effective_max_total_num_tokens`。

- python/sglang/srt/disaggregation/prefill.py (模块 解耦预填充; 类别 source; 类型 core-logic) : 一处引用 max_token_pool_size 更新为 effective_max_total_num_tokens。

关键符号: add_mla_attention_backend, RankZeroFilter.init, RankZeroFilter.filter, effective_max_total_num_tokens, max_token_pool_size

关键源码片段

python/sglang/srt/model_executor/model_runner.py

核心变更文件: 删除 MLA_ATTENTION_BACKENDS、add_mla_attention_backend、RankZeroFilter 和 alloc_memory_pool 中的遗留默认值; 重命名属性 effective_max_total_num_tokens。

```
# python/sglang/srt/model_executor/model_runner.py

# 删除前:
# MLA_ATTENTION_BACKENDS = [
# "aiter", "flashinfer", "fa3", "fa4", "triton", "flashmla",
# "cutedsl_mla", "cutlass_mla", "trtlm_mla", "tokenspeed_mla",
# "ascend", "dsa", "nsa", "intel_xpu",
# ]
# def add_mla_attention_backend(backend_name):
# if backend_name not in MLA_ATTENTION_BACKENDS:
# MLA_ATTENTION_BACKENDS.append(backend_name)
# logger.info(f"Added {backend_name} to MLA_ATTENTION_BACKENDS.")

# 删除前:
# class RankZeroFilter(logging.Filter):
# def __init__(self, is_rank_zero):
# super().__init__()
# self.is_rank_zero = is_rank_zero
# def filter(self, record):
# if record.levelno == logging.INFO:
# return self.is_rank_zero
# return True

# 重命名后:
@property
def effective_max_total_num_tokens(self):
    """Return the max token pool size considering hybrid swa settings."""
    if self.is_hybrid_swa:
        return self.full_max_total_num_tokens or self.swa_max_total_num_tokens
    # ... 其余逻辑保持不变
```

python/sglang/srt/managers/tp_worker.py

两处引用 max_token_pool_size 更新为 effective_max_total_num_tokens。

```
# python/sglang/srt/managers/tp_worker.py

# 在 alloc_memory_pool 中 (约第 345 行) :
```

```
max_req_len = min(
    self.model_config.context_len - 1,
    self.model_runner.effective_max_total_num_tokens - 1, # 原为 max_token_pool_size
)

# 在 get_worker_info 中 (约第 457 行) :
max_req_len = min(
    self.model_config.context_len - 1,
    self.model_runner.effective_max_total_num_tokens - 1, # 原为 max_token_pool_size
)
```

评论区精华

该 PR 没有 review 评论。由于是机械性清理且讨论为零，无明显技术争议。

- 暂无高价值评论线程

风险与影响

- 风险：低风险。变更主要是删除死代码和属性重命名，不涉及逻辑修改。属性重命名可能影响外部插件或未发现的内部引用，但搜索显示 `srt` 目录内所有使用均已更新。遗留默认值删除安全，因为调用链保证这些字段会被赋值。
- 影响：正面：减少代码量（-47 行），提高可读性。无功能影响。对用户透明；对开发者，需注意未来引用新属性名。
- 风险标记：属性重命名可能导致外部引用断裂

关联脉络

- PR #31169 Split `initialize()` into orchestration helpers: 同一系列 `ModelRunner` 拆分与清理工作的一部分。
- PR #31166 Narrow component dependencies to injected fields instead of `ModelRunner`: 后续清理依赖注入模式。