

PR #44558 完整报告

vllm-project/vllm

[Core] Add prefill step cadence for better non-PD DP balancing

合并时间: 2026-06-17 04:17

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/44558>

执行摘要

- 一句话: 新增 prefill step cadence 大幅提升 DP 场景 decode 延迟
- 推荐动作: 值得精读, 特别是 `schedule()` 中 `defer` 逻辑的实现和饱和保护机制。设计决策清晰, 性能收益显著且经过充分基准测试。适合希望深入理解 v1 调度器及 DP 平衡策略的工程师。

功能与动机

在非 PD (Pipeline Parallel) 的 DP 部署中, prefill 计算会与 decode 步骤竞争 GPU 资源, 导致 decode 延迟 (ITL) 大幅抖动。PR body 指出“interval=2 was optimal, but that a guard is needed to disable the throttling when saturated, to avoid overall perf degradation”, 目标是保持 DP 各 rank 上 prefill 的同步性, 同时利用 decode-only batch 极快的 CUDA Graph 执行速度。

实现拆解

1. 配置入口 (vllm/config/scheduler.py, vllm/engine/arg_utils.py): 在 SchedulerConfig 和 EngineArgs 中新增 prefill_schedule_interval 参数, 默认值在 SchedulerConfig 中定义 (未暴露, 但代码中从 DPEngineCoreProc 的行为推断为 1 即禁用)。CLI 添加 --prefill-schedule-interval 参数。
2. 接口扩展 (vllm/v1/core/sched/interface.py): 将 SchedulerInterface.schedule() 抽象方法增加可选参数 throttle_prefills: bool = False, 并在 docstring 中说明其语义。
3. 调度器核心逻辑 (vllm/v1/core/sched/scheduler.py): 在 Scheduler 的 __init__ 中添加 self.prefill_capacity_bound 标志位, 初始为 False。schedule() 方法首先检查 throttle_prefills 和 prefill_capacity_bound 计算 defer_prefills 布尔值, 条件为: throttle_prefills and not self.prefill_capacity_bound 且 any(not r.is_prefill_chunk for r in self.running) (即至少有一个正在 decode 的请求值得保护)。然后在遍历 RUNNING 和 WAITING 请求时, 若 defer_prefills 为 True, 则跳过新 prefill (num_computed_tokens == 0) 和正在进行的 prefill chunk (request.is_prefill_chunk), 但不影响 decode 请求和远程 KV 恢复 (num_external_computed_tokens > 0)。每次调度结束后, 若无 prefill 被延迟 (即当前步允许 prefill), 则记录 self.prefill_capacity_bound = bool(self.waiting), 表示此步是否仍有余量 (未排空等待队列), 供下一轮判断。

4. 引擎层控制 (vllm/v1/engine/core.py) : EngineCore 新增 `_should_throttle_prefills()` 方法, 默认返回 `False`。 `step()` 和 `step_with_batch_queue()` 中传入该值给 `scheduler.schedule()`。 DP 专用的 `DPEngineCoreProc` 重写 `_should_throttle_prefills()`: 当 `prefill_schedule_interval > 1` 且 `step_counter % prefill_schedule_interval != 0` 时返回 `True`, 确保同一 DP group 内 `step_counter` 同步, `prefill` 只在 `cadence` 步触发。
5. 测试配套: 单元测试 (`tests/v1/core/test_scheduler.py`) 覆盖 `gating` 行为、远程 KV 排除、inflight chunk 推迟、容量饱和和保护。集成测试 (`tests/v1/distributed/test_async_llm_dp.py`) 使用 MoE 模型在 DP 模式下验证 `prefill_schedule_interval` 的端到端正确性。

关键文件:

- `vllm/v1/core/sched/scheduler.py` (模块 调度器; 类别 `source`; 类型 `core-logic`; 符号 `schedule`) : 调度核心实现, 新增 `prefill` 节拍控制逻辑及饱和和保护标志。
- `vllm/v1/engine/core.py` (模块 引擎; 类别 `source`; 类型 `core-logic`; 符号 `_should_throttle_prefills`) : 引擎层控制 `_should_throttle_prefills`, DP 模式下重写为基于 `step_counter` 的节拍判断。
- `tests/v1/core/test_scheduler.py` (模块 测试; 类别 `test`; 类型 `test-coverage`; 符号 `test_schedule_prefills_gating`, `test_throttle_prefills_excludes_remote_kv_resume`, `test_throttle_defers_inflight_prefill_chunk`, `test_throttle_capacity_bound_guard_admits`) : 新增 4 个单元测试, 全面覆盖 DP `prefill` 节拍逻辑的各个分支。

关键符号: `schedule`, `_should_throttle_prefills`, `test_schedule_prefills_gating`, `test_throttle_prefills_excludes_remote_kv_resume`, `test_throttle_defers_inflight_prefill_chunk`, `test_throttle_capacity_bound_guard_admits`, `test_dp_prefill_schedule_interval`

关键源码片段

`vllm/v1/core/sched/scheduler.py`

调度核心实现, 新增 `prefill` 节拍控制逻辑及饱和和保护标志。

```
def schedule(self, throttle_prefills: bool = False) -> SchedulerOutput:
    self.current_step += 1
    # ... 原有逻辑 ...
    # 容量饱和和保护: 当上次 cadence 步仍未排空等待队列时,
    # 禁用节流, 防止饱和时吞吐下降。
    defer_prefills = (
        throttle_prefills and not self.prefill_capacity_bound
    ) and any(not r.is_prefill_chunk for r in self.running)

    # 在遍历 RUNNING 请求时, 如果是 prefill chunk 且需要延迟, 则跳过
    # 但 decode 请求正常调度。
    if defer_prefills and request.is_prefill_chunk:
        req_index += 1
        continue

    # 在遍历 WAITING 请求时, 跳过 num_computed_tokens == 0 的新 prefill
```

```

elif defer_prefills and request.num_computed_tokens == 0:
    break

# ... 原有调度逻辑 ...

# 调度结束后，如果本步未延迟 prefill（即允许 prefill 的步），
# 记录是否饱和（等待队列非空）。
if not defer_prefills:
    self.prefill_capacity_bound = bool(self.waiting)

```

vllm/v1/engine/core.py

引擎层控制 `_should_throttle_prefills`，DP 模式下重写为基于 `step_counter` 的节拍判断。

```

# 在 EngineCore 中（默认，非 DP 原样返回 False）：
def _should_throttle_prefills(self) -> bool:
    """是否延迟新 prefill（仅 DP 场景覆盖）。"""
    return False

# 在 DP EngineCoreProc 中（DP 场景）：
def __init__(self, *args, **kwargs):
    super().__init__(*args, **kwargs)
    scheduler_config = vllm_config.scheduler_config
    self.prefill_schedule_interval = scheduler_config.prefill_schedule_interval
    # step_counter 在所有 DP rank 上同步

def _should_throttle_prefills(self) -> bool:
    # 仅在 cadence 步（step_counter % interval == 0）时才允许 prefill
    return (
        self.prefill_schedule_interval > 1
        and self.step_counter % self.prefill_schedule_interval != 0
    )

```

评论区精华

- API 设计决策：WoosukKwon 建议将 `throttle_prefills` 作为 `schedule()` 的参数而非独立的 setter 方法（`set_throttle_prefills`）。njhill 采纳并提交了修改，但表示略有保留——认为 `schedule` 方法原本无参数，这种设计更分离。最终该方案被合并，WoosukKwon LGTM。
- API 设计：`throttle_prefills` 作为参数还是 setter (design)：采用参数方案，保持 `schedule()` 接口显式。

风险与影响

- 风险：
 1. 非 DP 场景无影响：默认 `_should_throttle_prefills()` 返回 `False`，调度行为不变。
 2. 饱和保护脆弱性：`prefill_capacity_bound` 标志位仅在非 `throttle` 步更新，如果等待队列在单步内被排空但仍有新请求到达，会导致 `guard` 滞后。但 PR body 中给出了实测数

据，饱和状态下不会出现性能回退。

3. prefill-heavy 负载退化：body 中明确指出了当输入极长（16k vs 200 decode）时，throttle 会降低 prefill 吞吐、增加 TTFT，此时应关闭此功能。

4 API 兼容性：SchedulerInterface.schedule() 增加可选参数，所有子类（包括外部实现）需要适配，否则会因签名不匹配报错。- 影响：面向用户：使用 DP 部署且 prefill 计算较重的用户，通过 --prefill-schedule-interval 可大幅降低 decode 延迟（实测 -88% P50 ITL），TTFT 同步改善。非 DP 用户无影响。prefill-heavy 负载用户应避免启用。系统资源：throttle 步骤中 GPU 执行 decode-only batch 时因 CUDA Graph 效率更高，整体 ITL 改善。团队协作：该设计为后续更通用的 weighted max-batch-tokens 方案（见 Issue 评论）提供了基础，调度器接口预留了 throttle 参数。- 风险标记：核心调度路径变更，饱和度保护依赖状态，prefill-heavy 负载需注意，接口扩展需兼容子类

关联脉络

- 暂无明显关联 PR