

PR #51865 完整报告

vllm-project/vllm

[Bugfix][MRV2] Require all requests to be decoding for uniform-decode dispatch

合并时间: 2026-08-12 08:15

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51865>

执行摘要

- 一句话: 统一解码批须校验请求状态, 修复投机解码全图误重放
- 推荐动作: 值得精读。核心看点有三: 一是 `prepare_inputs` 拆分出 `gather_batch_req_state` 并前置到 DP 同步之前的设计——把 CPU 侧轻量预判与 GPU 输入装配解耦, 既服务于跨 DP 图模式选择又避免了重复计算; 二是“形状 + 状态”双条件分类的语义分层, `get_uniform_token_count` 与 `get_uniform_decode_token_count` 的边界通过 docstring 固化; 三是 AST 扫描回归测试, 针对 speculator 复制粘贴文化做的防退化设计, 可推广到其他同类高风险调用点。

功能与动机

PR body 说明这是 #50532 的复制并追加了避免冗余计算的返工。commit message 指出根因: Model runner v2 仅凭形状 (`num_tokens == num_reqs * max_query_len` 的 `get_uniform_token_count`) 判定 uniform decode 批, 而完全不看请求是否真的在 decoding。恰好 `1 + num_speculative_tokens` 个 token 的 prompt chunk 具有相同形状, 于是 prefill 被以 `cg_mode=FULL` 调度并重放了捕获的 spec-verify decode 图到 prompt tokens 上, 产生静默数据损坏; 无投机解码时 1-token prompt 与 1-token decode 也存在同样的形状碰撞 (测试注释明确关联 issue #49918)。

实现拆解

实现分四步推进:

1. 判定逻辑分层 (`vllm/v1/worker/utils.py` + `cudagraph_utils.py`): 新增 `is_uniform_query_len` (纯形状判定: `num_reqs > 0 and num_tokens == max_query_len * num_reqs`) 与 `get_uniform_decode_token_count` (在形状基础上额外要求 `not has_prefill`)。 `cudagraph_utils.get_uniform_token_count` 重写为委托 `is_uniform_query_len`, 并在 docstring 中明确其为 shape-only 判定, 仅适用于构造上就是 decode 的批 (如 dummy run), 调度批必须使用 `get_uniform_decode_token_count`。
2. `gather_batch_req_state` 拆分 (`vllm/v1/worker/gpu/model_runner.py`): 把 `prepare_inputs` 前半段的 CPU 状态收集 (`sort_batch_req_ids` 排序、`idx_mapping` 构建、prefill 状态计算) 抽成独立方法 `gather_batch_req_state(scheduler_output, dummy_run)`, 返回 `BatchReqState` (NamedTuple, 含 `req_ids`、`num_scheduled_tokens`、`idx_mapping_np`、`prefill_len_np`、`num_computed_prefill_tokens_np`、`is_prefilling_np`、`has_prefill`) 与 `uniform_decode_token_count`。 `prepare_inputs` 改为直接消费

`batch_req_state`, 不再重复 `gather` 与索引映射; `execute_model` 在 `dp token` 数同步 / `cuda-graph dispatch` 之前调用该方法, 使分类结果参与跨 DP 的图模式选择, 并新增 `assert batch_req_state is not None` 保护真实路径。

- dummy run 特例与 `speculator` 切换: `dummy` 批没有请求状态可查, `gather_batch_req_state` 走 `dummy_run` 分支直接返回 `get_uniform_token_count` 的形状分类, 保证 `capture` 与 `replay` 判定一致。`AutoRegressiveSpeculator.propose` 与 `multi-module MTP` 的 `propose` 均把 `get_uniform_token_count` 换成 `get_uniform_decode_token_count` 并传入 `input_batch.has_prefill`。为此 `InputBatch` 新增 `has_prefill` 字段, `make_dummy` 补 `has_prefill=False`, `pcp_manager.partition_batch` 在分区后同步补 `has_prefill=bool(local_is_prefilling_np.any())`。
- 测试配套: `tests/v1/worker/test_gpu_batch_ordering.py` 用 `GPUModelRunner.__new__` 构造轻量 runner stub, 围绕 `gather_batch_req_state` 覆盖五种场景 (纯 decode 批、prompt chunk 撞 decode 形状、dummy 批、state 索引与 batch 位置错位、纯谓词), 另用 AST 扫描 `vllm/v1/worker/gpu/spec_decode/` 下所有 `speculator.py`, 断言没有任何调用点退化为纯形状分类。`tests/v1/spec_decode/test_dynamic_sd_cug.py` 新增 `test_prompt_chunks_shaped_like_spec_decode_miss_the_full_graph`, 验证 `K+1 token` 的 prompt 块在 `CudaGraphManager.dispatch` 层面命中 `PIECEWISE` 而非 `FULL` 图。

关键文件:

- `vllm/v1/worker/gpu/model_runner.py` (模块 模型执行; 类别 source; 类型 data-contract; 符号 `prepare_inputs`, `gather_batch_req_state`, `BatchReqState`): 核心变更文件: 将 `prepare_inputs` 的 CPU 状态收集拆为 `gather_batch_req_state`, 新增 `BatchReqState` 数据结构, `execute_model` 在 DP 同步前调用新方法完成 `uniform-decode` 判定。
- `tests/v1/worker/test_gpu_batch_ordering.py` (模块 批排序; 类别 test; 类型 test-coverage; 符号 `_make_runner`, `_uniform_token_count`, `test_spec_decode_batch_is_uniform_decode`, `test_prompt_chunk_of_decode_query_len_is_not_uniform_decode`): 最主要的测试文件: 围绕 `gather_batch_req_state` 新增五个场景测试, 并用 AST 扫描 `spec_decode` 包防止纯形状分类调用点回归。
- `vllm/v1/worker/utils.py` (模块 批分类; 类别 source; 类型 core-logic; 符号 `is_uniform_query_len`, `get_uniform_decode_token_count`): 新增 `is_uniform_query_len` 与 `get_uniform_decode_token_count`, 是整个修复的判定逻辑核心。
- `vllm/v1/worker/gpu/cudagraph_utils.py` (模块 图管理; 类别 source; 类型 dependency-wiring; 符号 `get_uniform_token_count`): `get_uniform_token_count` 重写为委托 `is_uniform_query_len`, 并在 `docstring` 中明确其为 `shape-only` 判定, 仅适用于 `dummy runs`。
- `tests/v1/spec_decode/test_dynamic_sd_cug.py` (模块 图调度; 类别 test; 类型 test-coverage; 符号 `test_prompt_chunks_shaped_like_spec_decode_miss_the_full_graph`): 新增 `dispatch` 层端到端测试, 验证 `K+1 token` 的 prompt 块命中 `PIECEWISE` 而非 `FULL` 图, 直接守护 `CudaGraphManager` 的调度行为。

- `vllm/v1/worker/gpu/spec_decode/autoregressive/speculator.py` (模块 投机解码; 类别 `source`; 类型 `dependency-wiring`; 符号 `propose`) : `AutoRegressiveSpeculator.propose` 切换为 `get_uniform_decode_token_count`, 并传入 `input_batch.has_prefill`。
- `vllm/v1/worker/gpu/spec_decode/multi_module_mtp/speculator.py` (模块 多模投机; 类别 `source`; 类型 `dependency-wiring`; 符号 `propose`) : multi-module MTP speculator 同样切换判定 API, 该文件是 `autoregressive` 的复制衍生物, 一并对齐修复。
- `vllm/v1/worker/gpu/input_batch.py` (模块 输入批; 类别 `source`; 类型 `data-contract`; 符号 `InputBatch, make_dummy`) : `InputBatch` 新增 `has_prefill` 字段并贯通 `make_dummy`, 是跨模块传递判定状态的数据契约变更。
- `vllm/v1/worker/gpu/pcp_manager.py` (模块 图分区; 类别 `source`; 类型 `core-logic`; 符号 `partition_batch`) : PCP 分区后同步补齐 `has_prefill` 字段, 保证分区子批的判定状态不丢失。

关键符号: `gather_batch_req_state`, `get_uniform_decode_token_count`, `is_uniform_query_len`, `prepare_inputs`, `get_uniform_token_count`, `test_no_speculator_dispatches_on_query_length_alone`, `test_prompt_chunks_shaped_like_spec_decode_miss_the_full_graph`

关键源码片段

`tests/v1/worker/test_gpu_batch_ordering.py`

最主要的测试文件: 围绕 `gather_batch_req_state` 新增五个场景测试, 并用 AST 扫描 `spec_decode` 包防止纯形状分类调用点回归。

```
def test_no_speculator_dispatches_on_query_length_alone():
    """任何 speculator 都不得仅凭形状分类 decode 批。

    之所以扫描整个包: 新 speculator 常通过复制现有实现而来,
    历史上就出现过修复后的调用点被“原样复制回去”的回归。
    """
    import vllm.v1.worker.gpu.spec_decode as spec_decode

    # 这两个是纯形状 API, 出现在任何 speculator 调用点都视为违规。
    shape_only = {"get_uniform_token_count", "is_uniform_query_len"}
    root = Path(spec_decode.__file__).parent
    offenders = []
    for path in sorted(root.rglob("speculator.py")):
        # 用 AST 而非文本匹配, 避免注释或字符串误报。
        for node in ast.walk(ast.parse(path.read_text())):
            if not isinstance(node, ast.Call):
                continue
            func = node.func
            name = func.id if isinstance(func, ast.Name) else getattr(func, "attr", "")
            if name in shape_only:
                offenders.append(f"{path.relative_to(root)}:{node.lineno} {name}")
```

```
assert not offenders, (
    "these speculators classify a decode batch by query length alone; use "
    f"get_uniform_decode_token_count instead: {offenders}"
)
```

评论区精华

唯一实质 review 讨论来自 LucasWilkinson 在 vllm/v1/worker/utils.py:623 的 nit:

```
"i find having get_uniform_token_count and get_uniform_decode_token_count
confusing, why can't we deprecate get_uniform_token_count and use
get_uniform_decode_token_count exclusively?"
```

从最终实现看，两个函数被刻意保留并做了语义分层：`get_uniform_token_count` 是 shape-only 判定，仍被 dummy run 路径使用（dummy 批没有请求状态可查，只能按形状分类，且必须与 capture 阶段一致）；`get_uniform_decode_token_count` 面向真实调度批，需要 `has_prefill` 状态。PR 以 docstring 明确区分二者适用场景，并未废弃前者。WoosukKwon 最终 APPROVED ("LGTM. Thanks!")，claude[bot] 因 fork 未自动审查。

- `get_uniform_token_count` 与 `get_uniform_decode_token_count` 命名易混淆，能否废弃前者 (design): 最终保留两个函数并做语义分层：`get_uniform_token_count` 是 shape-only 判定，dummy run 等构造上已知为 decode 的批仍需要它；`get_uniform_decode_token_count` 面向真实调度批。以 docstring 明确区分适用场景。

风险与影响

- 风险：
 1. 核心执行路径变更：`prepare_inputs` 签名变化（新增 `batch_req_state` 参数）触及 MRV2 每步执行的 hot path，`execute_model` 中新增断言 `batch_req_state is not None`，若未来调用顺序调整（如未先 `gather` 就 `prepare`）会直接断言失败，属于显式契约保护但需注意后续演进。
 2. 数据结构契约扩散：`InputBatch` 新增 `has_prefill` 字段，所有构造点（`make_dummy`、`pcp_manager.partition_batch`）已同步，但未来新增 `InputBatch` 构造路径时可能遗漏该字段。
 3. 行为变化带来的性能特征变化：此前“碰巧”命中 FULL 图的 `prefill` 批现在正确退回 `PIECEWISE`，结果正确但延迟特征可能与旧版本不同；混合批由于 `gather_batch_req_state` 内 $O(1)$ 形状预检先行，不会构建被立即丢弃的 `gather` 结构，性能开销可控。
 4. 回归风险：`speculator` 常通过复制现有实现扩展，本次专门用 AST 扫描测试防止 `get_uniform_token_count` 调用点被复制回去，属于针对性的防回归设计。
 - 影响：影响所有启用投机解码（MTP、EAGLE 等 auto-regressive 系）且存在 chunked prefill 与 decode 混合调度的 MRV2 部署，以及无投机解码下 1-token prompt chunk 与 decode 形状碰撞的边界场景。修复前这些场景会静默产出损坏结果，修复后正确回退 `PIECEWISE`。对纯 decode 批与纯 prefill 批无行为变化；对团队而言，新增了两个公共判定函数并明确了 shape-only 与 decode 判定的语义边界，后续 `speculator` 实现必须遵循新 API。
 - 风险标记：核心执行路径变更，数据契约扩展，修复前为静默数据损坏，

关联脉络

- PR #50532 [Bugfix][MRV2] Require all requests to be decoding for uniform-decode dispatch (原版): PR body 明确说明本 PR 是 50532 的复制, 并追加了避免冗余计算的返工 (gather_batch_req_state 拆分) 。
- PR #50020 [Bugfix][MRV2] Support encoder timing stats in model runner V2: 同为 MRV2 model_runner 路径的 bugfix, 说明 MRV2 重构期模型执行路径处于高频改动状态。
- PR #51854 [CI][Bugfix][V1] Remove stale FlashAttention metadata arguments: 同为 V1 worker 测试路径的 CI 回归修复, 与本 PR 共同反映 V1 执行路径近期活跃演进。