

PR #52329 完整报告

vllm-project/vllm

[Performance][MRV2] Cache logits-processing request state

合并时间: 2026-08-17 10:36

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52329>

执行摘要

- 一句话: 缓存 logits 处理判定, MRV2 采样门控提速约 70%
- 推荐动作: 值得精读。这是一次典型的“热路径判定预计算”优化: 把每步重复扫描 7 个数组的工作折叠为 `add_request()` 一次的 $O(1)$ 写入, 且测试通过参数化把组合谓词的每个分支都钉死, 并专门覆盖槽位复用与活跃子集过滤两个最容易踩坑的点。关注点: 未来扩展采样状态时如何保证缓存同步, 可考虑把谓词收敛到单一派生结构或在测试中增加“缓存与实际状态一致性”的断言。

功能与动机

PR body 明确说明性能动机: `Sampler._requires_logits_processing()` 在每一步采样都会运行, 且“currently scans multiple per-request NumPy/Python states”, 即每次采样都要按活跃请求逐行扫描 7 个状态数组, 构成 CPU 热路径开销。同时该 PR 依赖 #52284, 因为“an active thinking budget must participate in the combined predicate”——若思考预算不参与组合谓词, `temperature=0.0` 的 greedy 请求会绕过 thinking budget 的强制结束符注入 (第一个提交本身就是这个 bugfix 的 rebase)。

实现拆解

本 PR 将 `Sampler` 的 logits 处理判定从“每个采样步骤实时扫描”改为“请求加入时预计算缓存”, 分 5 步落地:

1. 新增缓存数组: 在 `Sampler.__init__` (`vllm/v1/worker/gpu/sample/sampler.py`) 中新增 `self.needs_logits_processing = np.zeros(max_num_reqs, dtype=bool)`, 与 `SamplingStates`、`PenaltiesState` 等子状态并行存放, 每个请求槽位一个 boolean 值。
2. 组合谓词前移: 在 `add_request()` 登记完所有子状态后, 逐个读取 `logit_bias_state.use_logit_bias`、`penalties_state.use_penalty`、`bad_words_state.num_bad_words`、`thinking_budget_state.enabled / use_thinking_budget`、`sampling_states` 的 `temperature / min_p / top_k / top_p`, 用 OR 组合成单一判定写入 `needs_logits_processing[req_idx]`。槽位被新请求复用时, 新值直接覆盖旧值, 避免残留上一次请求的判定。
3. 热路径简化: `apply_sampling_params()` 的门控从 `_requires_logits_processing(idx_mapping_np)` (每次按活跃行扫描 7 个 NumPy/Python 数组) 改为 `np.any(self.needs_logits_processing[idx_mapping_np])` (只扫描 1 个紧凑 boolean 数组)

的活跃行)。该判定通过 `input_batch.idx_mapping_np` 天然只覆盖本步活跃请求，保留原 `active-subset` 过滤语义。

4. 行为修正与死代码清理：组合谓词首次纳入思考预算（依赖 #52284），修复 `temperature=0.0` 的 greedy 请求绕过 thinking budget 强制结束符注入的缺陷；按 review 反馈移除只剩单次调用的 `_requires_logits_processing()` 方法，检查直接内联在 `apply_sampling_params()` 中。
5. 测试配套：新增 `tests/v1/worker/test_gpu_sampler_flags.py` (90 行)，用 11 组参数化用例把 default、greedy、thinking-budget、logit-bias、penalty、bad-words、temperature、min-p、top-k、top-p、warmup 与缓存预期一一钉死，并覆盖槽位复用覆盖与活跃子集过滤两个边界；`tests/v1/worker/test_gpu_thinking_budget.py` 新增 `test_v2_greedy_sampling_applies_thinking_budget` 验证 greedy + budget 场景下 END 标记仍被置为 `1e9`。作者本地运行两个测试文件共 26 个用例通过，并通过 ruff、mypy 与 Qwen3.5-9B 固定种子 A/B 矩阵、GSM8K 评估。

关键文件：

- `vllm/v1/worker/gpu/sample/sampler.py` (模块 采样器；类别 source；类型 core-logic；符号 `Sampler`, `add_request`, `apply_sampling_params`, `_requires_logits_processing`)：核心变更文件：新增 per-request boolean 缓存数组，在 `add_request()` 组合所有 logits 处理判定，热路径门控从扫描 7 个状态数组改为单数组 `np.any()`，并移除 `_requires_logits_processing()`。
- `tests/v1/worker/test_gpu_sampler_flags.py` (模块 采样测试；类别 test；类型 test-coverage；符号 `MockReasoningConfig`, `_make_sampler`, `test_logits_processing_cache_matches_request_features`, `test_logits_processing_cache_is_overwritten_when_slot_is_reused`)：新增测试文件：用 11 组参数化用例把组合谓词的每个分支与缓存预期一一对应，并专门覆盖槽位复用覆盖与活跃子集过滤两个边界场景。
- `tests/v1/worker/test_gpu_thinking_budget.py` (模块 思考预算；类别 test；类型 test-coverage；符号 `test_v2_greedy_sampling_applies_thinking_budget`)：新增 greedy + thinking budget 组合回归测试，验证思考预算参与组合谓词后，greedy 请求不会绕过强制结束符注入。

关键符号：`Sampler.add_request`, `Sampler.apply_sampling_params`, `_requires_logits_processing`, `test_logits_processing_cache_matches_request_features`, `test_logits_processing_cache_is_overwritten_when_slot_is_reused`, `test_logits_processing_cache_only_checks_active_requests`, `test_v2_greedy_sampling_applies_thinking_budget`

关键源码片段

`vllm/v1/worker/gpu/sample/sampler.py`

核心变更文件：新增 per-request boolean 缓存数组，在 `add_request()` 组合所有 logits 处理判定，热路径门控从扫描 7 个状态数组改为单数组 `np.any()`，并移除 `_requires_logits_processing()`。

```
# vllm/v1/worker/gpu/sample/sampler.py
```

```
class Sampler:
```

```
    def __init__(
```

```
        self,
```

```
        max_num_reqs: int,
```

```
        vocab_size: int,
```

```
        device: torch.device,
```

```
        req_states: RequestState,
```

```
        reasoning_config: ReasoningConfig | None = None,
```

```
        ...
```

```
    ):
```

```
        ...
```

```
        self.thinking_budget_state = ThinkingBudgetState(req_states, reasoning_config)
```

```
        # 每个请求槽位一个 boolean 值，预计算 " 该请求是否需要 logits 处理 "。
```

```
        # 热路径只需按 idx_mapping_np 取活跃行做一次 np.any() 判定，
```

```
        # 避免每次采样都扫描多个 NumPy/Python 状态数组。
```

```
        self.needs_logits_processing = np.zeros(max_num_reqs, dtype=bool)
```

```
        ...
```

```
    def add_request(
```

```
        self, req_idx: int, prompt_len: int, sampling_params: SamplingParams
```

```
    ) -> None:
```

```
        # 先登记各子状态（内部可能暂存 staged writes）。
```

```
        self.sampling_states.add_request(req_idx, sampling_params)
```

```
        self.penalties_state.add_request(req_idx, sampling_params)
```

```
        self.logit_bias_state.add_request(req_idx, prompt_len, sampling_params)
```

```
        self.bad_words_state.add_request(req_idx, sampling_params)
```

```
        self.logprob_token_ids_state.add_request(req_idx, sampling_params)
```

```
        self.thinking_budget_state.add_request(req_idx, sampling_params)
```

```
        # 组合谓词前移：各子状态的开关在请求加入时即确定，
```

```
        # 逐项读取后写入槽位；槽位被新请求复用时会覆盖，不会残留旧值。
```

```
        states = self.sampling_states
```

```
        temperature = states.temperature_np[req_idx]
```

```
        self.needs_logits_processing[req_idx] = (
```

```
            self.logit_bias_state.use_logit_bias[req_idx]
```

```
            or self.penalties_state.use_penalty[req_idx]
```

```
            or self.bad_words_state.num_bad_words_np[req_idx] > 0
```

```
            # 思考预算必须参与谓词：否则 temperature=0.0 的 greedy 请求
```

```
            # 会绕过 thinking budget 的强制结束符注入（配合 #52284 修复）。
```

```
            or (
```

```
                self.thinking_budget_state.enabled
```

```
                and self.thinking_budget_state.use_thinking_budget[req_idx]
```

```
            )
```

```
            or (temperature != 0.0 and temperature != 1.0)
```

```
            or states.min_p_np[req_idx] != 0.0
```

```
            or states.top_k_np[req_idx] != states.vocab_size
```

```
            or states.top_p_np[req_idx] != 1.0
```

```
        )
```

```

def apply_sampling_params(
    self,
    logits: torch.Tensor,
    ...
    skip_top_k_top_p: bool = False,
) -> torch.Tensor:
    # 热路径门控：只扫描 1 个紧凑 boolean 数组的活跃行。
    if not np.any(self.needs_logits_processing[idx_mapping_np]):
        return logits
    # 以下仍按需执行：logit bias、penalties、bad words、thinking budget、
    # temperature、min-p、top-k/top-p（逐项 in-place 或返回新张量）。
    ...

```

tests/v1/worker/test_gpu_sampler_flags.py

新增测试文件：用 11 组参数化用例把组合谓词的每个分支与缓存预期一一对应，并专门覆盖槽位复用覆盖与活跃子集过滤两个边界场景。

```

# tests/v1/worker/test_gpu_sampler_flags.py
def _make_sampler() -> Sampler:
    req_states = RequestState(
        max_num_reqs=4,
        max_model_len=64,
        max_num_batched_tokens=16,
        num_speculative_steps=1,
        vocab_size=VOCAB_SIZE,
        device=DEVICE,
    )
    return Sampler(
        max_num_reqs=4,
        vocab_size=VOCAB_SIZE,
        device=DEVICE,
        req_states=req_states,
        reasoning_config=MockReasoningConfig(),
    )

@pytest.mark.parametrize(
    ("sampling_params", "expected"),
    [
        pytest.param(SamplingParams(), False, id="defaults"),
        pytest.param(SamplingParams(temperature=0.0), False, id="greedy"),
        pytest.param(
            SamplingParams(thinking_token_budget=3), True, id="thinking-budget"
        ),
        pytest.param(SamplingParams(logit_bias={1: 1.0}), True, id="logit-bias"),
        pytest.param(SamplingParams(frequency_penalty=0.1), True, id="penalty"),
        pytest.param(
            SamplingParams(_bad_words_token_ids=[[1]]), True, id="bad-words"
        ),
    ],

```

```

        pytest.param(SamplingParams(temperature=0.7), True, id="temperature"),
        pytest.param(SamplingParams(min_p=0.1), True, id="min-p"),
        pytest.param(SamplingParams(top_k=10), True, id="top-k"),
        pytest.param(SamplingParams(top_p=0.9), True, id="top-p"),
        pytest.param(
            SamplingParams.for_sampler_warmup(), True, id="all-logits-processors"
        ),
    ],
)

def test_logits_processing_cache_matches_request_features(
    sampling_params: SamplingParams, expected: bool
):
    # 参数化钉死组合谓词的每一个分支，防止缓存与真实状态脱节。
    sampler = _make_sampler()
    sampler.add_request(3, prompt_len=1, sampling_params=sampling_params)
    assert sampler.needs_logits_processing[3] == expected

def test_logits_processing_cache_is_overwritten_when_slot_is_reused():
    # 同一槽位先放 warmup（需要处理）再放默认参数（不需要处理），
    # 验证缓存被覆盖，而不是残留上一次请求的旧值。
    sampler = _make_sampler()
    sampler.add_request(3, 1, SamplingParams.for_sampler_warmup())
    sampler.add_request(3, 1, SamplingParams())
    assert not sampler.needs_logits_processing[3]

def test_logits_processing_cache_only_checks_active_requests():
    # 热路径按 idx_mapping_np 只取活跃行，非活跃请求的脏缓存不影响判定。
    sampler = _make_sampler()
    sampler.add_request(0, 1, SamplingParams(temperature=0.0))
    sampler.add_request(2, 1, SamplingParams.for_sampler_warmup())

    sampling_only = np.array([0], dtype=np.int32)
    with_processing = np.array([0, 2], dtype=np.int32)

    assert not np.any(sampler.needs_logits_processing[sampling_only])
    assert np.any(sampler.needs_logits_processing[with_processing])

```

评论区精华

njhill 提出两点 review 意见：(1) `_requires_logits_processing()` 在移除多数组扫描后只剩单次调用，建议直接内联为 `if not np.any(self.needs_logits_processing[idx_mapping_np])`；(2) 字段名 `_needs_logits_processing` 带下划线前缀，与邻近 `sampler` 状态字段风格不一致，建议去掉 ("nit - let's avoid the underscore for consistency with other fields here")。作者回复 "Done, renamed it for consistency. Thanks!"，两处修改均已在提交 16713c9 中合入。

- 移除单次调用的 `_requires_logits_processing` 方法，内联检查 (design): 已采纳，检查直接内联在 `apply_sampling_params()` 中，方法体删除。
- 字段命名去掉下划线前缀，与其他 sampler 状态字段保持一致 (style): 已重命名，作者回复 "Done, renamed it for consistency. Thanks!", 随提交 16713c9 合入。

风险与影响

- 风险:
 1. 手写组合谓词的同步风险: `needs_logits_processing` 是对真实子状态的镜像，未来若新增采样状态（如新的惩罚类型、新的采样参数）而忘记同步 `add_request()` 中的 OR 表达式，热路径会静默跳过该处理，比旧实现（每次实时扫描）更隐蔽，且难以通过现有测试发现——这是结构性风险。
 2. 缓存时序耦合: 缓存值在 `add_request()` 时读取各子状态的 `.np` 数组；若未来某个子状态改成 staged-write 延迟提交且读取时机在提交前，缓存取值可能与实际采样阶段不一致。当前各子状态 `add_request` 后即更新数组，且 `test_v2_greedy_sampling_applies_thinking_budget` 覆盖了 `apply_staged_writes()` 之后采样的路径，风险较低但需保持警觉。
 3. 行为变更的合并顺序依赖: 组合谓词纳入思考预算是一个行为修复 (greedy 请求不再绕过 thinking budget)，该修复内容来自 #52284 并以 rebase 方式带入本分支；若 #52284 未先合入，单独合入本 PR 会改变 greedy 行为。PR body 已说明分支结构保证合并顺序。
 4. 收益限于 CPU 决策路径: 端到端收益只在 CPU 采样占比高的场景显著 (Qwen3-0.6B 高并发 +4.53% 吞吐)，GPU-bound 场景 (Qwen3.5-9B) 无增益，性能预期需按负载特征校准。- 影响: 影响 MRV2 (Model Runner V2) 默认执行路径的采样器模块: 每次采样步骤的 logits 处理判定从多数组扫描降为单数组 `np.any()`，对默认参数、纯 greedy 且无思考预算的请求收益最大 (决策路径 -70%，微基准 56.7 us -> 16.8 us)。对用户透明: Qwen3.5-9B 固定种子 A/B 矩阵 (default/greedy/temperature/top-p/top-k/min-p/penalties/logit bias/bad words/thinking budget) 与 GSM8K 100 题评估完全一致。对团队而言，本 PR 建立了“请求级预计算标志位 + 槽位复用覆盖”的模式，可复用于其他 per-request 分支判定，同时也暴露了手写谓词与真实状态之间的同步责任。- 风险标记: 核心采样热路径，手写组合谓词与子状态同步风险，thinking budget 行为变更，端到端收益依赖 CPU 采样占比

关联脉络

- PR #52284 [Bugfix][MRV2] Apply thinking budget to greedy requests: PR body 声明本 PR 依赖 #52284 (active thinking budget 必须参与组合谓词)，本分支第一个提交即 rebase 的 #52284 内容；该修复是组合谓词新增 thinking budget 判据的前提。
- PR #52311 [Bugfix][Model Runner V2][Spec Decode] Fix off-by-one in bad_words draft-prefix matching: 同属 MRV2 采样器子模块 (vllm/v1/worker/gpu/sample/) 的近期正确性修复，修改 bad_words.py 与对应单测；与本 PR 的 bad words 谓词判定在同一文件域内。

- PR #49613 [Bugfix][Sampling] Clear empty side on thinking-budget asymmetric SWAP: 修改 thinking_budget_state.py 与对应测试，与本 PR 组合谓词中的 thinking budget 判据同属思考预算状态机，体现 MRV2 采样器正确性收敛期。