

PR #52164 完整报告

vllm-project/vllm

[Attention][DSA] Take the native decode path for MTP=3 on SM90

合并时间: 2026-08-14 18:09

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52164>

执行摘要

- 一句话: SM90 上 MTP=3 改用 DeepGEMM 原生 decode, 免 KV 重复读取
- 推荐动作: 值得精读。三个设计亮点: 1) 把“内核能力”从 indexer 中抽出到 DeepGEMM 封装层 (`native_next_n_supported`), 并用架构族矩阵测试锁定行为, 避免重复的硬编码分支; 2) “合法性是 step 的属性而非配置的属性”这一洞见 (内核看到的是 `max_decode_len` 行 Q), 避免了一批恰好 3 token 深的批次走上不存在的 SM90 内核; 3) 调度 metadata 槽数由封装层自推导、`num_sms` 保持字面含义的设计哲学, 从 API 形态上杜绝调用方算错。需要留意的是门控与 DeepGEMM 内核版本强耦合, 升级 `deepgemm` 子模块时必须同步验证 SM90 的 {1, 2, 4} 支持面。

功能与动机

Issue #35878 明确指出: DeepGEMM 的 NVIDIA 维护分支 `nv_dev` 已支持 MTP3 ('This would allow us to avoid batch expansion and get speedups for `num_speculative_tokens=3`, when deploying DeepSeek Sparse Attention in DeepSeek V3.2 and GLM5'), 且 TRTLLM 已基于该分支实现, 而 'Current implementation uses batch expansion, which may not be as efficient, especially for long sequences'。DSA indexer 原先对 `next_n` 不在 {1, 2} 的配置一律展开, MTP=3 时每个请求的 KV tile 被读 4 次, 长上下文低延迟解码场景损失明显。PR body 还说明 #45322 已为 SM100 打通等价路径, 本 PR 是 issue 明确要求的 SM90 补齐。

实现拆解

按 5 个步骤拆解实现:

1. 能力判定下沉到 DeepGEMM 封装层 (`vllm/utils/deep_gemm.py`): 新增 `native_next_n_supported(next_n)`, SM90 家族仅返回 {1, 2, 4} (明确不含 3, 因为 3 没有对应内核), 其余架构族一律返回 `True`; 新增 `_paged_mqa_logits_schedule_slots(num_sms, next_n)`, 对 SM90 且 `next_n == 4` 按 2-CTA multicast 集群把调度槽数减半。`get_paged_mqa_logits_metadata` 改为从 `context_lens` 的二维形状推导 `next_n`, 内部换算 `num_slots` 后传给底层实现, `num_sms` 参数保持“SM 总数”的字面含义, 调用方不可能算错。
2. indexer 门控重构 (`vllm/v1/attention/backends/mla/indexer.py`): 新增 `_supports_native_decode(next_n)`, 在 CUDA + DeepGEMM 前提下按架构族分派——SM100 全放行、SM90 走 `native_next_n_supported`、其余仅放行 {1, 2}; 替换原先

硬编码的“非 SM100 且 next_n 不在 {1, 2} 即 flatten”逻辑, `self.use_flattening = not _supports_native_decode(next_n)`。

3. 运行时合法性改为 step 属性: `build()` 中新增 `step_next_n_ok = max_decode_len <= 1 or _supports_native_decode(max_decode_len)`, 因为内核实际收到的是 `max_decode_len` 行 Q 而非配置的 `next_n`; SM90 上一批恰好 3 token 深的批次没有原生内核, 必须动态回退到 flatten。
4. 调度元数据槽数自适应: `scheduler_metadata_buffer` 保持按 `(num_sms + 1, 2)` 预分配以兼容 CUDA graph, `build()` 里用 `get_paged_mqa_logits_metadata` 返回张量的形状把 buffer 窄化为前缀视图再拷贝, 满足内核“metadata 槽数 = num_sms / num_kv_multicast”的断言, 避免 `fp8_fp4_paged_mqa_logits` 直接崩溃。
5. 测试配套: 新增 `tests/v1/attention/test_indexer_native_next_n.py`, 用 monkeypatch 模拟 SM90/100/120 架构族与有无 DeepGEMM 的组合, 覆盖门控行为矩阵、SM90 槽数减半与 multicast 仅 SM90; `tests/kernels/attention/test_deepgemm_attention.py` 参数化为 (4, 1)、(2, 2)、(2, 4) 三组, 其中 `next_n=4` 对照参考实现验证数值, 非支持架构自动 skip。无部署或配置配套改动, 启动日志保留 `use_flattening` 标记便于确认实际路径。

关键文件:

- `vllm/v1/attention/backends/mla/indexer.py` (模块 注意力后端; 类别 source; 类型 core-logic; 符号 `_supports_native_decode`): DSA indexer 的核心门控位置: 新增 `_supports_native_decode` 按架构族分派原生 decode 能力, 替换硬编码的 flatten 判断; `build()` 内新增按 step 动态判断的 `step_next_n_ok`, 并把调度 metadata 窄化为内核实际槽数后再拷贝, 是本次行为变更的入口与最核心文件。
- `vllm/utils/deep_gemm.py` (模块 算子封装; 类别 source; 类型 core-logic; 符号 `native_next_n_supported`, `_paged_mqa_logits_schedule_slots`): 内核能力判定与调度槽数计算的唯一权威来源: 新增 `native_next_n_supported` (SM90 仅 {1,2,4}) 与 `_paged_mqa_logits_schedule_slots` (SM90 `next_n=4` 槽数减半), 并让 `get_paged_mqa_logits_metadata` 从 `context_lens` 二维形状自推导 `next_n`, 保证 `num_sms` 语义不被调用方误用。
- `tests/v1/attention/test_indexer_native_next_n.py` (模块 注意力后端; 类别 test; 类型 test-coverage; 符号 `_set_arch`, `test_native_decode_gate_per_architecture`, `test_native_decode_gate_without_deepgemm`, `test_sm90_next_n_4_halves_the_schedule_slots`): 新增测试文件, 用 monkeypatch 模拟 SM90/100/120 架构族与有无 DeepGEMM 的组合, 把 `_supports_native_decode` 的行为矩阵、SM90 槽数减半、multicast 仅 SM90 等关键决策固化为回归防线; 测试 docstring 还揭示了“门控错误 = 崩溃而非慢路径”的严重性。
- `tests/kernels/attention/test_deepgemm_attention.py` (模块 内核测试; 类别 test; 类型 test-coverage; 符号 `test_deepgemm_fp8_fp4_paged_mqa_logits`): 内核级数值验证: 把 `test_deepgemm_fp8_fp4_paged_mqa_logits` 参数化为 (4,1)、(2,2)、(2,4), 新增 `next_n=4` 对照参考实现, 直接覆盖本次 SM90 原生路径的核心内核行为。

关键符号: `_supports_native_decode`, `native_next_n_supported`,
`_paged_mqa_logits_schedule_slots`, `get_paged_mqa_logits_metadata`,
`DeepseekV32IndexerMetadataBuilder.build`

关键源码片段

vllm/v1/attention/backends/mla/indexer.py

DSA indexer 的核心门控位置：新增 `_supports_native_decode` 按架构族分派原生 decode 能力，替换硬编码的 `flatten` 判断；`build()` 内新增按 `step` 动态判断的 `step_next_n_ok`，并把调度 `metadata` 窄化为内核实际槽数后再拷贝，是本次行为变更的入口与最核心文件。

```
# vllm/v1/attention/backends/mla/indexer.py: decode 路径门控与 step 级合法性
```

```
def _supports_native_decode(next_n: int) -> bool:
```

```
    """decode 是否可以把每个请求的 next_n 行 Q 直接交给内核。
```

```
    展开 (flatten) 把一个请求变成 next_n 个单 token 伪请求，同一块
    KV tile 被重复读取 next_n 次；原生路径让内核看到真实的 (B, next_n)
    批次。门控按架构族分派，而不是简单的阈值比较。
```

```
    """
```

```
    if not (current_platform.is_cuda() and has_deep_gemm()):
        # 没有 DeepGEMM 内核时，只放行所有 backend 都能处理的形状
        return next_n in (1, 2)
    if current_platform.is_device_capability_family(100):
        return True # SM100: multi-atom 磁贴调度任意 next_n
    if current_platform.is_device_capability_family(90):
        return native_next_n_supported(next_n) # SM90: 仅 {1, 2, 4}
    return next_n in (1, 2)
```

```
# build() 中：合法性是 step 的属性而非配置的属性
```

```
max_decode_len = int(decode_lens_cpu.max().item())
```

```
next_n = 1 + self.num_speculative_tokens
```

```
# 内核实际收到 max_decode_len 行 Q，而不是配置的 next_n；SM90 上一批
```

```
# 恰好 3 token 深的批次没有原生内核，必须动态回退到 flatten
```

```
step_next_n_ok = max_decode_len <= 1 or _supports_native_decode(max_decode_len)
```

```
use_native = (
```

```
    not (self.use_flattening or self.supports_varlen)
```

```
    and max_decode_len <= next_n
```

```
    and step_next_n_ok
```

```
)
```

```
# 调度 metadata 窄化为内核实际需要的槽数再拷贝；buffer 保持预分配地址，
```

```
# 兼容 CUDA graph 重放，同时满足内核的槽数断言
```

```
schedule_metadata = self.scheduler_metadata_buffer
```

```
if current_platform.is_cuda() and has_deep_gemm():
```

```
    metadata = get_paged_mqa_logits_metadata(
```

```
        seq_lens, self.kv_cache_spec.storage_block_size, self.num_sms,
```

```
        indices=decode_indices,
```

```
    )
```

```
    schedule_metadata = self.scheduler_metadata_buffer[: metadata.shape[0]]
```

```
    schedule_metadata[:] = metadata
```

vllm/utils/deep_gemm.py

内核能力判定与调度槽数计算的唯一权威来源：新增 `native_next_n_supported` (SM90 仅 {1,2,4}) 与 `_paged_mqa_logits_schedule_slots` (SM90 next_n=4 槽数减半)，并让 `get_paged_mqa_logits_metadata` 从 `context_lens` 二维形状自推导 `next_n`，保证 `num_sms` 语义不被调用方误用。

vllm/utils/deep_gemm.py: 内核能力与调度槽数统一收敛在 DeepGEMM 封装层

```
def native_next_n_supported(next_n: int) -> bool:
```

```
    """判断 paged MQA logits 内核能否以原生 `next_n` 行 Q 处理请求。
```

```
    SM90 (Hopper) 只有 {1, 2, 4} 原生实现: 4 依赖 DeepGEMM nv_dev 分支的
    2-CTA multicast 集群, 3 没有对应内核, 因此这里不是简单的阈值判断。
```

```
    SM100 及以上通过 multi-atom 磁贴调度任意 next_n, 一律放行; 不支持
    的取值必须由调用方展开为单 token 行 (flatten)。
```

```
    """
```

```
    if current_platform.is_device_capability_family(90):
```

```
        return next_n in (1, 2, 4)
```

```
    return True
```

```
def _paged_mqa_logits_schedule_slots(num_sms: int, next_n: int) -> int:
```

```
    """paged MQA logits 内核实际启动的调度任务数。
```

```
    SM90 且 next_n == 4 时按 2-CTA multicast 集群调度, 一个任务覆盖两个
    SM, 槽数需要减半; `fp8_fp4_paged_mqa_logits` 会对 metadata 槽数做
    断言, 两者不一致会直接崩溃而不是走慢路径。
```

```
    """
```

```
    num_kv_multicast = (
```

```
        2 if next_n == 4 and current_platform.is_device_capability_family(90) else 1
```

```
)
```

```
    return num_sms // num_kv_multicast
```

```
def get_paged_mqa_logits_metadata(context_lens, block_size, num_sms, indices=None):
```

```
    """构建调度 metadata; 从 context_lens 的 2D 形状推导 next_n 并换算槽数。
```

```
    num_sms 保持“SM 总数”的字面含义, 调用方不需要也不应该自己除 2。
```

```
    """
```

```
    _lazy_init()
```

```
    if _get_paged_mqa_logits_metadata_impl is None:
```

```
        return _missing()
```

```
    next_n = context_lens.shape[1] if context_lens.dim() == 2 else 1
```

```
    num_slots = _paged_mqa_logits_schedule_slots(num_sms, next_n)
```

```
    kwargs = {} if indices is None else {"indices": indices}
```

```
    return _get_paged_mqa_logits_metadata_impl(
```

```
        context_lens, block_size, num_slots, **kwargs
```

```
)
```

tests/v1/attention/test_indexer_native_next_n.py

新增测试文件，用 monkeypatch 模拟 SM90/100/120 架构族与有无 DeepGEMM 的组合，把 `_supports_native_decode` 的行为矩阵、SM90 槽数减半、multicast 仅 SM90 等关键决策固化为回归防线；测试 docstring 还揭示了“门控错误 = 崩溃而非慢路径”的严重性。

```
# tests/v1/attention/test_indexer_native_next_n.py: 门控行为矩阵回归测试
NUM_SMS = 114 # H100 PCIe
```

```
def _set_arch(monkeypatch, family: int, *, cuda: bool = True, deep_gemm: bool = True):
    """把当前平台伪装成指定 CUDA 架构族，便于离线验证门控逻辑。"""
    monkeypatch.setattr(current_platform, "is_cuda", lambda: cuda)
    monkeypatch.setattr(
        current_platform,
        "is_device_capability_family",
        lambda capability, device_id=0: capability // 10 == family,
    )
    monkeypatch.setattr(indexer, "has_deep_gemm", lambda: deep_gemm)
```

```
@pytest.mark.parametrize(
    "family,expected_native",
    [
        # SM90 通过 2-CTA multicast 拿到了 next_n=4（即 MTP=3），但永远没有 3
        (9, {1, 2, 4}),
        # SM100 用 multi-atom 磁贴调度任意 next_n
        (10, {1, 2, 3, 4, 5, 8}),
        # SM120 也声称支持 multi-atom 但未在硬件上验证，先保持保守门控
        (12, {1, 2}),
    ],
)
def test_native_decode_gate_per_architecture(monkeypatch, family, expected_native):
    _set_arch(monkeypatch, family)
    for next_n in (1, 2, 3, 4, 5, 8):
        assert indexer._supports_native_decode(next_n) == (next_n in expected_native), (
            f"family={family} next_n={next_n}"
        )
```

```
def test_sm90_next_n_4_halves_the_schedule_slots(monkeypatch):
    """SM90 的 next_n=4 按 2-CTA 集群调度：一个任务覆盖两个 SM。"""
    _set_arch(monkeypatch, 9)
    assert _paged_mqa_logits_schedule_slots(NUM_SMS, 4) == NUM_SMS // 2
    for next_n in (1, 2, 3):
        assert _paged_mqa_logits_schedule_slots(NUM_SMS, next_n) == NUM_SMS
```

评论区精华

该 PR 没有实质性的 review 评论线程（review_comments_count 为 0）。评论区的交互主要是流程性的：claude[bot] 指出 fork PR 自动评审被禁用，ZJY0516 随后直接批准并触发 `/ci`

run, github-actions 启动了 Buildkite CI #83864。真正有信息量的设计讨论都写在 PR body 的方法学说明中，相当于作者对评审常见质疑的预答复：用 `--num-gpu-blocks-override` 钉住两侧 KV cache、用 `--num-prompts == --max-concurrency` 保证单波次 decode-bound，并主动披露 'The concurrency-4 cell does not (its TPOT spread is $\pm 9\%$, larger than the effect)', 以两次运行展示波动范围而非夸大收益。此外作者明确了两个关键设计取舍：一是 batch=1/4K 时内核级原生路径比 flatten 慢 3%，但端到端未出现该回归，故门控保持无条件放行；二是 next_n=3 没有 SM90 内核，合法性必须按 step 动态判断（14 万 decode steps 实测从未出现 max_decode_len=3）。

- fork PR 自动化评审状态 (other): 未触发一次性评审，随后由维护者 ZJY0516 直接批准合并。
- CI 触发与验证 (other): CI 已触发并随合并完成验证（PR 带 ready/verified 标签）。
- SM90 next_n=4 门控与 batch=1 回归取舍 (design): 接受无条件门控 + per-step 兜底的设计，作者用路径覆盖数据支撑该决策。

风险与影响

- 风险：
 1. 门控即崩溃点：新测试文件 docstring 明确 'Getting this wrong is not a slow path but a crash'。native_next_n_supported 的 {1, 2, 4} 假设与 DeepGEMM nv_dev 分支内部实现强耦合，未来 DeepGEMM 子模块升级若改变 SM90 的 next_n 支持面或 multicast 集群尺寸，会直接触发内核断言而非优雅降级。
 2. 调度槽减半假设：_paged_mqa_logits_schedule_slots 硬编码 2-CTA 减半；当前 H100 PCIe 的 114 SM 为偶数，若未来出现奇数 SM 或不同集群尺寸的机器，num_sms // 2 的取整可能与内核期望不一致。
 3. 兜底路径未实测：max_decode_len=3 的 step 在 SM90 走 flatten 防御路径，作者实测 14 万 steps 从未出现，该分支属于稳态之外的代码路径；遇到新模型行为变化时只影响性能、不影响正确性。
 4. 收益边界有限：端到端吞吐仅提升 0.9%–1.1%，且 concurrency=4 的 TPOT 改善落在各臂 $\pm 9\%$ 运行波动的范围内；paged MQA logits 内核只占 671B MoE decode step 的一小部分，这是收益的天花板。
 5. CUDA graph 兼容：metadata buffer 改为前缀视图窄化，buffer 基址仍是预分配的固定地址，cudagraph 两种模式 (eager 与 FULL_AND_PIECEWISE) 均已实测通过。
 - 影响：用户侧：H100/H200 上 DeepSeek-V3.2 + DSA + MTP=3 配置在长上下文、高并发下 decode 延迟改善（内核级 32K 上下文 batch=64 时快 1.46 倍）；其余平台与配置 (SM8x、SM12x、ROCm、next_n 为 1/2) 行为不变。系统侧：仅影响 DeepseekV32IndexerMetadataBuilder 构建的 DSA decode 路径，默认路径无变化；启动日志的 use_flattening 标记可确认实际走哪条路径，便于线上核对。团队侧：_supports_native_decode 加行为矩阵测试提供了跨架构能力扩展的可复制模式，后续放开 SM120 或新增内核只需改一处并附测量数据即可。
 - 风险标记：SM90 原生路径错误将触发内核断言崩溃，调度槽减半与 DeepGEMM 2-CTA 行为强耦合，max_decode_len=3 兜底路径无稳态实测，端到端收益有限（约 1%），DeepGEMM 子模块升级需同步验证门控

关联脉络

- PR #45322 [Attention][DSA] SM100 原生 next_n 路径 (PR body 提及的对照实现) : PR body 明确 '#45322 did the equivalent for SM100; this is the SM90 case the issue asks for'。本 PR 是其 SM90 补齐，行为矩阵中 SM100 列保持原生。
- PR #50062 [Model Runner V2][Spec Decode] Add KV cache support for multi-layer MTP: 同一 MTP speculative decode 功能线的基础设施 PR，为多模块 MTP 增加调度与 KV cache 支持，是 MTP=3 场景的前置能力。
- PR #51674 [Kernel][Perf] Add fused CUDA post-conv MTP decode kernel for Qwen3.5 GDN: 同为 speculative decode 加内核级性能优化的方向，可对照不同模型族 (Qwen GDN vs DeepSeek DSA) 的 MTP 解码优化手法。