

# PR #48597 完整报告

vllm-project/vllm

[Perf][GLM-5.2] Blackwell decode optimizations

合并时间: 2026-07-24 12:36

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/48597>

## 执行摘要

- 一句话: GLM-5.2 Blackwell 解码优化及后续拆分
- 推荐动作: 虽然该 PR 已被 revert, 但其设计决策和优化技术值得精读:
- CuTeDSL fused Q 核: 如何利用 CuTe DSL 将多个小算子融合为单 kernel, 减少全局内存往返。
- PDL 的使用: 通过 GPU 硬件支持的依赖启动机制减少内核启动延迟, 是 Blackwell 性能优化的关键之一。
- bf16\_skinny\_gemm dispatch: 基于 shape 的 GEMM 调度策略, 为小 M 问题选择定制实现, 可作为类似场景的参考。
- MTP 索引共享: 通过生命周期钩子实现在多步投机解码中复用预填充计算结果, 减少冗余计算。
- 拆分策略: 从大型 PR 到多个聚焦 PR 的拆分模式, 体现了如何管理高风险变更。

建议阅读该 PR 的 body (跟踪 checklist) 和后续拆分 PR (#49790-#49793, #50230) 的实现细节。

## 功能与动机

PR body 引用 benchmark 数据: 目标是在 Blackwell 上提升 GLM-5.2/DeepSeek-V3.2 低延迟 decode 路径的端到端性能。验证结果表明, 全系列优化后 output tok/s 从 446.7 提升至 542.5 (+21.4%), median TPOT 从 1.94 ms 降至 1.56 ms (-19.6%), GSM8K 500-shot 准确率保持 0.950。关联 Issue #35161 修复了 MoE 对齐填充问题, 是优化基础之一。

## 实现拆解

实现步骤按模块拆解如下:

1. 自定义 fused\_q CuTeDSL 算子([vllm/models/deepseek\\_v32/nvidia/ops/fused\\_q\\_cutedsl.py](#), [kernels.py](#)): 新增 CuTeDSL 融合 Q 核, 将 MLA 的 Q 投影、RoPE、FP8 量化和 indexer 操作合并为一个 GPU 内核。通过 [direct\\_register\\_custom\\_op](#) 注册为 PyTorch 自定义 op, 并为 torch.compile 提供 fake 实现。
2. fused\_norm\_rope 内核 PDL 支持([kernels.py](#)): 在 Triton 核函数中添加 [USE\\_PDL](#) 编译常量, 当硬件支持 (SM100) 时插入 [tl.extra.cuda.gdc\\_wait\(\)](#) 和

`tl.extra.cuda.gdc_launch_dependents()` 调用，实现程序化依赖启动，减少内核启动延迟。同时将 Q RMS Norm 提前到 kernel 的 `pid==2` 分支，使其不依赖缓存可用性，优化 profile 流程。

3. Decode-M GEMM 调度优化(`vllm/models/deepseek_v32/nvidia/attention.py`, `mtp.py`): 针对 bf16 A/B 投影 (`qkv_a`、`q_b`、`eh_proj`) 在 Blackwell 上引入 `bf16_skinny_gemm` 和 `dsv3_fused_a_gemm` 定制 GEMM，根据不同形状 (N, K) 配置 `skinny_max` 阈值，在小 M (2-3) 时使用皮肤 GEMM，中 M 时使用 `fused_a`，大 M 时回退 cuBLAS，经验测速选择最优路径。
4. MTP 投机解码优化(`vllm/v1/worker/gpu/spec_decode/mtp/speculator.py`, `vllm/models/deepseek_v32/nvidia/mtp.py`): 实现 MTP 迭代间的 top-k 索引共享 (`set_skip_topk`)，预填充步骤 0 计算索引，后续步骤复用。将 MTP 层的末尾 all-reduce 融合到最终 RMSNorm 中 (`fused_allreduce_rms_norm`)，并实现 `get_top_tokens` 支持本地 argmax 归约。
5. 生命周期钩子框架(`vllm/v1/worker/gpu/spec_decode/autoregressive/speculator.py`): 在 `AutoRegressiveSpeculator` 中新增 `on_prefill_begin/end`、`on_multi_step_decode_begin/end` 空方法，供子类覆盖。这些钩子在 `capture` 和 `propose` 的对应阶段调用，确保捕获和回放时状态一致。MTPSpeculator 利用该钩子实现索引共享的开关。
6. 注意力管理层配套(`vllm/models/deepseek_v32/nvidia/attention.py`, `vllm/v1/attention/backends/mla/sparse_utils.py`): 在 `DeepseekV32Attention` 初始化中强制 `cache_dtype='fp8'` (当 auto 时)，以使用 FP8 sparse cache。新增 `phys_shadow` 弱引用注册表，缓存 indexer 的物理索引转换结果，避免在 `skip_topk` 层重复计算。

关键文件:

- `vllm/models/deepseek_v32/nvidia/kernels.py` (模块 模型内核; 类别 source; 类型 core-logic; 符号 `_can_use_fused_q_cuteds1`, `_is_arch_support_pdl`, `_fused_q_cuteds1_impl`, `_fused_q_cuteds1_fake`): 核心 kernel 封装, 新增 `fused_q_cuteds1` 自定义 op 注册和 PDL 支持, 是解码优化的关键
- `vllm/models/deepseek_v32/nvidia/ops/fused_q_cuteds1.py` (模块 融合算子; 类别 infra; 类型 infrastructure; 符号 `_make_fake_tensor`, `is_fused_q_cuteds1_supported`, `fused_q_cuteds1`, `FusedQKernel`): 新增 CuTeDSL 融合 Q 核实现, 通过编译时特化生成高效 kernel, 是 decode 性能提升的核心
- `vllm/v1/worker/gpu/spec_decode/autoregressive/speculator.py` (模块 投机解码; 类别 source; 类型 core-logic; 符号 `on_prefill_begin`, `on_prefill_end`, `on_multi_step_decode_begin`, `on_multi_step_decode_end`): 新增生命周期钩子框架, 为模型特定优化提供可扩展接口, MTP 索引共享基于此实现
- `vllm/models/deepseek_v32/nvidia/attention.py` (模块 注意力层; 类别 source; 类型 core-logic; 符号 `_decode_m_gemm`): 实现 Decode-M GEMM 调度 (`bf16_skinny_gemm/fused_a`), 强制 FP8 sparse cache, 是性能关键路径
- `vllm/models/deepseek_v32/nvidia/mtp.py` (模块 MTP 头; 类别 source; 类型 core-logic; 符号 `get_top_tokens`): MTP 层优化: `bf16_skinny_gemm`、`fused allreduce+rms`、

get\_top\_tokens, 提升 MTP 解码效率

- vllm/v1/worker/gpu/spec\_decode/mtp/speculator.py (模块 MTP 投机器; 类别 source; 类型 core-logic; 符号 init, on\_prefill\_end, on\_multi\_step\_decode\_begin, on\_multi\_step\_decode\_end) : 利用生命周期钩子实现 MTP top-k 索引共享, 减少重复计算
- vllm/v1/attention/backends/mla/sparse\_utils.py (模块 稀疏注意力; 类别 source; 类型 core-logic; 符号 register\_phys\_shadow, phys\_shadow) : 新增 phys\_shadow 弱引用注册表, 缓存 indexer 物理索引转换结果, 避免重复计算
- vllm/\_custom\_ops.py (模块 自定义算子; 类别 source; 类型 infrastructure; 符号 bf16\_skinny\_gemm, bf16\_skinny\_gemm\_fake) : 注册 bf16\_skinny\_gemm 自定义 op, 使 GEMM 调度在 torch.compile 下可用
- tests/kernels/test\_fused\_deepseek\_v32\_norm\_rope.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 test\_platform\_capability\_queries\_are\_constant\_during\_compile, capability\_branches, test\_pdl\_is\_disabled\_before\_blackwell, test\_fused\_norm\_rope\_profile\_without\_cache\_compiles) : 验证 fused\_norm\_rope 在 PDL 开启 / 关闭时的正确性和编译常量行为
- tests/kernels/test\_bf16\_skinny\_gemm.py (模块 单元测试; 类别 test; 类型 test-coverage; 符号 \_rel\_err, test\_bf16\_skinny\_gemm\_matches\_reference, test\_bf16\_skinny\_gemm\_strided\_output) : 验证 bf16\_skinny\_gemm 与参考实现的数值一致性和 stride 输出支持
- vllm/compilation/passes/fusion/allreduce\_rms\_fusion.py (模块 编译优化; 类别 source; 类型 core-logic; 符号 unfused\_fallback) : 支持 fused allreduce\_rms\_norm 的大小回退和编译安全修复
- vllm/model\_executor/layers/fused\_allreduce\_gemma\_rms\_norm.py (模块 算子层; 类别 source; 类型 data-contract; 符号 \_fi\_ar\_max\_size\_mb) : 配合编译安全修复, 调整 max\_token\_num 计算方式

关键符号: \_can\_use\_fused\_q\_cutedsl, \_is\_arch\_support\_pdl, \_fused\_q\_cutedsl\_impl, fused\_q\_cutedsl, FusedQKernel.compile, FusedQKernel.call, fused\_norm\_rope, \_fused\_norm\_rope\_kernel, DeepseekV32Attention.init, DeepseekV32Attention.\_decode\_m\_gemm, MTPSpeculator.init, MTPSpeculator.on\_prefill\_end, MTPSpeculator.on\_multi\_step\_decode\_begin, MTPSpeculator.on\_multi\_step\_decode\_end, AutoRegressiveSpeculator.on\_prefill\_begin, AutoRegressiveSpeculator.on\_prefill\_end, AutoRegressiveSpeculator.on\_multi\_step\_decode\_begin, AutoRegressiveSpeculator.on\_multi\_step\_decode\_end, register\_phys\_shadow, phys\_shadow, bf16\_skinny\_gemm, get\_top\_tokens

## 关键源码片段

### vllm/models/deepseek\_v32/nvidia/kernels.py

核心 kernel 封装, 新增 fused\_q\_cutedsl 自定义 op 注册和 PDL 支持, 是解码优化的关键

```
# SPDX-License-Identifier: Apache-2.0
import torch
```

```

from vllm.platforms import current_platform
from vllm.utils import import_utils import has_cutedsl
from vllm.utils.torch_utils import direct_register_custom_op


# 检测是否可使用 CuTeDSL 融合 Q 核 (仅 Blackwell SM100 + cutedsl 可用)
@torch.compiler.assume_constant_result
def _can_use_fused_q_cutedsl() -> bool:
    return current_platform.has_device_capability(100) and has_cutedsl()


# 检测架构是否支持程序化依赖启动 (PDL), 用于减少内核启动延迟
@torch.compiler.assume_constant_result
def _is_arch_support_pdl() -> bool:
    return (
        current_platform.has_device_capability(100)
        and current_platform.is_arch_support_pdl()
    )


# fused_q_cutedsl 的真实实现, 委托给 ops/fused_q_cutedsl.py
def _fused_q_cutedsl_impl(
    positions: torch.Tensor,
    q_pe: torch.Tensor,
    rope_cache: torch.Tensor,
    ql_nope: torch.Tensor,
    q_scale: torch.Tensor,
    mqa_output: torch.Tensor,
    idx_q: torch.Tensor,
    idx_rope_cache: torch.Tensor,
    idx_weights: torch.Tensor,
    idx_weights_softmax_scale: float,
    idx_weights_head_scale: float,
    idx_q_fp8: torch.Tensor,
    idx_weights_out: torch.Tensor,
    has_indexer: bool,
    index_rope_interleave: bool,
) -> None:
    from .ops.fused_q_cutedsl import fused_q_cutedsl

    fused_q_cutedsl(
        positions,
        q_pe,
        rope_cache,
        ql_nope,
        q_scale,
        mqa_output,
        idx_q,
        idx_rope_cache,
    )

```

```

        idx_weights,
        idx_weights_softmax_scale,
        idx_weights_head_scale,
        idx_q_fp8,
        idx_weights_out,
        has_indexer=has_indexer,
        index_rope_interleave=index_rope_interleave,
    )

```

```

# fake 实现, 用于 torch.compile 的图捕获
def _fused_q_cuteds1_fake(*args, **kwargs) -> None:
    pass

```

```

# 注册为自定义 op, 声明变更为 mqa_output, idx_q_fp8, idx_weights_out
direct_register_custom_op(
    op_name="fused_q_cuteds1",
    op_func=_fused_q_cuteds1_impl,
    mutates_args=["mqa_output", "idx_q_fp8", "idx_weights_out"],
    fake_impl=_fused_q_cuteds1_fake,
    dispatch_key="CUDA",
)

```

同时, 在 `_fused_norm_rope_kernel` 中新增 `USE_PDL` 编译常量:

```

@triton.jit
def _fused_norm_rope_kernel(
    ...,
    USE_PDL: tl.constexpr, # 新增: 是否启用程序化依赖启动
):
    pid = tl.program_id(0)
    tok_idx = tl.program_id(1)
    if USE_PDL:
        tl.extra.cuda.gdc_wait() # 等待先前依赖完成
        tl.extra.cuda.gdc_launch_dependents() # 声明当前内核产生依赖

    if pid == 3:
        # topk 索引清理逻辑 (略)
        ...
    elif pid == 2:
        # Q RMS Norm: 早于缓存可用性, profile 阶段也需要
        q_block = tl.arange(0, Q_BLOCK_SIZE)
        q_mask = q_block < Q_DIM
        q_c = tl.load(q_c_ptr + tok_idx * q_c_stride + q_block, mask=q_mask, other=0.0)
        q_c_rms_w = tl.load(q_rms_norm_w_ptr + q_block, mask=q_mask)
        q_c = _rms_norm(q_c, q_c_rms_w, q_rms_eps, Q_DIM)
        tl.store(q_c_out_ptr + tok_idx * q_c_out_stride + q_block, q_c, mask=q_mask)
    elif pid == 1:
        # KV RMS Norm + KV RoPE + MLA concat_and_cache (略)

```

...

## vllm/v1/worker/gpu/spec\_decode/autoregressive/speculator.py

新增生命周期钩子框架，为模型特定优化提供可扩展接口，MTP 索引共享基于此实现

```
class AutoRegressiveSpeculator(DraftModelSpeculator):
    def __init__(self, vllm_config: VllmConfig, device: torch.device):
        super().__init__(vllm_config, device)
        # ... 其他初始化 ...

    # 生命周期钩子（空实现，子类覆盖）
    # 这些钩子在 capture 和 propose 的对应阶段被调用，
    # 确保捕获和回放时状态一致（例如 attention flags 必须相同）
    def on_prefill_begin(self, num_reqs: int, num_tokens: int) -> None:
        ...

    def on_prefill_end(self, num_reqs: int, num_tokens: int) -> None:
        ...

    def on_multi_step_decode_begin(self, num_reqs: int) -> None:
        ...

    def on_multi_step_decode_end(self, num_reqs: int) -> None:
        ...

    def capture(self) -> None:
        # ... 前置操作 ...
        self.on_prefill_begin(self.max_num_reqs, self.max_num_tokens)
        # 执行 prefill 图捕获
        self.prefill_cudagraph_manager.capture(
            self._prefill,
            self.model_state,
            self.target_input_buffers,
            self.block_tables,
            self.target_attn_groups,
            self.kv_cache_config,
            progress_bar_desc="Capturing prefill CUDA graphs",
        )
        self.on_prefill_end(self.max_num_reqs, self.max_num_tokens)

        if self.num_speculative_steps == 1:
            return

        self.on_multi_step_decode_begin(self.max_num_reqs)
        # 执行 decode 图捕获
        self.decode_cudagraph_manager.capture(
            self._generate_draft,
            self.model_state,
            self.input_buffers,
```



```
        self.block_tables,
        self.attn_groups,
        self.kv_cache_config,
        progress_bar_desc="Capturing decode CUDA graphs",
    )
    self.on_multi_step_decode_end(self.max_num_reqs)
```

## 评论区精华

Review 中主要讨论了以下关键点：

- 性能验证要求：yewentao256 要求提供 `vllm bench serve` 端到端性能和 `lm_eval` 准确率数据。zhou9402 在 GB300 上提供了 2 组对比数据，展示 +21.4% 吞吐提升且 GSM8K 准确率保持 0.950。
- PR 规模问题：yewentao256 指出“The PR becomes too large to merge, could you shrink diff as much as possible? Ideally < 500 LOC”。这最终导致 PR 被拆分。
- Revert 决定：WoosukKwon 评论“I've reverted this PR since it didn't go through the proper review process, and it includes some questionable design choices. Will discuss with zhou9402 offline”。
- `phys_shadow` 设计讨论：chaunceyjiang 建议将物理索引转换结果存储在 `attention metadata` 中而非全局弱引用注册表，并提供了 FlashMLA H20 的对比数据（+0.41% throughput, -1.73% TPOT）。zhou9402 认同该方案更清晰，但认为不应在当前 PR 中扩大范围，建议作为后续 follow-up。
- 性能验证要求 (correctness): 数据充分，验证了优化收益和准确率保持
- PR 规模与拆分 (design): PR 被 revert，优化通过拆分后的 PR 逐步引入 main
- Revert 决定 (other): PR 被 revert，后续拆分 PR 需要更充分的 review
- `phys_shadow` 设计讨论 (design): 承认 `metadata` 方案更好，但不在当前 PR 中实现，留作后续改进

## 风险与影响

- 风险：原始 PR 的主要技术风险包括：
- 架构局限性：优化仅针对 Blackwell SM100 (CUDA capability 10.0)，在其他平台（H100、AMD、Intel GPU）会 fallback 到通用路径。但 fallback 的 `deepseek_v32` 包入口在非 SM100 上直接 raise `NotImplementedError`，可能导致模型加载失败（被后续提交 #5a1043e 修正为 fallback 到 generic `deepseek_v2`）。
- 编译兼容性：CuTeDSL 依赖增加了编译复杂度；`torch.compile` 在非 SM100 上可能因 `decompose_triton_kernel_wrapper_functional` 断言而失败，因此需要架构守卫。
- 设计争议：全局 `phys_shadow` 弱引用注册表的使用被 review 者质疑，认为应存放在 `attention metadata` 中以减少耦合。
- 回归风险：MTP 索引共享的逻辑（`set_skip_topk`）如果与 future 的模型不兼容，可能静默产生错误结果。测试覆盖了核心路径但可能不全面。

- 性能反优化: `bf16_skinny_gemm` 的阈值依赖于实测 tuning, 如果未来模型 shape 变化或驱动更新, 可能选择次优路径。当前仅在 GLM-5.2 和 DSv3.2 上验证了特定 shape。

这些风险导致 PR 被 revert, 但后续拆分 PR 通过逐步审查和更充分的测试缓解了风险。

- 影响: 对用户的影响: 最终用户在 Blackwell GPU 上运行 GLM-5.2 或 DeepSeek-V3.2 模型时, 将获得约 21% 的解码吞吐提升和更低的 TPOT。但由于 PR 被 revert, 这些收益是通过后续拆分的 PR (如 #49790-#49793) 逐步以更可控的方式引入。用户需要关注每个拆分 PR 的发布说明。

对系统的影响: 引入了 CuTeDSL 作为可选编译依赖 (仅 SM100); 新增了

`fused_q_cutedsl` 和 `bf16_skinny_gemm` 两个自定义 ops; 在 `vllm._custom_ops` 中注册了新的 CUDA 算子。`deepseek_v32` 包成为 GLM-5.2 和 DSv3.2 在 SM100 上的主要路径。

对团队的影响: 该 PR 的拆分模式成为后续大型 PR 的范例——先集成跟踪 PR, 再拆分为多个 <500 LOC 的聚焦 PR。review 过程强调了早期性能验证和设计决策的提前沟通。

- 风险标记: 核心路径变更, 仅限 Blackwell, 设计争议, 大规模变更, 已被 revert

## 关联脉络

- PR #49768 Revert "[Perf][GLM-5.2] Blackwell decode optimizations": 直接 revert 了本 PR
- PR #49790 SM100 sparse-model integration and routing: 本 PR 拆分后的 follow-up, 负责 sparse model 路由
- PR #49791 Small-batch decode GEMM optimizations: 本 PR 拆分后的 follow-up, 负责小批量 GEMM 优化
- PR #49792 CuTeDSL fused-query kernel: 本 PR 拆分后的 follow-up, 负责 CuTeDSL 融合查询核
- PR #50230 Programmatic dependent launch for the decode kernels: 本 PR 拆分后的 follow-up, 负责 PDL 程序化依赖启动
- PR #49793 MTP/speculative-decoding optimizations: 本 PR 拆分后的 follow-up, 负责 MTP 投机解码优化
- PR #48335 FP32 router GEMM shape support and tuning: 本 PR 的组成部分, 已合并到 main
- PR #47973 BF16x3 router GEMM: 本 PR 的组成部分, 已合并到 main
- PR #47970 Preserve FP32 router weights for the relevant model family: 本 PR 的组成部分, 已合并到 main
- PR #45895 Generic MTP TopK index sharing and compaction 1/2: 本 PR 的组成部分, 已合并到 main
- PR #47238 Generic MTP TopK index sharing and compaction 2/2: 本 PR 的组成部分, 已合并到 main
- PR #35161 Fix expert\_ids padding values in moe\_align\_block\_size kernel: 本 PR 的组成部分 (MoE 对齐填充修复), 已合并到 main
- PR #46661 FlashInfer A2A support: 本 PR 的组成部分, 已合并到 main



- PR #49678 Store sparse physical indices in attention metadata: 与本 PR 的 phys\_shadow 设计讨论相关, 建议 metadata 方案