

PR #32327 完整报告

sgl-project/sglang

[DeepSeek-V4] Add Q8KV8 sparse MLA prefill runtime backend

合并时间: 2026-08-20 10:23

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/32327>

执行摘要

- 一句话: DeepSeek-V4 新增 Q8KV8 FP8 稀疏 MLA prefill 运行时后端
- 推荐动作: 值得精读。关注四个设计点: 一是「先内核后运行时」的两段式 PR 拆分 (#25751 → #32327), 降低大型特性 review 负担; 二是 fail-fast 输入契约校验范式, 把 CUDA 层的裸指针契约显式化到 Python 封装层, 有效避免静默垃圾输出或进程被 `exit(1)` 杀死; 三是 TP 场景下 Q 头按 64 头 CTA 粒度 padding + identity scale 屏蔽填充头的做法; 四是 dtype 感知的 workspace 复用策略, 兼顾 BF16/FP8 两路径与 CUDA graph 捕获。

功能与动机

这是 issue #25746 路线图的 PR2 (运行时集成部分)。issue 明确提出: NSA 风格稀疏 MLA prefill 是长上下文服务的性能关键路径, memory bandwidth 与 KV cache 访问模式是主要瓶颈, Q8KV8 FP8 设计可显著降低内存流量, 但当时缺少面向 NSA 风格 prefill 的专用稀疏 FP8 内核。PR body 补充了运行时侧的动机: DeepSeek-V4 已支持 FP8 KV cache, 但既有稀疏 prefill 后端仍主要走 BF16 稀疏 MLA 路径, 导致即使 KV 以 FP8 存储, 运行时也无法利用原生 FP8 Q8KV8 稀疏 prefill 内核的张量核计算能力。

实现拆解

1. 新增 CLI 配置入口 (server_args.py)

新增 `DSV4_PREFILL_BACKEND_CHOICES = ["auto", "flashmla_sparse", "flashmla_sparse_q8"]` 与 `dsv4_prefill_backend` ServerArgs 字段 (默认 `auto`), 通过 `--dsv4-prefill-backend` 暴露。默认值与 `flashmla_sparse` 都走既有 BF16 稀疏 prefill 路径, 只有 `flashmla_sparse_q8` 才启用 Q8KV8, 保证默认行为零变化。

2. 后端选择与 workspace 改造 (sparse_prefill_utils.py)

新增 `use_dsv4_q8kv8_sparse_prefill()` 选择函数: 以 CLI 配置为准, 同时保留 `SGLANG_DSV4_Q8KV8_PREFILL` 环境变量作为调试期强制开关 (`truthy/falsy` 直接覆盖, 便于运行时路径硬化前快速切换)。 `SparsePrefillWorkspace.get()` 增加 `dtype` 参数, 使 BF16 与 FP8 两种 workspace 复用同一分配逻辑, 仅在容量或 dtype 不匹配时才重新分配——对 CUDA graph 捕获期避免反复申请显存很关键。

3. 内核输入适配 (dequant_k_cache.py、sparse_mla_q8kv8_prefill_sm90.py)

`dequant_k_cache.py` 新增三个核心符号：`gather_dequant_requant_fp8_paged` 把 DeepSeek-V4 分页 KV cache (448 维 nope 为 UE8M0 量化 FP8 + 64 维 BF16 rope) 一次性完成 gather、反量化、重量化并输出连续 FP8 workspace, `extra_rows` 支持为越界稀疏索引提供零值落地区；`q8kv8_padded_num_heads` 把 TP 本地 Q 头数向上取整到 SM90 内核的 64 头 CTA 粒度 (64/128)；`cast_q_fp8_for_q8kv8_prefill` 完成 Q 的 FP8 cast 与头部填充。`sparse_mla_q8kv8_prefill_sm90.py` 为内核封装层补上启动前的 fail-fast 输入契约校验 (rank、dtype、CUDA 设备、连续性、scale 张量、`topk_length` 范围)。

4. DeepSeek-V4 运行时集成 (`deepseek_v4_backend.py`)

`DeepseekV4AttnBackend.__init__` 在启用 Q8KV8 时立即校验 `is_sm90_supported()` 与 `head_dim_v == 512`, 不满足直接抛错, 并初始化 `_q8kv8_qpad_buf`、`_q8kv8_attn_sink_pad`、`_q8kv8_identity_scale` 三个可复用缓冲区。稀疏 prefill 主路径上 `use_dsv4_q8kv8_sparse_prefill()` 命中时转入新增的 `_forward_prefill_sparse_q8kv8()`, 内部先由 `_prepare_q8kv8_q_and_sink()` 完成 Q 头 padding 与 attention sink padding, 再按 C0/C4/C128 三种压缩比 dispatch 到对应 FP8 稀疏 prefill 分支。

5. 测试配套

新增 `test/registered/kernels/ops/attention/test_q8kv8_sparse_prefill_backend.py` (68行) : 不启动完整 server, 构造最小 V4 metadata 与 token pool 表面, 把 BF16 路径 gather 出的 workspace 与 Q8 路径 FP8 workspace 反量化后逐元素对比; 覆盖 `compress_ratio` 0/4/128 三种分支, 并加入 `s_q=512`、`topk=256` 重复 10 次的跨迭代竞态回归。`test_server_args.py` 补充 `--dsv4-prefill-backend` 各取值与校验路径的 CLI 测试。测试注册到 CUDA CI (base-b-kernel-unit, 1-gpu-large)。

关键文件:

- `python/sglang/srt/layers/attention/deepseek_v4_backend.py` (模块 注意力后端; 类别 source; 类型 core-logic; 符号 `_prepare_q8kv8_q_and_sink`, `_forward_prefill_sparse_q8kv8`) : Q8KV8 运行时集成主入口: 新增 `_forward_prefill_sparse_q8kv8` 与 `_prepare_q8kv8_q_and_sink`, 负责后端路由、SM90/d_v 校验、Q 头与 attention sink padding、C0/C4/C128 dispatch, 是本 PR 的核心改动文件。
- `test/registered/kernels/ops/attention/test_q8kv8_sparse_prefill_backend.py` (模块 稀疏预填充; 类别 test; 类型 test-coverage; 符号 `test_q8kv8_sparse_prefill_backend_selector_uses_cli_value`, `_Pool`, `_Capture`, `_TokenToKVPool`) : 681 行新增测试, 不启动完整 server 即可覆盖 C0/C4/C128 三种压缩比、BF16 workspace 与 Q8 FP8 workspace 的数值一致性, 以及跨迭代竞态回归, 是本次 review 中测试缺口闭环的核心。
- `python/sglang/srt/layers/attention/dsv4/sparse_prefill_utils.py` (模块 预填充工具; 类别 source; 类型 core-logic; 符号 `use_dsv4_q8kv8_sparse_prefill`, `get`) : 新增 `use_dsv4_q8kv8_sparse_prefill` 后端选择函数 (含环境变量 debug 覆盖), 并让 `SparsePrefillWorkspace.get` 支持 dtype 感知分配, 是 FP8/BF16 两路径共用的关键工具层。
- `python/sglang/kernels/ops/attention/dsv4/dequant_k_cache.py` (模块 缓存量化; 类别 source; 类型 infrastructure; 符号 `gather_dequant_requant_fp8_paged`,

q8kv8_padded_num_heads, cast_q_fp8_for_q8kv8_prefill, _gather_dequant_requant_fp8_paged_kernel) : 新增gather_dequant_requant_fp8_paged 融合内核: 把分页 KV 一次完成 gather + 反量化 + 重量化为连续 FP8, 是 Q8KV8 prefill 的输入数据通路核心。

- python/sglang/kernels/ops/attention/sparse_mla_q8kv8_prefill_sm90.py (模块 内核封装; 类别 source; 类型 infrastructure; 符号 sparse_mla_q8kv8_prefill_fwd) : 给 SM90 Q8KV8 内核封装层补上 fail-fast 输入契约校验, 直接回应 review 中 topk 对齐与 tensor 契约两条意见, 防止进程被内核 exit(1) 杀死或产生静默垃圾输出。
- python/sglang/srt/server_args.py (模块 服务参数; 类别 source; 类型 configuration) : 新增 dsv4_prefill_backend ServerArgs 字段与 DSV4_PREFILL_BACKEND_CHOICES, 是 Q8KV8 功能的 CLI 入口, 默认 auto 保持既有行为。
- test/registered/unit/server_args/test_server_args.py (模块 参数测试; 类别 test; 类型 test-coverage; 符号 test_dsv4_prefill_backend_cli_choices) : 补充 --dsv4-prefill-backend 的 CLI 取值测试, 验证三个合法值与参数解析行为, 是配置入口的回归保障。

关键符号: _forward_prefill_sparse_q8kv8, _prepare_q8kv8_q_and_sink, use_dsv4_q8kv8_sparse_prefill, gather_dequant_requant_fp8_paged, q8kv8_padded_num_heads, cast_q_fp8_for_q8kv8_prefill, sparse_mla_q8kv8_prefill_fwd, test_q8kv8_sparse_prefill_backend_selector_uses_cli_value, test_dsv4_prefill_backend_cli_choices

关键源码片段

python/sglang/srt/layers/attention/deepseek_v4_backend.py

Q8KV8 运行时集成主入口: 新增 `_forward_prefill_sparse_q8kv8` 与 `_prepare_q8kv8_q_and_sink`, 负责后端路由、SM90/d_v 校验、Q 头与 attention sink padding、C0/C4/C128 dispatch, 是本 PR 的核心改动文件。

```
# deepseek_v4_backend.py: Q8KV8 运行时入口与前置校验
# 在 __init__ 中一旦用户显式选择 flashmla_sparse_q8, 立即校验硬件与模型约束,
# 避免问题拖到 prefill 阶段才暴露 (届时只会表现为挂起或集合通信失败)。
self.dsv4_prefill_backend = getattr(
    model_runner.server_args, "dsv4_prefill_backend", "auto"
)

if use_dsv4_q8kv8_sparse_prefill(self.dsv4_prefill_backend):
    # SM90 Q8KV8 内核是 Hopper 专用实现, 非 SM90 直接拒绝启动
    if not is_sm90_supported():
        raise ValueError(
            "DeepSeek-V4 flashmla_sparse_q8 prefill requires SM90 CUDA GPUs."
        )
    # 内核按 d_v=512 的 MLA 布局编写, d_v 不匹配时产出错误结果而非报错
    if self.head_dim_v != 512:
        raise ValueError(
            "DeepSeek-V4 flashmla_sparse_q8 prefill requires d_v=512, "
            f"got {self.head_dim_v}."
        )
```

```

# 预分配可复用缓冲区: Q 填充 / attention sink 填充 / 恒等 scale,
# 供 CUDA graph 捕获期复用, 避免反复申请显存
self._q8kv8_qpad_buf = None
self._q8kv8_attn_sink_pad = None
self._q8kv8_identity_scale = None

# 稀疏 prefill 主路径上, 命中 Q8KV8 时切换到 FP8 内核,
# 否则回落原有 BF16 flashmla_sparse 路径 (默认行为保持不变)
if use_dsv4_q8kv8_sparse_prefill(self.dsv4_prefill_backend):
    return self._forward_prefill_sparse_q8kv8(
        q=q,
        layer_id=layer_id,
        compress_ratio=compress_ratio,
        forward_batch=forward_batch,
        token_to_kv_pool=token_to_kv_pool,
        core_attn_metadata=core_attn_metadata,
        attn_sink=attn_sink,
    )

```

python/sglang/srt/layers/attention/dsv4/sparse_prefill_utils.py

新增 `use_dsv4_q8kv8_sparse_prefill` 后端选择函数 (含环境变量 debug 覆盖), 并让 `SparsePrefillWorkspace.get` 支持 dtype 感知分配, 是 FP8/BF16 两路径共用的关键工具层。

```

# sparse_prefill_utils.py: Q8KV8 后端选择与 dtype 感知的 workspace 复用

# 生产配置以 --dsv4-prefill-backend 为准; 环境变量作为调试期强制开关,
# 在运行时路径尚未完全硬化时允许快速切换, 无需改启动参数。
DSV4_Q8KV8_PREFILL_ENV = "SGLANG_DSV4_Q8KV8_PREFILL"

```

```

def use_dsv4_q8kv8_sparse_prefill(dsv4_prefill_backend: str = "auto") -> bool:
    """返回 DeepSeek-V4 sparse prefill 是否使用 Q8KV8。"""
    env_value = os.getenv(DSV4_Q8KV8_PREFILL_ENV)
    if env_value is not None:
        # 显式环境变量优先: truthy 强制开启, falsy 强制关闭
        return env_value.lower() in {"1", "true", "yes", "on"}
    return dsv4_prefill_backend == "flashmla_sparse_q8"

```

```

class SparsePrefillWorkspace:

```

```

    """Backend 自持的 sparse prefill KV 暂存区。

```

```

    workspace 内容在每次 attention 前被完全覆盖, 因此不同 token 桶与
    压缩比可以安全共享同一块 buffer; sparse prefill 在受支持的路径上
    是 eager 串行执行的, 需要更大容量时替换分配也是安全的。
    """

```

```

    def __init__(self, device: torch.device):
        self.device = device
        self._buffer: Optional[torch.Tensor] = None

```

```

def get(
    self,
    num_tokens: int,
    dtype: torch.dtype = torch.bfloat16,
) -> torch.Tensor:
    assert num_tokens > 0
    current_capacity = self._buffer.shape[0] if self._buffer is not None else 0
    current_dtype = self._buffer.dtype if self._buffer is not None else None
    # Q8KV8 需要 FP8 workspace (容量减半), BF16 路径仍用原 dtype;
    # 仅当容量或 dtype 不匹配时才重新分配, 避免 CUDA graph 捕获期反复申请显存
    if num_tokens > current_capacity or dtype != current_dtype:
        self._buffer = torch.empty(
            (num_tokens, 1, WORKSPACE_DIM),
            dtype=dtype,
            device=self.device,
        )
    return self._buffer[:num_tokens]

```

python/sglang/kernels/ops/attention/dsv4/dequant_k_cache.py

新增 `gather_dequant_requant_fp8_paged` 融合内核: 把分页 KV 一次完成 gather + 反量化 + 重量化为连续 FP8, 是 Q8KV8 prefill 的输入数据通路核心。

```

# dequant_k_cache.py: 分页 KV 的 gather + 反量化 + 重量化 FP8 融合适配
# paged 缓存布局: 448 维 nope 为 UE8M0 量化的 FP8, 尾部 64 维 rope 为 BF16。
# Q8KV8 sparse prefill 内核要求 kv 输入为连续 FP8, 因此不能直接复用
# BF16 路径的 dequantize_k_cache_paged, 这里在 gather 时同步完成
# 反量化 (nope 区域按每 64 元素一个 scale) + rope 段 BF16→FP8 重量化。

```

```

def gather_dequant_requant_fp8_paged(
    quant_k_cache: torch.Tensor,
    page_table_1_flattened: torch.Tensor,
    page_size: int,
    extra_rows: int = 0,
    out: Optional[torch.Tensor] = None,
) -> torch.Tensor:
    assert quant_k_cache.is_contiguous()
    assert page_table_1_flattened.dtype in (torch.int32, torch.int64)
    assert extra_rows >= 0

```

```

    quant_k_cache_u8 = quant_k_cache.view(torch.uint8)
    num_tokens = page_table_1_flattened.shape[0]
    total_rows = num_tokens + extra_rows
    bytes_per_page = quant_k_cache_u8.shape[-1]
    s_offset_bytes = page_size * NOPE_ROPE_BYTES

```

```

    # 同一段 uint8 内存按 FP8 / BF16 / uint8 三种视图复用,
    # 让 Triton 内核按区域语义直接读写, 避免数据搬移
    buf_fp8 = quant_k_cache_u8.view(fp8_dtype).reshape(-1)
    buf_bf16 = quant_k_cache_u8.view(torch.bfloat16).reshape(-1)
    buf_uint8 = quant_k_cache_u8.reshape(-1)

```

```

if out is None:
    out = torch.zeros(
        (total_rows, 1, DIM_NOPE + DIM_ROPE),
        dtype=fp8_dtype,
        device=quant_k_cache.device,
    )
else:
    assert out.shape == (total_rows, 1, DIM_NOPE + DIM_ROPE)
    assert out.dtype == fp8_dtype
    if extra_rows:
        # extra_rows 是内核把越界稀疏索引映射到的零值落地区
        out[num_tokens:].zero_()

if num_tokens == 0:
    return out

# 单 token 一个线程的 gather 内核，按常量参数编译，
# 规避动态 shape 带来的 JIT 重编译与索引计算开销
_gather_dequant_requant_fp8_paged_kernel[(num_tokens,)](
    out,
    buf_fp8,
    buf_bf16,
    buf_uint8,
    page_table_1_flattened,
    out.stride(0),
    BYTES_PER_PAGE=bytes_per_page,
    PAGE_SIZE=page_size,
    DIM_NOPE=DIM_NOPE,
    DIM_ROPE=DIM_ROPE,
    TILE_SIZE=TILE_SIZE,
    NUM_SCALE_TILES=NUM_SCALE_TILES,
    NOPE_ROPE_BYTES=NOPE_ROPE_BYTES,
    PADDED_SCALE_PER_TOKEN=PADDED_SCALE_PER_TOKEN,
    S_OFFSET_BYTES=s_offset_bytes,
)
return out

```

评论区精华

三处 review 意见全部闭环，焦点集中在「内核启动前的 fail-fast 校验」与「测试覆盖缺口」：

- topk 对齐收紧：BBuf 在 H200 上实测发现 topk=64 能通过 Python 侧 %64 校验，但 CUDA launcher 要求 `params.topk % (2 * B_TOPK) == 0` (B_TOPK=64)，随后 KU_ASSERT 触发 `exit(1)` 直接杀死进程。作者将校验改为 %128，让 topk=64 在 launch 前 fail-fast，并新增拒绝路径的校验测试。
- 输入契约校验：BBuf 指出 `entry.cuh` 把 q/kv 当作连续 FP8 裸指针直接 launch，dtype、连续性、设备与 scale 张量均未检查，坏调用方会得到静默垃圾或非法访问；同时建议守卫

0 <= topk_length <= topk。作者补齐 rank、dtype、设备、连续性、scale 张量、attention sink 契约与 topk_length 范围的全套校验。

- C4/C128 与竞态回归：BBuf 发现 backend 测试全部只传 `compress_ratio=0`，C4/C128 的独立 gather/index 路径未被覆盖；并指出该内核修复本身涉及跨迭代竞态，建议加 `s_q>=512`、`topk>=256` 的重复启动用例，还分享了自己在 H200 上跑 512/256 十次稳定的经验。作者两处都已补齐。
 - topk 对齐校验从 %64 收紧为 %128 (correctness): 作者将 Python 侧校验改为 `topk % 128 == 0`，使 `topk=64` 在 launch 前 fail-fast，并新增覆盖该拒绝路径的校验测试。
 - 内核输入 tensor 契约 fail-fast 校验 (correctness): 作者补齐 q/kv/indices 的 rank、dtype、CUDA 设备、同设备、连续性校验，`q_scale/kv_scale` 张量类型与形状校验，`attn_sink` 契约校验，以及 `topk_length` 范围校验。
 - C4/C128 分支与跨迭代竞态回归测试 (testing): 作者补充 `compress_ratio=4` 与 128 的 backend 覆盖，并新增 `s_q=512`、`topk=256` 重复 10 次的跨迭代竞态回归测试。

风险与影响

- 风险：
 1. 硬件约束：flashmla_sparse_q8 仅在 SM90 上可用，且要求 `d_v=512`；虽在 `__init__` 阶段 fail-fast 拒绝，但用户若在非 Hopper 环境显式指定会直接启动失败，需文档明确。
 2. 核心 attention 路径变更：deepseek_v4_backend.py 是 DeepSeek-V4 稀疏 prefill 的关键路径，Q 头 padding 与 identity scale 若出差错会静默产出错误注意力输出（不会崩溃），依赖新增单测与真实 serving shape 回归兜底。
 3. 跨迭代竞态：FP8 workspace 复用 + CUDA graph 捕获场景下存在跨迭代 buffer 复用竞态，已在 `s_q=512`、`topk=256` 重复 10 次的用例中覆盖，但 CI 只在 1 卡 SM90 上跑。
 4. 兼容性：sparse_mla_q8kv8_prefill_sm90.py 收紧校验后，任何未按 64 头粒度 padding 的调用方都会显式报错，可能影响其他复用该封装的上游调用，需要同步升级。
 5. 精度：FP8 量化引入额外精度损失，虽 GSM8K 与 LongBench-v2 无实质回退，但其他长上下文任务仍需用户侧验证。
 - 影响：
 1. 用户：DeepSeek-V4 (Flash) 长上下文 serving 用户可通过 `--dsv4-prefill-backend flashmla_sparse_q8` 获得 4%-8.5% 的 TTFT/ 输入吞吐提升（H20 实测，chunk 8192/16384、c=1/c=16 全场景一致正向），且默认路径不变。
 2. 系统：改动横跨 server_args、注意力后端、JIT kernel 封装与 dequant 内核，新增约 1.3k 行，但入口显式 opt-in，默认行为零变化，风险面可控。
 3. 团队：作为 issue #25746 路线图 PR2，与 PR1 (#25751 独立内核) 构成 kernel → runtime 的两段式落地范式；后续 GLM 等 NSA 风格模型可复用该运行时适配模式。
 4. 维护：fail-fast 校验范式和 dtype 感知 workspace 复用的引入，提升了 JIT kernel 封装层的可调用安全性，也为后续 Q8KV8 下沉到 decode 或其他稀疏注意力模型铺路。
- 风险标记：核心 attention 路径变更，SM90 硬件约束，新增 CLI 配置入口，跨迭代竞态风险，CUDA graph 捕获兼容

关联脉络

- 暂无明显关联 PR