

PR #40177 完整报告

vllm-project/vllm

Add nvfp4 kv cache support

合并时间: 2026-05-01 12:55

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/40177>

执行摘要

- 一句话: FlashInfer 后端添加 NVFP4 KV cache 端到端支持
- 推荐动作: 值得精读, 尤其对于关注 NVIDIA Blackwell 量化和 FlashInfer 集成的研究者。注意 PR 中的 NVFP4 输出处理 (FP8 缓冲区和反量化) 是关键权衡设计; 文档生成脚本的修改也值得借鉴。需关注后续精度报告的补充。

功能与动机

这是对 PR #37332 (引入 NVFP4 KV cache 基础设施) 的延续, 目标是完全启用 NVFP4 KV cache, 使得在 Blackwell GPU 上可以通过 FlashInfer TRTLLM 后端进行端到端 NVFP4 KV cache 推理, 从而降低显存开销并提升吞吐量。

实现拆解

1. FlashInfer 后端注册 NVFP4: 在 `vllm/v1/attention/backends/flashinfer.py` 中将 "nvfp4" 加入 `supported_kv_cache_dtypes` 列表; 新增 `get_dtype_for_flashinfer` 方法 (从 `get_fp8_dtype_for_flashinfer` 重命名) 处理 `nvfp4` $\rightarrow$ `torch.uint8` 映射; 在 `_get_prefill_wrapper` 和 `_get_decode_wrapper` 中强制使用 `backend="trtllm-gen"` 以绕过 FA2/FA3 的限制。
2. FP8 输出处理: NVFP4 kernel 只输出 FP8, 因此在 `FlashInferImpl` 的 `__init__` 中预分配 FP8 缓冲区 `_nvfp4_fp8_out`; 在 `build` 方法中将 `o_dtype` 设为 FP8, 并在 `forward` 中将输出反量化回模型 dtype。
3. Block scale 传递: 通过 `nvfp4_kv_cache_split_views` 从打包的 NVFP4 cache 中分离出 FP8 block scale, 并作为 `kv_cache_sf` 参数传入 wrapper 的 `run()` 及 `trtllm_batch_decode_with_kv_cache`。
4. ModelOpt 与配置更新: 在 `vllm/model_executor/layers/quantization/modelopt.py` 中将 `KV_CACHE_QUANT_ALGOS` 扩展为 `["FP8", "NVFP4"]`, 并将类 `ModelOptFp8KVCacheMethod` 重命名为 `ModelOptKVCacheMethod`, 统一处理 FP8 和 NVFP4 的 scale 加载。
5. 兼容性校验与文档: 在 `vllm/config/vllm.py` 中添加 `validate_nvfp4_kv_cache_with_mla` 校验, 防止 NVFP4 与 MLA 后端共用; 更新 `tools/pre_commit/generate_attention_backend_docs.py`, 在非 SM100 的 native FlashInfer 列表中排除 `nvfp4`。

关键文件:

- `vllm/v1/attention/backends/flashinfer.py` (模块 注意力后端; 类别 `source`; 类型 `core-logic`; 符号 `get_fp8_dtype_for_flashinfer`, `get_dtype_for_flashinfer`, `supported_kv_cache_dtypes`) : 核心实现: 注册 `nvfp4` dtype、强制 `trtllm-gen` 后端、FP8 输出处理、`block scale` 传递。
- `tests/v1/attention/test_trtllm_attention_integration.py` (模块 集成测试; 类别 `test`; 类型 `test-coverage`; 符号 `_run_trtllm_integration`, `_create_nvfp4_hnd_kv_cache`, `test_trtllm_gen_nvfp4_kv_integration`) : 集成测试: 验证完整 `FlashInferImpl` + `MetadataBuilder` 流程与 `nvfp4 KV cache` 的正确性。
- `tests/kernels/attention/test_flashinfer_trtllm_attention.py` (模块 单元测试; 类别 `test`; 类型 `test-coverage`; 符号 `build_paged_kv_metadata`, `make_nvfp4_kv_cache`, `make_quantized_kv_cache`) : 单元测试: 重构辅助函数, 新增 `make_nvfp4_kv_cache` 和 `build_paged_kv_metadata`, 支持 `nvfp4` 参数化。
- `vllm/model_executor/layers/quantization/modelopt.py` (模块 量化配置; 类别 `source`; 类型 `data-contract`; 符号 `ModelOptFp8KVCacheMethod`, `ModelOptKVCacheMethod`, `KV_CACHE_QUANT_ALGOS`) : 数据契约: 扩展 `KV_CACHE_QUANT_ALGOS` 并统一 `KVCacheMethod` 类以支持 `NVFP4` 检查点。
- `vllm/config/vllm.py` (模块 配置校验; 类别 `source`; 类型 `core-logic`; 符号 `validate_nvfp4_kv_cache_with_mla`) : 配置校验: 添加 `nvfp4` 与 `MLA` 的兼容性检查, 防止非法组合。
- `tools/pre_commit/generate_attention_backend_docs.py` (模块 文档生成; 类别 `source`; 类型 `core-logic`) : 文档生成: 确保 `nvfp4` 仅出现在 `TRTLLM` 后端特性表中, 避免误导用户。

关键符号: `get_dtype_for_flashinfer`, `_get_prefill_wrapper`, `_get_decode_wrapper`, `build`, `forward`, `validate_nvfp4_kv_cache_with_mla`, `ModelOptKVCacheMethod.init`, `make_nvfp4_kv_cache`, `_create_nvfp4_hnd_kv_cache`, `test_trtllm_gen_nvfp4_kv_integration`

关键源码片段

`vllm/v1/attention/backends/flashinfer.py`

核心实现: 注册 `nvfp4` dtype、强制 `trtllm-gen` 后端、FP8 输出处理、`block scale` 传递。

```
class FlashInferBackend(AttentionBackend):
    # 支持的 KV cache dtype 列表, 新增 "nvfp4" 支持
    supported_kv_cache_dtypes: ClassVar[list[CacheDType]] = [
        "auto",
        "float16",
        "bfloat16",
        "fp8",
        "fp8_e4m3",
        "fp8_e5m2",
        "nvfp4", # NVFP4 为 4-bit 量化格式, 仅支持 SM100
    ]

    @staticmethod
```

```

def get_dtype_for_flashinfer(kv_cache_dtype: str) -> torch.dtype:
    # 将 KV cache dtype 字符串映射为 PyTorch 数据类型
    if kv_cache_dtype in ("fp8", "fp8_e4m3"):
        return torch.float8_e4m3fn
    elif kv_cache_dtype == "fp8_e5m2":
        return torch.float8_e5m2
    elif kv_cache_dtype == "nvfp4":
        # NVFP4 使用 uint8 存储打包后的数据 (包括 FP4 数据和 FP8 block scale)
        return torch.uint8
    else:
        raise ValueError(f"Unrecognized dtype: {kv_cache_dtype}")

@staticmethod
def get_kv_cache_shape(
    num_blocks: int, block_size: int,
    num_kv_heads: int, head_size: int,
    cache_dtype_str: str = "auto",
) -> tuple[int, ...]:
    # NVFP4 使用特殊打包维度: head_size/2 + head_size/16
    if cache_dtype_str == "nvfp4":
        last_dim = nvfp4_kv_cache_full_dim(head_size)
        return (num_blocks, 2, block_size, num_kv_heads, last_dim)
    return (num_blocks, 2, block_size, num_kv_heads, head_size)

```

tests/v1/attention/test_trtllm_attention_integration.py

集成测试: 验证完整 FlashInferImpl + MetadataBuilder 流程与 nvfp4 KV cache 的正确性。

```

def _create_nvfp4_hnd_kv_cache(
    k_contexts, v_contexts, block_size, num_kv_heads, head_size,
    dtype, device, num_blocks, common_attn_metadata, kv_scale_val,
):
    """通过 reshape_and_cache_flash 将 bf16 context 量化为 nvfp4
    KV cache, 并返回符合 HND stride order 的 uint8 tensor.
    """
    # 先创建 bf16 的 HND cache 以填充 block table
    bf16_cache = _create_hnd_kv_cache(
        k_contexts, v_contexts, block_size, num_kv_heads,
        head_size, dtype, device, num_blocks, common_attn_metadata,
    )

    # 分配 nvfp4 cache: 维度从 head_size 变为 full_dim
    full_dim = nvfp4_kv_cache_full_dim(head_size)
    hnd_order = (0, 1, 3, 2, 4)
    nvfp4_cache = torch.zeros(
        (num_blocks, 2, num_kv_heads, block_size, full_dim),
        dtype=torch.uint8, device=device,
    ).permute(*hnd_order) # 转换为 NHD 顺序以便 reshape_and_cache_flash 使用

    # 展开为 token 维度 [N*T, H, D] 并调用量化 kernel

```

```

num_tokens = num_blocks * block_size
k_tokens = (bf16_cache[:, 0].permute(0, 2, 1, 3)
            .reshape(num_tokens, num_kv_heads, head_size))
v_tokens = (bf16_cache[:, 1].permute(0, 2, 1, 3)
            .reshape(num_tokens, num_kv_heads, head_size))
slot_mapping = torch.arange(num_tokens, dtype=torch.long, device=device)
torch.ops._C_cache_ops.reshape_and_cache_flash(
    k_tokens, v_tokens,
    nvfp4_cache[:, 0], nvfp4_cache[:, 1],
    slot_mapping, "nvfp4",
    kv_scale_val, kv_scale_val,
)
# 转回 HND 顺序 (trtllm kernel 期望的物理布局)
return nvfp4_cache.permute(*hnd_order)

```

vllm/model_executor/layers/quantization/modelopt.py

数据契约：扩展 KV_CACHE_QUANT_ALGOS 并统一 KVCacheMethod 类以支持 NVFP4 检查点。

```

# 扩展 KV cache 量化算法列表，新增 NVFP4 支持
KV_CACHE_QUANT_ALGOS = ["FP8", "NVFP4"]

```

```

class ModelOptKVCacheMethod(BaseKVCacheMethod):
    """
    支持从 FP8 或 NVFP4 检查点加载 KV cache 缩放因子。
    """

    def __init__(self, quant_config: "ModelOptQuantConfigBase"):
        super().__init__(quant_config)

# 统一所有 ModelOpt 配置类的 KVCacheMethodCls
ModelOptFp8Config.KVCacheMethodCls = ModelOptKVCacheMethod
ModelOptNvFp4Config.KVCacheMethodCls = ModelOptKVCacheMethod
ModelOptMxFp8Config.KVCacheMethodCls = ModelOptKVCacheMethod

```

评论区精华

- 性能对比(@vadiklyutiy 请求)：作者提供了 Qwen3-8B 在 B100 上的测试结果，显示大 batch 下 NVFP4 的 TPOT 明显优于 FP8 (bs=256 时 1.24 ms vs 1.55 ms)。
- 精度评估(@pavanmajety 要求 GSM8K 测试)：作者运行了 GSM8K 评估，并对 Qwen3-8B 报告了精度结果；但指出 Qwen3.5 模型存在精度问题待调查。
- SM120 兼容性(@wangqia0309 提问)：作者说明 NVFP4 KV cache 的架构支持取决于 FlashInfer，目前仅限 SM100。
- MLA 检查(@pavanmajety 要求)：作者添加了 validate_nvfp4_kv_cache_with_mla 验证，当 NVFP4 与 MLA 一起使用时抛出 ValueError，该请求已解决。
- FP8 缓冲分配位置(@mgoin 指出)：初始版本在 forward 中分配，后通过 commit 改为在 init 时分配，避免运行时开销。

- 文档生成(@mgoin 与 @pavanimajety 反馈): nvfp4 仅应在 TRTLLM 后端表中展示, 作者修复了文档生成脚本。
- 性能测试 (FP4 vs FP8) (performance): 作者提供了 Qwen3-8B 在 B100 上的 TPOT 比较: 大 batch (bs=256) 时 nvfp4 (1.24 ms) 优于 fp8 (1.55 ms), 小 batch 相近。
- 精度验证 (GSM8K evals) (testing): 部分解决: 作者提供了部分精度数据, 但 Qwen3.5 的精度问题被标记为待调查。
- SM120 架构支持 (question): 已确认当前仅限于 SM100。
- MLA 兼容性检查 (correctness): 已解决: 添加了 ValueError 校验。
- FP8 输出缓冲区分配时机 (performance): 已解决: 通过 commit 2e9da76 将分配移到 `__init__` 中。
- 文档中 nvfp4 的展示范围 (documentation): 已解决: 修改了 `generate_attention_backend_docs.py`, 在 `native` 列表中排除 nvfp4。

风险与影响

- 风险:
 1. 兼容性风险: NVFP4 仅受 FlashInfer TRTLLM 后端支持且需要 Blackwell SM100, 在非 SM100 GPU 上无法使用但会优雅降级 (不选择 nvfp4 即可)。
 2. 精度风险: NVFP4 为 4-bit 量化, 精度低于 FP8, 且 kernel 输出 FP8 后反量化可能引入额外舍入误差; 作者添加了 unittest 与 bf16 对比, 但覆盖有限。
 3. MLA 冲突: 已通过验证阻止, 但未来若 MLA 后端适配 NVFP4 需移除该检查。
 4. 依赖风险: 依赖 FlashInfer 的 `reshape_and_cache_flash` kernel 与 `trtllm-gen` 后端, 这些接口未来可能变更。
 5. 性能风险: 小 batch 下 NVFP4 通过量可能不及 FP8, 但该场景已通过 benchmark 确认可接受。
 - 影响: 用户: 支持 SM100 GPU 的用户可以通过 `--kv-cache-dtype nvfp4` 启用更低精度的 KV cache, 减少显存占用并提升大 batch 吞吐量; 其他用户无影响。系统: 新增一个 KV cache dtype 及配套转换逻辑, 增加代码复杂度约 400 行 (源码 + 测试)。
 - 团队: 需维护 NVFP4 与 FlashInfer 后端的兼容性, 并监控精度退化报告。
 - 风险标记: 仅支持 SM100, 实验性量化格式, 依赖 FlashInfer TRTLLM 后端, FP8 反量化精度损失, MLA 不兼容 (已添加校验), 小 batch 性能无提升

关联脉络

- PR #40033 [NVFP4][Hopper/AMD Instinct] Add Triton kernels for NVFP4 dequantization and QDQ emulation: 同为 NVFP4 量化基础设施, 提供了反量化工具函数 (`dequant_nvfp4_kv_cache` 等) 被本 PR 测试代码引用。
- PR #41326 Faster per-token fp8 group quant packed kernel for blackwell: 同为 Blackwell GPU 上的量化优化, 与 NVFP4 关注同一硬件平台, 可视为性能互补。