

PR #47857 完整报告

vllm-project/vllm

[Model] Add LongCat-Flash-Lite (n-gram embedding)

合并时间: 2026-07-10 22:17

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/47857>

执行摘要

- 一句话: 添加 LongCat-Flash-Lite n-gram 嵌入变体模型
- 推荐动作: 该 PR 值得阅读, 特别是 ModelState 隔离的架构模式、CUDA 自定义算子集成方式以及 fix 中的 guard 变量技巧。对于计划集成类似模型或需要理解 MRV2 的开发者有参考价值。

功能与动机

支持 n-gram 嵌入变体, 是 LongCat-Flash 系列的扩展。PR body 说明: 'Adds LongcatFlashNgramForCausalLM, the n-gram-embedding variant of LongCat-Flash, as a Model-Runner-V2 model. The n-gram input layer's per-request token history is isolated in a ModelState.' 之前尝试 #33611 采用了不同方式, 本次通过 ModelState 隔离更干净。

实现拆解

1. 新增 vllm/model_executor/models/longcat_flash_ngram.py, 定义 NgramEmbedding 和 FlashNgramModel, 实现 n-gram 嵌入逻辑, 利用 CUDA 内核 (ngram_embedding_kernels.cu) 计算 n-gram id, ModelState 类隔离每个请求的 token 历史。
2. 在 vllm/model_executor/models/config.py 中添加 LongcatFlashNgramForCausalLMConfig, 设置默认编译模式为 NONE、CUDA Graph 模式为 FULL, 避开 torch.compile 和 PIECEWISE 不兼容问题, 并将其注册到 MODELS_CONFIG_MAP。
3. 在 vllm/_custom_ops.py 和 CUDA 绑定层 (ops.h、torch_bindings.cpp) 添加 ngram_compute_ngram_ids 自定义算子, 封装 CUDA 内核调用。
4. 在 vllm/model_executor/models/registry.py 中注册新模型以启用测试。
5. 修改 vllm/v1/worker/gpu/attn_utils.py, 支持双注意力模块的 KV 缓存绑定 (通过 model_type 为 longcat_flash 和 longcat_flash_ngram 时设置 num_attn_module=2)。
6. 修复 longcat_flash.py 和 longcat_flash_mtp.py 中 MLA 缩放的 double-bake bug: 添加 _mla_q_lora_scaled 等 guard 变量避免重复缩放。
7. 修改 longcat_flash_mtp.py 中使用 object.__setattr__ 绕过严格配置验证, 并修复初始化填充语法错误。

8. 修改 `vllm/config/speculative.py` 中与 MTP 相关的初始化逻辑。

9. 调整测试工具 `tests/models/utils.py` 以支持新模型。

关键文件：

- `vllm/model_executor/models/longcat_flash_ngram.py` (模块 模型实现; 类别 source; 类型 core-logic; 符号 `uses_ngram_embedding`, `_config_dtype`, `NgramEmbedding`, `init`) : 新增文件, 包含 n-gram 嵌入模型核心实现: `NgramEmbedding`、`FlashNgramModel`、`LongcatNgramModelState` 及 `ModelState` 集成, 是 PR 的主要变更。
- `vllm/model_executor/models/config.py` (模块 模型配置; 类别 source; 类型 data-contract; 符号 `LongcatFlashNgramForCausalLMConfig`, `verify_and_update_config`) : 新增 `LongcatFlashNgramForCausalLMConfig` 配置类, 设置默认编译模式避免不兼容问题, 并注册到 `MODELS_CONFIG_MAP`。同时修复了 `LlamaBidirectionalConfig` 中 `pooling_type` 默认值的潜在 bug。
- `vllm/_custom_ops.py` (模块 自定义算子; 类别 source; 类型 core-logic; 符号 `ngram_compute_n_gram_ids`) : 新增 `ngram_compute_n_gram_ids` 自定义算子, 封装 CUDA 内核调用, 是 n-gram 嵌入推理的关键操作。
- `vllm/model_executor/models/longcat_flash_mtp.py` (模块 MTP 实现; 类别 source; 类型 bugfix) : 修复 MLA 缩放 `double-bake` 和 `n_routed_experts` 设置绕过, 是伴随修复的重要文件。
- `vllm/model_executor/models/longcat_flash.py` (模块 原始模型; 类别 source; 类型 bugfix) : 修复 MLA 缩放 `double-bake`, 添加 `guard` 变量。
- `vllm/v1/worker/gpu/attn_utils.py` (模块 注意力工具; 类别 source; 类型 core-logic) : 修改 KV 缓存绑定支持双注意力模块, 通过模型类型判断设置 `num_attn_module=2`。
- `csrc/libtorch_stable/ops.h` (模块 CUDA 头文件; 类别 source; 类型 core-logic) : 声明 n-gram 内核函数, 与 `torch_bindings.cpp` 配合注册自定义算子。
- `csrc/libtorch_stable/ngram_embedding_kernels.cu` (模块 CUDA 内核; 类别 other; 类型 core-logic) : 新增 CUDA 内核文件, 实现 n-gram id 计算逻辑 (从 SGLang 适配)。
- `vllm/model_executor/models/registry.py` (模块 模型注册; 类别 source; 类型 data-contract) : 注册新模型 `LongcatFlashNgramForCausalLM`, 启用集成测试。
- `csrc/libtorch_stable/torch_bindings.cpp` (模块 PyTorch 绑定; 类别 source; 类型 core-logic) : 绑定 n-gram 内核到 Torch, 注册为自定义操作。
- `vllm/config/speculative.py` (模块 推测解码; 类别 source; 类型 core-logic) : 修改 MTP 相关的初始化逻辑, 配合 ngram 模型。
- `tests/models/utils.py` (模块 测试工具; 类别 test; 类型 test-coverage) : 扩展测试工具以支持新模型的测试。

关键符号: `LongcatFlashNgramForCausalLM`, `NgramEmbedding.init`, `NgramEmbedding._init_ngram_embeddings`, `NgramEmbedding.embed_batched`, `ngram_compute_n_gram_ids`, `LongcatFlashNgramForCausalLMConfig.verify_and_update_config`, `LongcatNgramModelState.add_request`

关键源码片段

vllm/model_executor/models/longcat_flash_ngram.py

新增文件，包含 n-gram 嵌入模型核心实现：NgramEmbedding、FlashNgramModel、LongcatNgramModelState 及 ModelState 集成，是 PR 的主要变更。

longcat_flash_ngram.py - NgramEmbedding 部分实现

```
class NgramEmbedding(nn.Module):
    """Token embedding fused with hashed n-gram embeddings.

    TP-sharded: the k*(n-1) per-embedder tables are concatenated into one
    :class:`VocabParallelEmbedding` (oe_embedder) with per-embedder offsets,
    and the projections are stacked into one oe_projection applied with a
    single bmm. Hashing math is ported from the HF reference.
    """

    def __init__(self, config: FlashConfig, base_embeddings: nn.Module) -> None:
        super().__init__()
        self.config = config
        self.word_embeddings = base_embeddings

        self.m = config.ngram_vocab_size_ratio * config.vocab_size
        self.k = config.emb_split_num
        self.n = config.emb_neighbor_num
        self.pad_id = config.pad_token_id
        self.eos_token_id = config.eos_token_id
        self._dtype = _config_dtype(config)

        self._init_ngram_embeddings()

    def _init_ngram_embeddings(self) -> None:
        self.num_embedders = self.k * (self.n - 1)
        oe_dim = self.config.hidden_size // self.num_embedders
        self.oe_dim = oe_dim

        # 每个 embedder 的 table 大小 = m + i*2 + 1, 偏移量累加
        sizes = [int(self.m + i * 2 + 1) for i in range(self.num_embedders)]
        offsets = [0]
        for s in sizes:
            offsets.append(offsets[-1] + s)
        self._offsets = offsets
        self._sizes = sizes

        # All embedder tables concatenated into one VocabParallelEmbedding
        self.oe_embedder = VocabParallelEmbedding(
            offsets[-1], oe_dim, params_dtype=self._dtype
        )
        # Stacked projections
        self.oe_projection = nn.Parameter(
```

```

        torch.empty(
            self.num_embedders, oe_dim, self.config.hidden_size, dtype=self._dtype
        ),
        requires_grad=False,
    )

    # Precomputed tables for CUDA kernel
    vocab = self.config.vocab_size
    ne_weights = torch.zeros(self.n - 1, self.k, self.n, dtype=torch.int32)
    ne_mods = torch.zeros(self.n - 1, self.k, dtype=torch.int32)
    for i in range(self.n - 1):
        for j in range(self.k):
            mod = int(self.m + 2 * (i * self.k + j) + 1)
            ne_mods[i, j] = mod
            for delta in range(self.n):
                ne_weights[i, j, delta] = pow(vocab, delta, mod)
    self.register_buffer("ne_weights", ne_weights, persistent=False)
    self.register_buffer("ne_mods", ne_mods, persistent=False)
    self.register_buffer(
        "exclusive_sizes",
        torch.tensor(offsets, dtype=torch.int32),
        persistent=False,
    )

```

vllm/model_executor/models/config.py

新增 LongcatFlashNgramForCausalLMConfig 配置类，设置默认编译模式避免不兼容问题，并注册到 MODELS_CONFIG_MAP。同时修复了 LlamaBidirectionalConfig 中 pooling_type 默认值的潜在 bug。

vllm/model_executor/models/config.py - 新增配置类

```

class LongcatFlashNgramForCausalLMConfig(VerifyAndUpdateConfig):
    @staticmethod
    def verify_and_update_config(vllm_config: "VllmConfig") -> None:
        # LongCat-Flash-Lite 的 zero-expert MoE 在 torch.compile 下触发数据依赖 assert,
        # n-gram inputs_embeds 只适配了 FULL CUDAGraph (PIECEWISE 会中断)。
        # 因此默认关闭编译且使用 FULL CUDAGraph, 除非用户明确指定。
        from vllm.config.compilation import CompilationMode, CUDAGraphMode

        compilation_config = vllm_config.compilation_config
        if compilation_config.mode is None:
            compilation_config.mode = CompilationMode.NONE
        if compilation_config.cudagraph_mode is None:
            compilation_config.cudagraph_mode = CUDAGraphMode.FULL

```

评论区精华

- Claude[bot]指出 LongcatNgramModelState.add_request 中 n-gram 左上下文在 resume 路径（抢占、prefix-cache 命中、KV 转移）下可能使用空 / 截断数据，导致 CUDA 内核静

默错误。该问题未被本次修复，列为已知问题。

- LucasWilkinson对 `attn_utils.py` 中通过硬编码模型名称判断双注意力模块表示担忧 ('it would be nice to remove model specific stuff from here')，但接受当前方案。mgoin 回应同意后续改进。
- LucasWilkinson赞扬了 `ModelState` 的使用 ('nice use of ModelState!')。
- LucasWilkinson提出 nit 建议：如果 `remove_request` 时填充 -1，可以消除 `_req_id_to_index` 映射。但非必选。
- Resume 路径 n-gram 上下文安全性 (correctness): 问题已记录但未修复，作为已知问题留给后续 PR。
- `attn_utils` 中模型特定代码的抽象 (design): 接受当前方式，后续考虑改进。
- `ModelState` 实现赞扬 (style): 设计得到认可。
- `_req_id_to_index` 映射优化 (design): 非强制优化，可后续考虑。

风险与影响

- 风险:
 - Resume 路径数据错误: claudel[bot] 指出的 n-gram 上下文截断问题未修复，可能导致静默推理错误，影响需要 resume 的场景（如 preemption、KV 转移）。
 - CUDAGraph 兼容性限制: 模型强制禁用 `torch.compile` 并使用 FULL CUDAGraph，无法使用 PIECEWISE 模式，限制了在某些输入长度下的优化收益。
 - Zero-expert MoE Dynamo 不兼容: `torch.compile` 因数据依赖 assert 失败，可能需要额外工作支持。
 - 模型特定代码侵入: `attn_utils.py` 中硬编码的模型类型判断降低了通用性。
- 影响:
 - 用户: 新模型可直接通过 `vllm serve` 部署，无需手动配置编译或 CUDAGraph 参数。验证支持 GSM8K 5-shot ~84% 正确率，0% 无效输出。
 - 系统: 新增约 500 行源码 (Python + CUDA)，编译依赖 CUDA 工具链。模型注册和配置修改影响全局模型加载路径。
 - 团队: 为后续 MRV2 模型集成提供了参考模式，特别是 `ModelState` 隔离和自定义算子集成。
 - 风险标记: resume 路径未修复，CUDAGraph 兼容限制，模型特定逻辑侵入，`torch.compile` 不兼容

关联脉络

- PR #33611 Abandoned attempt at n-gram embedding (different approach): 被当前 PR 替代，之前尝试通过共享 runner 而非 `ModelState` 隔离，已废弃。
- PR #46623 LongCat-Next (different multimodal model, not duplicate): 同一作者的不同模型 PR，但功能不同，PR body 说明不是重复。