

PR #49686 完整报告

vllm-project/vllm

[Multimodal] Expose mm hash algorithm selection to cli args

合并时间: 2026-07-31 13:25

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49686>

执行摘要

- 一句话: 将多模态哈希算法暴露为 `--mm-hasher-algorithm` CLI 参数
- 推荐动作: 值得精读。核心看点有三个: 一是环境变量向 CLI 参数迁移的兼容性处理 (`get_from_deprecated_env_if_set`) ; 二是把『分组是否可批』从密码学哈希比较重构为直接数值比较的设计权衡, 既消除对哈希算法的隐式依赖又降低开销; 三是 `hash_kwargs` 使用仅位置参数强制调用方显式传算法, 属于用接口设计防止配置漂移的典型案列。建议结合 `tests/multimodal/test_utils.py` 的 `dtype` 拆分测试一起阅读。

功能与动机

PR body 明确指出: 目前 mm hash 算法只能通过环境变量 `VLLM_MM_HASHER_ALGORITHM` 配置, 本 PR 的目标是暴露为 CLI 参数, 让用户可以直接通过 `--mm-hasher-algorithm` 指定。在 review 中 DarkLight1337 提出『Should we deprecate the environment variable then?』, 作者确认并实现弃用机制。此外, review 中 euisuh 发现分组哈希被硬编码为 `sha256` 属于意外行为变更, 驱动了后续『remove hashing necessity for grouping mm kwargs』的重构提交。背景还关联 issue #18334: 为 FIPS 合规部署提供 `sha256/sha512` 选择。

实现拆解

1. 配置层新增类型与默认值: 在 `vllm/config/multimodal.py` 中新增 `MMHasherAlgorithm = Literal["blake3", "sha256", "sha512"]` 类型, 并新增默认工厂 `_get_mm_hasher_algorithm()`, 通过 `get_from_deprecated_env_if_set` 读取已弃用的 `VLLM_MM_HASHER_ALGORITHM` (弃用版本 v0.27), 未设置时回退到 `blake3`。`MultiModalConfig` 新增 `mm_hasher_algorithm` 字段, 成为所有下游读取的唯一入口。
2. CLI 接线: 在 `vllm/engine/arg_utils.py` 的 `EngineArgs` 中增加 `mm_hasher_algorithm` 字段并从 `MultiModalConfig` 继承默认值; `add_cli_args()` 注册 `--mm-hasher-algorithm` 参数; `create_model_config()` 将该值传入 `MultiModalConfig`, 实现 env var 与 CLI 共用同一配置源。
3. 哈希接口显式化: `vllm/multimodal/hasher.py` 中 `MultiModalHasher.hash_kwargs()` 增加仅位置参数 `algorithm`, `_get_hasher_factory()` 不再自行读取环境变量或做小写归一化。所有调用方 (`ProcessorInputs.get_mm_hashes`、`processor.py` 的 `_apply_hf_processor`、`terratorch.py`、`transformers/multimodal.py`、`renderers/hf.py` 等) 均改为从 `mm_config.mm_hasher_algorithm` 显式透传。

4. 分组逻辑去哈希化: vllm/multimodal/utils.py 删除 `_get_group_hash()` 与 `groupby(group_ids)` 方案, 新增 `_can_batch_mm_items()` 相邻比较函数; `group_and_batch_mm_items()` 改为贪心扫描、相邻可合并则延展当前组。
vllm/multimodal/inputs.py 的 `nested_tensors_equal()` 新增 `check_dtype` 参数, shared field 比较时强制 dtype 一致, 分组结果不再受哈希算法配置影响, 且省去逐项哈希的开销。
5. 测试与文档配套: tests/multimodal/test_hasher.py 新增 `test_hash_algorithm` (与 hashlib 结果对照) 与 `test_hash_algorithm_required` (缺省 algorithm 抛 TypeError), 存量用例统一显式传 `blake3`; tests/multimodal/test_utils.py 新增 `test_group_and_batch_mm_items_splits_shared_data_by_dtype` 回归测试; tests/config/test_multimodal_config.py 补充非法算法校验测试; 修复 moss_audio 等下游调用点; docs/usage/security.md 更新 FIPS 用法与环境变量弃用说明。

关键文件:

- vllm/multimodal/utils.py (模块 多模态组批; 类别 source; 类型 core-logic; 符号 `_can_batch_mm_items`, `group_and_batch_mm_items`): 分组批处理核心逻辑: 删除基于哈希的 `_get_group_hash` 与 `groupby` 方案, 新增 `_can_batch_mm_items` 逐项比较, 分组结果不再受哈希算法配置影响
- vllm/multimodal/hasher.py (模块 多模态哈希; 类别 source; 类型 core-logic; 符号 `MultiModalHasher.hash_kwargs`, `_get_hasher_factory`): 哈希接口核心变更: `hash_kwargs` 新增仅位置参数 `algorithm`, 不再隐式读取环境变量, 所有调用方必须显式透传配置
- vllm/config/multimodal.py (模块 配置层; 类别 source; 类型 core-logic; 符号 `MMHasherAlgorithm`, `_get_mm_hasher_algorithm`, `MultiModalConfig.mm_hasher_algorithm`): 配置层入口: 新增 `MMHasherAlgorithm` 类型、`_get_mm_hasher_algorithm` 默认工厂与 `mm_hasher_algorithm` 配置字段, 负责环境变量弃用回退
- vllm/engine/arg_utils.py (模块 参数解析; 类别 source; 类型 dependency-wiring; 符号 `EngineArgs.mm_hasher_algorithm`, `EngineArgs.add_cli_args`): CLI 接线: `EngineArgs` 增加 `mm_hasher_algorithm` 字段, `add_cli_args` 注册 `--mm-hashers-algorithm`, `create_model_config` 完成配置透传
- tests/multimodal/test_utils.py (模块 组批测试; 类别 test; 类型 test-coverage; 符号 `test_group_and_batch_mm_items_splits_shared_data_by_dtype`): 新增按 dtype 拆分 shared data 的回归测试, 是分组逻辑去哈希化后最重要的行为保障
- tests/multimodal/test_hasher.py (模块 哈希测试; 类别 test; 类型 test-coverage; 符号 `test_hash_algorithm`, `test_hash_algorithm_required`): 哈希接口变更的测试配套: 新增算法对照测试与必填参数测试, 存量用例全部显式传入 `blake3`

关键符号: `MultiModalHasher.hash_kwargs`, `_get_hasher_factory`, `group_and_batch_mm_items`, `_can_batch_mm_items`, `nested_tensors_equal`, `get_mm_hashes`, `_get_mm_hasher_algorithm`, `EngineArgs.add_cli_args`

关键源码片段

vllm/multimodal/utils.py

分组批处理核心逻辑：删除基于哈希的 `_get_group_hash` 与 `groupby` 方案，新增 `_can_batch_mm_items` 逐项比较，分组结果不再受哈希算法配置影响

```
def _can_batch_mm_items(
    left: MultiModalKwargsItem,
    right: MultiModalKwargsItem,
) -> bool:
    # 判断两个相邻条目能否安全合并到同一批。
    # 原实现基于 MultiModalHasher.hash_kwargs() 对 shared field 计哈希来比较，
    # 既昂贵又受哈希算法配置影响；改为直接按字段类型与嵌套张量值比较，
    # shared field 命中时额外检查 dtype，保证不同 dtype 的组数据不会被合并。
    if left.keys() != right.keys():
        return False

    for key, left_elem in left.items():
        right_elem = right[key]
        left_field, right_field = left_elem.field, right_elem.field
        is_shared_field = isinstance(left_field, MultiModalSharedField) and isinstance(
            right_field, MultiModalSharedField
        )
        if (type(left_field) is not type(right_field)) or (
            is_shared_field
            and not nested_tensors_equal(
                left_elem.data, right_elem.data, check_dtype=True
            )
        ):
            return False

    return True


def group_and_batch_mm_items(
    items: Sequence[MultiModalKwargsItem],
    *,
    device: torch.types.Device = None,
    pin_memory: bool = False,
) -> Generator[tuple[int, BatchedTensorInputs]]:
    # 将连续条目（可能来自不同请求）按可否合并分组，并批量归约数据。
    # 只做相邻比较（贪心遍历），不再对整批条目预计算哈希 group_ids，
    # 分组耗时从 O(N) 次哈希下降为 O(N) 次轻量比较；shared field 语义
    # 由字段类型一致性 + 嵌套张量相等（含 dtype）保证。
    start_idx = 0
    for end_idx in range(1, len(items) + 1):
        # 相邻条目可合并则延展当前组，否则就在 end_idx 处切分一批
        if end_idx < len(items) and _can_batch_mm_items(
            items[end_idx - 1], items[end_idx]
        ):
            continue
        yield start_idx, items[start_idx:end_idx]
        start_idx = end_idx
```

```

        continue

    group_data = _batch_mm_items(
        items[start_idx:end_idx],
        device=device,
        pin_memory=pin_memory,
    )
    yield end_idx - start_idx, group_data
    start_idx = end_idx

assert start_idx == len(items)

```

vllm/multimodal/hasher.py

哈希接口核心变更：hash_kwargs 新增仅位置参数 algorithm，不再隐式读取环境变量，所有调用方必须显式透传配置

```

@functools.lru_cache(maxsize=3)
def _get_hasher_factory(
    algorithm: MMHasherAlgorithm,
) -> Callable[[], 'hashlib._Hash']:
    # 按算法名返回可创建哈希器实例的工厂函数，结果带 LRU 缓存。
    # 算法由调用方显式传入（默认来自 MultiModalConfig.mm_hasher_algorithm），
    # 工厂不再内部读取环境变量；调用前已完成小写归一化与合法性校验。
    if algorithm == 'blake3':
        from blake3 import blake3

        return blake3
    elif algorithm == 'sha256':
        return hashlib.sha256
    elif algorithm == 'sha512':
        return hashlib.sha512
    else:
        # 配置层已用 Literal 类型约束，走到这里说明调用方传入了非法值
        raise ValueError(f'Unsupported hash algorithm: {algorithm}')

```

```

class MultiModalHasher:
    @classmethod
    def hash_kwargs(
        cls,
        algorithm: MMHasherAlgorithm,
        /,
        **kwargs: object,
    ) -> str:
        # 按关键字参数生成多模态输入缓存用的哈希摘要。
        # algorithm 设计为仅位置参数，防止与模型自定义字段 algorithm 混淆，
        # 也强制所有调用点（processor、renderer、model）显式透传 CLI 配置；
        # 序列化规则沿用 iter_item_to_bytes 的递归降维逻辑。
        hasher_factory = _get_hasher_factory(algorithm)

```

```

hasher = hasher_factory()

for k, v in sorted(kwargs.items(), key=lambda kv: kv[0]):
    for bytes_ in cls.iter_item_to_bytes(k, v):
        hasher.update(bytes_)

return hasher.hexdigest()

```

vllm/config/multimodal.py

配置层入口：新增 MMHasherAlgorithm 类型、_get_mm_hasher_algorithm 默认工厂与 mm_hasher_algorithm 配置字段，负责环境变量弃用回退

```
MMHasherAlgorithm = Literal['blake3', 'sha256', 'sha512']
```

```

def _get_mm_hasher_algorithm() -> MMHasherAlgorithm:
    # 作为 MultiModalConfig.mm_hasher_algorithm 的默认工厂。
    # 优先取 CLI 传入的值；若未配置则回退到旧环境变量
    # VLLM_MM_HASHER_ALGORITHM（将在 v0.27 弃用），最终默认 blake3。
    # 这样既给用户新的 CLI 入口，又保持旧部署的向后兼容。
    env_value = get_from_deprecated_env_if_set(
        'VLLM_MM_HASHER_ALGORITHM',
        'v0.27',
        'mm_hasher_algorithm',
    )
    env_value = 'blake3' if env_value is None else env_value
    return cast(MMHasherAlgorithm, env_value.lower())

```

```

@config
class MultiModalConfig:
    # ... 其他多模态配置字段 ...
    mm_hasher_algorithm: MMHasherAlgorithm = Field(
        default_factory=_get_mm_hasher_algorithm
    )
    # 多模态输入缓存使用的哈希算法；FIPS 合规部署可设 sha256 或 sha512。

```

评论区精华

1. 环境变量是否弃用 (design) : DarkLight1337 在 PR 评论中直接提问『Should we deprecate the environment variable then?』，作者回复『Sure, let's deprecate it.』，最终通过 `get_from_deprecated_env_if_set(..., "v0.27", "mm_hasher_algorithm")` 落地，保留向后兼容的过渡期。
2. 分组哈希硬编码 sha256 的意外变更 (correctness) : euisuh 在 `vllm/multimodal/utils.py` 的 diff 上指出，`_get_group_hash()` 原本通过 `MultiModalHasher.hash_kwargs()` 遵循 `VLLM_MM_HASHER_ALGORITHM`，而 PR 中间版本硬编码 `sha256`，导致 shared-field 分组既不跟随环境变量也不跟随新 CLI 参数，且默认行为从 `blake3` 变为 `sha256`。作者承认『Oh, good catch!』。

3. 用直接比较替代哈希 (performance)：作者随后在 `utils.py` 评论中提出：shared fields 通常很小，可以直接比较数值与 dtype 来避免昂贵哈希，最终以 `_can_batch_mm_items` 加 `nested_tensors_equal(check_dtype=True)` 落地，并新增按 dtype 拆分的回归测试。
- 是否弃用 `VLLM_MM_HASHER_ALGORITHM` 环境变量 (design): 通过 `get_from_deprecated_env_if_set` 保留环境变量回退能力，并标记 v0.27 弃用，CLI 参数成为首选配置方式。
 - 分组哈希被硬编码 sha256 的意外行为变更 (correctness): 没有把算法透传给哈希，而是直接移除分组哈希，改为 `_can_batch_mm_items` 的数值与 dtype 比较，从根本上消除该问题。
 - shared field 比较能否避免昂贵哈希 (performance): 新增 `nested_tensors_equal(check_dtype=True)` 与贪心相邻比较，并补充按 dtype 拆分的回归测试。

风险与影响

- 风险：
 1. 接口破坏性变更：MultiModalHasher.hash_kwargs() 从 (`**kwargs`) 改为 (`algorithm, /, **kwargs`)，任何未更新的第三方调用都会立即抛 `TypeError`；仓库内调用点已全部同步，但外部直接使用该 API 的代码需要适配。
 1. 分组语义变化：shared-field 分组由哈希比较改为字段类型 + 值 + dtype 精确比较。不同 dtype 的相邻条目现在会被强制拆分（旧逻辑下可能因哈希碰撞被合并），`group_and_batch_mm_items` 产出的批形状可能变化，影响依赖该函数的上游批处理逻辑。
 2. 缓存键失效：用户从默认 blake3 切换到 sha256/sha512 后，多模态处理器缓存（LRU/shm）的所有既有哈希键全部失效，缓存命中率短期归零，需要重新处理输入。
 3. 迁移期歧义：环境变量与 CLI 参数在 v0.27 前并存，若两者同时设置且结果不同，以 CLI 为准，日志中的弃用告警可能造成运维混淆。
 4. 多调用点透传遗漏：哈希算法需要穿透 `get_mm_hashes`、多个 processor 与 renderer 路径，改动横跨 17 个文件且经历了 5 次 main 合并，存在个别调用点未同步的风险（PR 中已修复 `moss/test` 等遗漏）。- 影响：对用户而言，获取了更直观的配置入口，FIPS 场景可通过 `--mm-hasher-algorithm sha256/sha512` 一键启用，不再依赖环境变量。对系统而言，多模态缓存键的行为更可控，分组批处理去哈希化后性能更好、语义更清晰，但切换算法会带来缓存冷启动成本。对团队而言，该 PR 为 v0.27 移除环境变量铺平了道路，并确立『配置集中在 MultiModalConfig、哈希算法显式透传』的约定；后续新增调用方必须从 `mm_config` 取算法，否则 `hash_kwargs` 会直接报错。整体影响面主要局限在多模态输入处理链。- 风险标记：hash_kwargs 接口破坏性变更，分组语义从哈希改为值比较，切换算法导致缓存键全部失效，环境变量与 CLI 并存期迁移歧义，多调用点透传易遗漏

关联脉络

- 暂无明显关联 PR