

# PR #41205 完整报告

vllm-project/vllm

Fix Nano Nemotron text-only weight loading

合并时间: 2026-05-05 02:43

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/41205>

## 执行摘要

- 一句话: 修复 Nano Nemotron VL 文本模式下多模态权重加载
- 推荐动作: 该 PR 设计清晰, 通过复用 `multimodal_config` 接口统一判断文本仅模式, 避免了引入额外配置字段。新增测试覆盖全面, 值得精读。建议关注类似多模态模型 (如 Qwen-VL) 的文本仅模式实现, 可复用此模式。

## 功能与动机

Fix Nano Nemotron VL weight loading when all multimodal prompt limits are zero, including `--language-model-only`. In this mode the multimodal tower modules are replaced by missing-stage placeholders, so adapter, vision, and audio checkpoint weights should be skipped instead of inspected or loaded.

## 实现拆解

1. 加载入口判断: 在 `vllm/model_executor/models/nano_nemotron_vl.py` 的 `load_weights` 起始处, 从 `self.model_config.multimodal_config` 获取配置, 通过检查 `image`、`video`、`audio` 三种模态的 `get_limit_per_prompt` 返回值是否全为零, 得到布尔值 `load_multimodal_weights`。
2. 条件跳过多模态权重: 在遍历权重时, 对 `mlp1` (适配器)、`vision_model.radio_model.*` (视觉)、`sound` (音频) 权重的处理分支中, 若 `load_multimodal_weights` 为 `False`, 则直接 `continue`, 不执行任何加载或断言。
3. 条件调用子模块加载: 在收集完权重后, 仅当 `load_multimodal_weights` 为 `True` 时, 才调用 `self.vision_model.load_weights(vision_weights)` 和 `self.sound_encoder.load_weights(sound_weights)`。语言模型权重 `self.language_model.load_weights(llm_weights)` 始终执行。
4. 测试覆盖: 新增 `tests/models/multimodal/test_nano_nemotron_vl.py`, 使用 `stub` 对象模拟三种配置场景: 文本仅模式 (验证多模态权重被跳过)、仅图像模式 (验证视觉权重加载且音频权重缺失触发断言)、以及仅图像模式但提供音频权重 (验证严格断言)。

关键文件:

- `vllm/model_executor/models/nano_nemotron_vl.py` (模块 模型加载; 类别 `source`; 类型 `core-logic`; 符号 `load_weights`): 核心权重加载逻辑, 新增 `load_multimodal_weights` 标志控制多模态权重加载

- tests/models/multimodal/test\_nano\_nemotron\_vl.py (模块 模型测试; 类别 test; 类型 test-coverage; 符号 \_TextOnlyMultiModalConfig, get\_limit\_per\_prompt, \_ImageOnlyMultiModalConfig, \_ModelConfig) : 新增测试覆盖文本仅模式跳过多模态权重等三种场景
- vllm/transformers\_utils/configs/hyperclovax.py (模块 配置工具; 类别 source; 类型 style) : 格式调整, 不影响功能

关键符号: load\_weights

## 关键源码片段

### vllm/model\_executor/models/nano\_nemotron\_vl.py

核心权重加载逻辑, 新增 load\_multimodal\_weights 标志控制多模态权重加载

```
def load_weights(self, weights: Iterable[tuple[str, torch.Tensor]]):
    # 从模型配置中获取多模态配置, 并判断是否所有模态的 prompt limit 都为零
    mm_config = self.model_config.multimodal_config
    load_multimodal_weights = not all(
        mm_config.get_limit_per_prompt(modality) == 0
        for modality in ("image", "video", "audio")
    )
    # 适配器参数字典始终构建, 因为 mlp1 独立于多模态加载
    adapter_dict = dict(self.mlp1.named_parameters())

    def is_llm(name: str) -> bool:
        return name.startswith("language_model")

    def is_adapter_weights(weight: tuple[str, torch.Tensor]):
        return weight[0].startswith("mlp1")

    def is_vision_weights(name: str) -> bool:
        return name.startswith("vision_model.radio_model.")

    def is_sound_weights(name: str) -> bool:
        return name.startswith("sound")

    llm_weights, vision_weights, sound_weights = [], [], []

    for name, w in weights:
        if is_llm(name):
            llm_weights.append((".".join(name.split(".")[:-1]), w))
        elif is_adapter_weights((name, w)):
            if not load_multimodal_weights:
                continue # 跳过适配器权重
            trimmed_name = ".".join(name.split(".")[:-1])
            param = adapter_dict[trimmed_name]
            with torch.no_grad():
                default_weight_loader(param, w)
```

```

elif is_vision_weights(name):
    if not load_multimodal_weights:
        continue # 跳过视觉权重
    hf_key = name[len("vision_model."):]
    vision_weights.append((hf_key, w))
elif is_sound_weights(name):
    if not load_multimodal_weights:
        continue # 跳过音频权重
    assert self.sound_encoder is not None
    sound_weights.append((name, w))

self.language_model.load_weights(llm_weights) # 始终加载 LLM 权重
if load_multimodal_weights:
    self.vision_model.load_weights(vision_weights)
    if self.sound_encoder is not None and len(sound_weights) > 0:
        self.sound_encoder.load_weights(sound_weights)

```

## tests/models/multimodal/test\_nano\_nemotron\_vl.py

新增测试覆盖文本仅模式跳过多模态权重等三种场景

```

def test_nano_nemotron_vl_skips_multimodal_weights_in_text_only_mode():
    # 通过 object.__new__ 绕过 __init__, 以便手动设置内部状态
    model = object.__new__(NemotronH_Nano_VL_V2)
    language_model = _LanguageModel() # 记录加载的权重
    object.__setattr__(model, "model_config", _ModelConfig()) # 文本仅配置
    object.__setattr__(model, "language_model", language_model)
    object.__setattr__(model, "mlp1", _AdapterModule())
    object.__setattr__(model, "vision_model", _MissingMultiModalModule()) # 占位符, 应跳过
    object.__setattr__(model, "sound_encoder", None)

    language_weight = object()
    model.load_weights([
        ("language_model.layers.0.weight", language_weight),
        ("mlp1.0.weight", object()),
        ("vision_model.radio_model.encoder.weight", object()),
        ("sound_encoder.encoder.weight", object()),
    ])

    # 期望仅语言模型权重被加载, 其他被跳过
    assert language_model.loaded_weights == [("layers.0.weight", language_weight)]

```

## 评论区精华

- 声音编码器缺失时的断言处理: gemini-code-assist 指出当 load\_multimodal\_weights 为 True 但 sound\_encoder 为 None 时, assert self.sound\_encoder is not None 会导致崩溃, 建议改为条件跳过。netanel-haber 倾向于保留严格断言以尽早暴露配置问题, 建议将断言提前。Baekpica 最终保留严格断言但调整了条件顺序: 先断言再加载, 仅在多模态启用且遇到音频权重时触发。

- `adapter_dict` 构建位置: netanel-haber 询问是否即使在 `load_multimodal_weights` 为 `False` 时也需要构建 `adapter_dict`。Baekpica 确认 `mlp1` 独立存在, 始终构建字典, 并通过后续的 `continue` 避免实际加载适配器权重。
- Sound encoder 缺失时的 `assert` 处理 (correctness): Baekpica 采纳了保留严格断言但调整条件顺序的方案: 先断言 `sound_encoder` 不为 `None`, 再加载权重。
- `adapter_dict` 构建位置 (design): Baekpica 确认 `mlp1` 独立于多模态加载, 始终构建 `adapter_dict`。

## 风险与影响

- 风险: 主要风险依赖 `multimodal_config.get_limit_per_prompt` 的返回值, 若返回非预期值 (如负数) 可能导致误判。严格断言在多模态启用但音频编码器缺失时直接崩溃, 但这也是主动防御, 避免静默错误。新增的三个单元测试覆盖了文本仅模式和图像仅模式的正反场景, 降低了回归风险。变更仅影响 Nano Nemotron VL 模型, 且代码修改量小 (+15/-3), 性能影响可忽略。
- 影响: 影响范围限定于 `NemotronH_Nano_VL_V2` 模型, 在文本仅模式下避免启动崩溃; 对已有正常多模态使用无影响。测试新增 114 行, 覆盖了关键边界场景, 提高了该模型权重加载的健壮性。对其他多模态模型无直接影响, 但修复思路可借鉴。
- 风险标记: 多模态配置依赖, 核心路径变更

## 关联脉络

- 暂无明显关联 PR