

PR #36916 完整报告

sgl-project/sglang

[Diffusion] Detect quantized transformer replacements

合并时间: 2026-08-31 13:55

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/36916>

执行摘要

- 一句话: 从替换权重自身探测量化声明, 修复预量化 transformer 覆盖加载
- 推荐动作: 值得精读, 尤其是 `_resolve_weight_override_quantization` 的“声明合并 + 张量级探测 + fails closed”三段式设计: 以物化权重为唯一事实来源、多来源声明递归合并并拒绝冲突、无声明时通过 dtype 探测兜底并拒绝加载, 这套模式对任何需要从 checkpoint 自描述量化格式的加载器都有借鉴价值。建议后续将 `assert metadata_spec is not None` 替换为显式错误, 并为探测循环增加分片采样以控制启动开销。

功能与动机

PR body 提出三个目标: 从物化后的替换 checkpoint 推导原生 transformer 量化而非基础组件; 在模型构建前拒绝冲突或不支持的序列化声明; 保持在线量化与 Nunchaku 和序列化替换权重分离。作者在 issue 评论中补充说明: B200 Flux.2 NVFP4 启动失败根因是 layout-only 序列化 header (只有 `quant_algo` 没有 `quant_method`) 被误当作完整声明处理, 修复后这类 header 保持为声明 checkpoint 并进入既有的 safetensors 布局推断, 不支持的布局仍在构建前拒绝。

实现拆解

实现按 5 步展开:

1. 数据契约扩展: `python/sglang/multimodal_gen/configs/models/dits/base.py` 中 `DiTArchConfig` 新增 `quant_ignore_remap: dict` 字段, 为量化配置的 ignore 层提供架构级重映射表, 供替换权重量化解析时使用。
2. 替换权重量化解析核心: `python/sglang/multimodal_gen/runtime/loader/transformer_load_utils.py` 新增 `_resolve_weight_override_quantization`, 以“物化后的替换权重集”为唯一事实来源, 按优先级合并三路信息: 替换权重目录下 `config.json` 的 `quantization_config`、safetensors 文件头 metadata (`_quantization_metadata / quantization_config`)、以及张量级 dtype 探测 (`.weight_scale / .input_scale / .comfy_quant` 后缀, 或 dtype 为 `F8_E4M3 / I8 / U8` 的 `.weight` 张量)。同时新增 `_merge_quant_declaration` 做递归声明合并, 冲突字段直接抛 `ValueError`。
3. 解析入口重构: `_resolve_quant_config` 新增 `arch_config` 参数 (透传来自 `DiTArchConfig` 的 reverse 映射与 ignore 重映射), 在存在 `transformer_weights_path` 时优先采用替换权重的解析结果; 当替换权重已声明或已探测到量化时, 拒绝再叠加在线

--quantization (fails closed) 。

4. 加载规格与互斥校验: `resolve_transformer_quant_load_spec` 透传 `arch_config`, 并新增“替换 checkpoint 量化与 Nunchaku 互斥”的校验; `component_loaders/transformer_loader.py` 的 `load_customized` 把 `dit_config.arch_config` 接入调用链。
5. 测试配套: `test_transformer_quant.py` 新增 5 个用例, 覆盖相邻量化配置优先、未量化替换不继承 base 配置、无 `quant_method` 的 header 延迟到布局推断、声明量化拒绝在线量化、未声明量化张量 fails closed; 并调整 `_make_server_args` 的 `arch_config` 默认字段以匹配新数据契约。

关键文件:

- `python/sglang/multimodal_gen/runtime/loader/transformer_load_utils.py` (模块 加载器; 类别 source; 类型 core-logic; 符号 `_merge_quant_declaration`, `_resolve_weight_override_quantization`, `_resolve_quant_config`, `resolve_transformer_quant_load_spec`) : 核心变更文件: 新增 `_resolve_weight_override_quantization` 与 `_merge_quant_declaration`, 重构 `_resolve_quant_config` 的优先级与冲突校验, 并在 `resolve_transformer_quant_load_spec` 中新增与 Nunchaku 的互斥检查。
- `python/sglang/multimodal_gen/test/unit/test_transformer_quant.py` (模块 量化测试; 类别 test; 类型 test-coverage; 符号 `test_weight_override_uses_adjacent_quantization_config`, `test_unquantized_weight_override_does_not_inherit_base_config`, `test_weight_override_defers_header_without_quant_method_to_layout`, `test_declared_weight_override_rejects_online_quantization`) : 新增 5 个单测覆盖替换权重量化解析的各类组合, 是验证 fails closed 与延迟布局推断行为的关键测试配套。
- `python/sglang/multimodal_gen/configs/models/dits/base.py` (模块 模型配置; 类别 source; 类型 data-contract; 符号 `DiTArchConfig`) : `DiTArchConfig` 新增 `quant_ignore_remap` 数据契约字段, 为量化 ignore 层重映射提供架构级配置位。
- `python/sglang/multimodal_gen/runtime/loader/component_loaders/transformer_loader.py` (模块 组件加载; 类别 source; 类型 dependency-wiring; 符号 `load_customized`) : `load_customized` 将 `dit_config.arch_config` 透传给 `resolve_transformer_quant_load_spec`, 使量化解析能拿到 reverse 映射与 ignore 重映射。

关键符号: `_resolve_weight_override_quantization`, `_merge_quant_declaration`, `_resolve_quant_config`, `resolve_transformer_quant_load_spec`

关键源码片段

`python/sglang/multimodal_gen/test/unit/test_transformer_quant.py`

新增 5 个单测覆盖替换权重量化解析的各类组合, 是验证 fails closed 与延迟布局推断行为的关键测试配套。

```
def test_weight_override_defers_header_without_quant_method_to_layout(self):
    # 场景: safetensors 文件头只有 quant_algo (NVFP4), 没有 quant_method。
    # 这种“布局级”声明不应被当作完整量化声明解析, 而应保留为已声明状态,
```

```

# 交给既有的 safetensors 布局推断流程（修复 B200 Flux.2 NVFP4 启动失败的关键）。
with tempfile.TemporaryDirectory() as directory:
    weights = f"{directory}/model.safetensors"
    save_file(
        {"block.weight": torch.ones((2, 2))},
        weights,
        metadata={"quantization_config": json.dumps({"quant_algo": "NVFP4"})},
    )

    quant_config, declared = _resolve_weight_override_quantization(
        [weights], {}, {}
    )

self.assertIsNone(quant_config)
self.assertTrue(declared)

def test_undeclared_quantized_weight_override_fails_closed(self, _build_nvfp4):
    # 场景：替换权重含 weight_scale 张量（明显已量化），但 config.json 与
    # metadata 都没有任何声明。此时必须 fails closed 拒绝，而不是静默按
    # 未量化权重加载，否则推理结果会是错误的。
    with tempfile.TemporaryDirectory() as directory:
        weights = f"{directory}/model.safetensors"
        save_file(
            {
                "block.weight": torch.ones((2, 2)),
                "block.weight_scale": torch.ones(2),
            },
            weights,
        )
        server_args = self._make_server_args(transformer_weights_path=weights)

        with self.assertRaisesRegex(ValueError, "no supported native"):
            _resolve_quant_config(
                hf_config={},
                server_args=server_args,
                safetensors_list=[weights],
                component_model_path="/base",
            )

```

评论区精华

该 PR 没有外部 Review 评论，review 讨论为空；两条 issue 评论均为作者自审记录，属于关键决策说明：

- B200 Flux.2 NVFP4 启动失败根因：layout-only 序列化 header（quant_algo 无 quant_method）之前被当作完整量化声明处理导致启动失败。修复后这类 header 保持为声明 checkpoint，进入既有 safetensors 布局推断，不支持的布局仍在构建前拒绝。

- CI 失败归因：唯一 NVIDIA 根失败是 FastWan `DmdDenoisingStage` 1-GPU 性能阈值，该测试以 `component_weights_paths: {}` 启动默认模型且不传 `--transformer-weights-path`，不执行新准入路径；组件精度与 2-GPU 失败是 fast-fail 级联，无需 GT 更新或盲目 rerun。
 - B200 Flux.2 NVFP4 启动失败根因 (correctness): 已修复；不支持的布局仍在模型构建前拒绝。
 - CI 根失败归因 (testing): 无需 GT 更新或盲目 rerun。

风险与影响

- 风险：存在以下技术风险：
- 启动性能：`_resolve_weight_override_quantization` 对每个 safetensors 分片做一次全量 key 遍历并读取 dtype (`safe_open + get_slice().get_dtype()`)，超大 DiT checkpoint 下启动耗时可能上升，后续可考虑只探测关键分片。
- 断言依赖：`assert metadata_spec is not None` 依赖 `resolve_checkpoint_quant_spec` 总是返回 spec 的契约，若未来出现无法识别的 metadata 结构会以断言失败而非清晰错误暴露。
- 行为收紧：新增“替换 checkpoint 量化与 Nunchaku 互斥”和“声明量化拒绝在线量化”会直接拒绝一批此前可能跑通的组合配置，属有意收紧，但需要在文档中同步说明。
- Fails closed 影响：未声明但含量化张量（如 `weight_scale`）的替换权重会被拒绝，可能误伤使用自定义 scale 命名的非标准 checkpoint。
- 影响：影响范围集中在 diffusion 子系统，但有明确的用户侧收益：
- 用户侧：通过 `--transformer-weights-path` 加载预量化 transformer（如 NVFP4）的用户获得更准确的量化配置来源，B200 Flux.2 NVFP4 启动失败得到修复；未使用权重覆盖的现有 diffusion 推理不受影响。
- 系统侧：加载链路新增一次全量张量探测，启动时间略增；错误信息更早暴露（模型构建前）。
- 团队侧：提交历史中有 3 个 merge commit 解决与主线冲突（涉及 `transformer_loader.py`、`transformer_load_utils.py`、`test_transformer_quant.py` 三个文件），说明本 PR 与近期 config 重构、组件加载系列改动有交织，后续跟进需注意合并轨迹。
- 风险标记：核心加载路径变更，全量张量探测带来启动开销，与 Nunchaku 互斥为行为变化，未声明量化 fails closed 可能误伤非标准 checkpoint

关联脉络

- PR #36902 [Diffusion] Delegate recognized quantized components to Transformers: 同一 diffusion 量化加载链路：识别出的量化组件委托 Transformers 加载，本 PR 的替换权重量化检测会与之交互。
- PR #36875 [diffusion] Preserve exact component identity during loading: 组件加载身份保持，涉及 component loader 与 weight source 的契约，与本 PR 的替换权重物化路径相关。
- PR #35739 [multimodal] Fix NVFP4 diffusion models on sm_120 (RTX PRO 6000 / RTX 50xx): 同一 ModelOpt NVFP4 量化路径的修复，本 PR 的 B200 Flux.2 NVFP4 修复

是其延续。