

PR #35370 完整报告

sgl-project/sglang

[diffusion] feat: load GGUF transformer checkpoints (MiniMax-H3)

合并时间: 2026-08-20 15:46

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/35370>

执行摘要

- 一句话: MiniMax-H3 新增 GGUF 量化 transformer 加载路径
- 推荐动作: 值得精读。核心设计决策包括: 以头部元数据先行 + mmap-backed 加载避免全量克隆; 按张量级元数据做 per-layer 量化方法分派并复用 SRT 反量化内核; 用实测数据否决 fused MMVQ/MMQ 而选择反量化 + 原生 GEMM; 以及把配置冲突全部前置到下载前 fail-fast。特别建议关注 FP32 精度岛与 BF16 性能权衡的取舍逻辑, 以及 `_reject_lora_on_packed_weights` 把校验放在 offload 物化之前的防御式设计。

功能与动机

MiniMax-H3 的 transformer 在 BF16 下为 61.7 GiB, FL2VA 分区合计 135 GiB。

Layerwise offload 已让模型能在消费级 GPU 上运行, 剩余瓶颈是 checkpoint 下载、主机内存与逐层流式传输的字节量。本 PR 新增预量化 GGUF 路径, Q4_K_M DiT 仅 17.5 GiB, 与在线 `kitchen_int8` 互补: GGUF 避免先下载并固定完整 BF16 DiT, 而 `kitchen_int8` 仍是更快的计算导向路径。

实现拆解

1. 入口与选择协议: `transformer_load_utils.py` 新增 `resolve_transformer_gguf_to_load` 与 `_resolve_gguf_quant_load_spec`, 支持本地文件、`owner/repo/path/file.gguf` 与 `owner/repo:QUANT_TYPE` (仅当后缀唯一匹配时) 三种选择方式; `_validate_gguf_runtime_support` 在下载前拒绝非 CUDA、显式 `--quantization`、SVDQuant、GGUF 组件上的 FSDP、AdaLN cache/online 路径等冲突组合。
2. 元数据先行: 新增 `gguf_weights.py`, `read_gguf_tensor_meta` 在参数对象存在前从文件头解析每个张量的逻辑形状与 packed 形状: 量化张量要求为 2D `.weight` 且 inner 维度对齐 GGML 块大小, 映射为 `(out, row_bytes)` 的 `uint8` 布局并命名为 `qweight`; `gguf_weights_iterator` 保持 packed 数据 mmap-backed, 直到通用 loader 拷贝 TP 分片, 避免加载期间克隆完整的 17.5 GiB 文件。
3. 量化方法与内核选择: 新增 `quantization/gguf.py`, `GGUFConfig.get_quant_method` 按张量元数据逐层分派——未量化类型走 `UnquantizedLinearMethod`, 量化类型走新增的 `GGUFLinearMethod`, 后者复用 SRT 的 `dequantize_gguf_weight` 反量化后执行原生 `nn.functional.linear`。内核选择以 H200 实测为准: 256/2048 token 下 SRT MMQ 耗时 2.60 ms/20.61 ms, 而反量化 + GEMM 仅 0.37 ms/0.91 ms, 故不使用面向低 token LLM 的 fused MMVQ/MMQ。

4. 模型契约适配: `minimax_h3.py` 在 GGUF 量化下跳过 safetensors 的 qkv 重排 (GGUF 已按 `[q_all, k_all, v_all]` 存储); 修剪版社区 checkpoint (含 `adaln_curve_grid`) 的 AdaLN 投影保持 FP32 精度岛以对齐发布的 ComfyUI 实现; `MiniMaxH3MLP` 对 GGUF 量化放开 `fc1` 激活缓冲复用, 重新进入 fused `silu_and_mul` 内核。
5. LoRA 与承载边界: `lora_pipeline.py` 新增 `_reject_lora_on_packed_weights`, 在 `convert_to_lora_layers` 与动态 `set_lora` 两个入口前置检查, 目标层只有 `qweight` 而无 `weight` 时直接报错; 校验放在关闭 offload 之前, 避免内存受限部署被物化层 OOM。同时 `is_target_layer` 改为 `getattr` 兜底, 恢复默认 LoRA 目标准入。
6. 测试与文档: 新增 `test_gguf_diffusion.py` (793 行, 手写最小 GGUF v3 二进制), 覆盖端序 / 损坏头、量化 packed 行布局、pruned 曲线形状、块对齐拒绝与 TP row/column/merged 切片; `test_hf_transformers.py` 覆盖 Hub 量化类型解析与歧义拒绝; `test_minimax_h3_dit_contract.py` 覆盖修剪曲线插值与 FP32 精度岛契约。文档同步更新 CLI、量化兼容性与 MiniMax-H3 用法。

关键文件:

- `python/sglang/multimodal_gen/runtime/loader/gguf_weights.py` (模块 加载器; 类别 source; 类型 core-logic; 符号 `GGUFTensorMeta`, `read_gguf_tensor_meta`, `gguf_weights_iterator`, `_tensor_to_torch`): 新增的 GGUF 元数据读取核心: 在参数对象存在前解析每个张量的逻辑与 packed 形状, 并派生 `qweight` 命名, 是整条加载路径的地基。
- `python/sglang/multimodal_gen/runtime/layers/quantization/gguf.py` (模块 量化层; 类别 source; 类型 core-logic; 符号 `GGUFConfig`, `GGUFLinearMethod`): GGUF 量化适配器: 按张量元数据逐层选择未量化或反量化 + 原生 GEMM 方法, 并完成 TP 分片块对齐校验。
- `python/sglang/multimodal_gen/runtime/loader/transformer_load_utils.py` (模块 加载调度; 类别 source; 类型 core-logic; 符号 `TransformerQuantLoadSpec`, `resolve_transformer_gguf_to_load`, `_resolve_gguf_quant_load_spec`, `_validate_gguf_runtime_support`): GGUF 加载计划的编排与前置校验所在: `TransformerQuantLoadSpec` 增加 `gguf_file`, 下载前拒绝冲突配置。
- `python/sglang/multimodal_gen/runtime/models/dits/minimax_h3.py` (模块 模型适配; 类别 source; 类型 data-contract; 符号 `MiniMaxH3MLP`, `_time_embedding`, `MiniMaxH3DiT`): 模型契约适配: GGUF 的 qkv 布局跳过 safetensors reorder, 修剪版 AdaLN 曲线走 FP32 精度岛, 并放开 `fc1` 激活缓冲复用。
- `python/sglang/multimodal_gen/runtime/pipelines_core/lora_pipeline.py` (模块 LoRA 管线; 类别 source; 类型 core-logic; 符号 `_reject_lora_on_packed_weights`, `convert_to_lora_layers`, `set_lora`): LoRA 与 packed 权重的兼容性防线: 转换或 `set_lora` 前预检, 避免模型被转换到一半或 offload 物化时 OOM。
- `python/sglang/multimodal_gen/test/unit/test_gguf_diffusion.py` (模块 单元测试; 类别 test; 类型 test-coverage; 符号 `_write_gguf`, `TestGGUFTensorMeta`, `read_gguf_tensor_meta`, `gguf_weights_iterator`): 793 行 CPU 单测: 手写最小 GGUF v3 二进制, 覆盖端序、损坏头、量化布局、TP 切片与量化方法选择, 是加载路径质量的主要保障。

- python/sclang/srt/layers/quantization/gguf.py (模块 SRT 量化层; 类别 source; 类型 core-logic; 符号 dequantize_gguf_weight) : SRT 侧反量化辅助函数扩展, 供 diffusion 复用, 是跨模块复用的接缝。
- python/sclang/srt/utils/hf_transformers/common.py (模块 HF 工具; 类别 source; 类型 dependency-wiring; 符号 resolve_hf_gguf_reference) : Hub 侧 GGUF 引用解析能力 (resolve_hf_gguf_reference) 由 SRT 暴露给 diffusion, 消除重复的下载解析逻辑。

关键符号: read_gguf_tensor_meta, gguf_weights_iterator,
GGUFConfig.get_quant_method, GGUFLinearMethod.create_weights,
GGUFLinearMethod.apply, resolve_transformer_gguf_to_load,
_validate_gguf_runtime_support, _reject_lora_on_packed_weights,
dequantize_gguf_weight

关键源码片段

python/sclang/multimodal_gen/runtime/pipelines_core/lora_pipeline.py

LoRA 与 packed 权重的兼容性防线: 转换或 `set_lora` 前预检, 避免模型被转换到一半或 offload 物化时 OOM。

```
def _reject_lora_on_packed_weights(self) -> None:
    """在任何层被替换前, 先拒绝指向 packed 权重的 LoRA 请求。

    `BaseLayerWithLoRA` 读取 `base_layer.weight`, 而 GGUF 这类打包量化
    只注册 `qweight`。提前检查可以避免:
    1. 模型被转换到一半后失败;
    2. `set_lora` 在内存受限部署上先物化全部 offload 层导致 OOM。
    """

    for module_name in ("transformer", "transformer_2"):
        module = self.modules.get(module_name)
        if module is None:
            continue
        for name, layer in module.named_modules():
            if not self.is_target_layer(name):
                continue
            params = dict(layer.named_parameters(recurse=False))
            if "weight" not in params and "qweight" in params:
                raise ValueError(
                    f"LoRA 不支持 {module_name}.{name}: 其权重以 packed "
                    "形式存储 (GGUF), 无法合并或叠加 adapter。请改为加载 "
                    "未量化 checkpoint 以使用 LoRA。"
                )

def set_lora(self, lora_name: str, lora_path: str) -> None:
    # ... 前面的参数校验逻辑省略 ...
    # 先校验再关闭 offload: offload 物化会把每层都搬到显存,
    # 在内存受限部署上会 OOM 而不是返回“不支持 LoRA”的错误。
    self._reject_lora_on_packed_weights()
```

Disable layerwise offload before convert_to_lora_layers ...

评论区精华

唯一的实质 review 交锋围绕 FP32 精度契约展开。

zijiexia: `params_dtype` 按 `adaln_curve_grid` 切到 `_FP32_DTYPE` 会让 pruned checkpoint 转入 FP32, 且静默退出只接受 BF16 的 fused in-place modulation 内核, 这是否为有意行为?

mickqian: 这是有意的。发布的 pruned 实现 (ComfyUI 源码) 将曲线表和缩减的 AdaLN 投影都保持在 FP32, 我们在 bdc6e6 显式化了该精度契约并加了回归测试。隔离 BF16 A/B (1x GB300, 1344x768, 107 帧, 50 步) 显示 BF16 可启用 fused modulation 路径, denoise 从 95.77 秒降到 84.69 秒 (-11.6%), 即 FP32 是保真优先、性能让步的取舍。

该线程最终以设计确认并附带回归测试收尾, 结论是刻意保留精度岛而非缺陷。

- pruned AdaLN 检查点的 FP32 精度契约 (design): mickqian 确认有意为之: 对齐发布的 ComfyUI 修剪实现 (曲线表与缩减 AdaLN 投影均保持 FP32), 已显式化精度契约并添加回归测试; 隔离 BF16 A/B 显示 denoise 95.77 秒 → 84.69 秒 (-11.6%), 即 FP32 是保真优先的性能让步。

风险与影响

- 风险:
 1. 未端到端验证路径: ref2va、Q4_K/Q4_K_M 之外的量化类型、多 GPU GGUF 均未做端到端验证, TP packed 切片只有单元测试覆盖; PR 中提到的 2x H100 验证机位仍在基础设施 provisioning 中, 属于已知未闭环项。
 2. 软依赖新增: 读取 GGUF 需要 gguf 包, `_gguf_module()` 在缺失时抛出 `ImportError`; 若部署环境的依赖声明未同步更新, 用户会在加载点才收到错误。
 3. 精度契约变更: pruned 检查点的 FP32 精度岛会绕过 BF16 fused modulation, 端到端性能下降约 11.6%; 这是有意的对齐行为, 但容易被后续开发者误判为遗漏并 "修复"。
 4. 校验依赖外部表: GGML_QUANT_SIZES、SRT 的 DEQUANT_TYPES/UNQUANTIZED_TYPES 与 _SUPER_BLOCK_DEQUANT_TYPES 三套集合需要与 GGUF 规范及 SRT 内核同步演进, 否则可能产生假阳性或假阴性校验。
 5. LoRA 兼容性边界: GGUF packed 权重不支持 LoRA, 已通过显式拒绝避免半转换状态, 但用户必须切换到未量化 checkpoint 才能使用 LoRA, 属于功能性限制。- 影响: 用户侧: MiniMax-H3 现在可通过 `--transformer-weights-path` 加载 GGUF DiT, 下载体积与主机内存占用显著下降 (61.7 GiB → 17.5 GiB), 配合 layerwise offload 可在消费级 GPU (如 RTX 5090, 峰值 26.3 GiB) 运行。系统侧: 加载构建时间从约 19 秒降到约 1 秒, 集群冷启动成本下降; 量化 DiT 在长序列上避免 fused MMVQ/MMQ 回归 (H200 上 256 token 从 2.60 ms 降到 0.37 ms)。团队侧: 确立了 diffusion 复用 SRT GGUF 能力的边界与模式, Hub 解析与反量化内核不再重复实现; 同时新增了 LoRA、多 GPU 等限制项, 需要文档持续维护。- 风险标记: GGUF 多 GPU 路径未端到端验证, 新增 gguf 包软依赖, FP32 精度岛牺牲约 11.6% 性能, LoRA 与 GGUF 不兼

容显式拒绝，TP 切片与量化块对齐限制，量化类型表需与 GGUF 规范同步

关联脉络

- PR #35618 [diffusion] UX: report where a component's weights are: 同属 diffusion 组件加载与驻留体系，涉及 component loader 与 layerwise offload，与本 PR 的 GGUF 加载 + offload 交互直接相关。
- PR #35614 [diffusion] chore: reduce per-request log noise: 同属 diffusion 运行管线（denoising、stage）的维护脉络，本 PR 对 transformer 加载路径的扩展与其相邻。
- PR #35612 [diffusion] fix: keep Cosmos3 T=1 fusion on Blackwell: 同属 diffusion 内核融合相关问题，与本文 fused MMVQ/MMQ 内核取舍的讨论一脉相承。