

PR #29911 完整报告

sgl-project/sglang

[XPU] Remove redundant xpu graph backend and make xpu graph opt-in by default

合并时间: 2026-07-03 15:58

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29911>

执行摘要

- 一句话: 移除冗余 XPU Graph 后端并默认禁用
- 推荐动作: 值得精读。本 PR 展示了在多硬件后端中如何安全地废弃旧实现、调整默认配置以及处理不同设备的 API 兼容性问题。server_args.py 中的 _lock 机制和 _handle_xpu_backends 的配置回退逻辑, 是构建健壮配置系统的良好参考。同时, mem_get_info 的 try-except 模式值得在类似场景下复用。

功能与动机

PR body 说明: Fix merge conflicts with #23180 and disable XPU Graph by default. 原有 XPU CUDA Graph Backend 与重构后的通用框架存在冲突, 且默认启用 XPU Graph 在部分设备上启动时容易因捕获失败而崩溃, 因此决定移除冗余后端并将 XPU Graph 改为 opt-in。

实现拆解

1. 删除冗余的 XPU Graph 后端: 移除 xpu_cudagraph_backend.py 文件及其 XPU CUDA Graph Backend 类, 并从 runner_backend/utils.py 的 resolve_decode_backend 函数中移除对应的导入和返回分支。
2. 默认禁用 XPU Decode Graph: 在 server_args.py 的 _handle_xpu_backends 方法中, 当 (Phase.DECODE, "backend") 未被用户锁定时, 将 cuda_graph_config.decode.backend 设为 Backend.DISABLED。仅当用户通过 --cuda-graph-backend-decode 或 --cuda-graph-config 显式指定时才会启用。
3. 改进 XPU 内存查询: 在 multimodal_gen/runtime/platforms/xpu.py 和 srt/utils/common.py 的 get_available_gpu_memory 方法中, 使用 torch.xpu.mem_get_info 获取真实空闲内存, 并添加 try-except 回退: 若设备不支持 mem_get_info, 则回落至 get_device_properties 与 memory_allocated 的计算方式。
4. 解除 XPU Graph Runner 的限制: 在 xpu_graph_runner.py 中移除对 SpeculativeAlgorithm 的 import 以及相应的断言, 使 XPU Graph Runner 可以支持投机推断。
5. 更新文档与测试: 在 xpu.mdx 文档中明确 XPU Graph 为 opt-in, 需添加 --cuda-graph-backend-decode full 参数; 测试用例同步调整为不依赖默认 Graph 捕获。

关键文件:

- python/sglang/srt/hardware_backend/xpu/xpu_cudagraph_backend.py (模块 XPU 后端；类别 source；类型 deletion；符号 XPUCudaGraphBackend, init, capture_session, capture_one)：核心变更：完全删除 106 行的 XPUCudaGraphBackend 类，该类的职责已被 XL 端通用框架取代。
- python/sglang/multimodal_gen/runtime/platforms/xpu.py (模块 多模态平台；类别 source；类型 core-logic；符号 get_available_gpu_memory)：关键逻辑调整：改进 get_available_gpu_memory，使用 try-except 包裹 mem_get_info 调用，防止不支持的设备崩溃。
- python/sglang/srt/server_args.py (模块 配置层；类别 source；类型 core-logic；符号 _handle_xpu_backends)：配置层核心：在 _handle_xpu_backends 中新增逻辑，默认禁用 decode graph，实现 opt-in 行为。
- python/sglang/srt/utils/common.py (模块 内存工具；类别 source；类型 core-logic；符号 get_available_gpu_memory)：通用内存工具：同步修改 get_available_gpu_memory 的 XPU 分支，加入 try-except 回退，与多模态平台保持一致。
- python/sglang/srt/model_executor/runner_backend/utils.py (模块 解码路由；类别 source；类型 data-contract；符号 resolve_decode_backend)：数据契约变更：移除对 XPUCudaGraphBackend 的导入和返回，简化后端路由。

关键符号：XPUCudaGraphBackend, _handle_xpu_backends, get_available_gpu_memory, resolve_decode_backend

关键源码片段

python/sglang/multimodal_gen/runtime/platforms/xpu.py

关键逻辑调整：改进 get_available_gpu_memory，使用 try-except 包裹 mem_get_info 调用，防止不支持的设备崩溃。

```

if empty_cache:
    torch.xpu.empty_cache()

# 使用 mem_get_info() 并带有一个合理性上限，以避免 KV-cache 过度分配
# 因为某些驱动错误地将总内存报告为空闲内存。
# 与回退一致：free = max(0, total - allocated)。
try:
    free_gpu_memory, total_gpu_memory = torch.xpu.mem_get_info(device_id)
    used_memory = float(torch.xpu.memory_allocated(device_id))
    free_gpu_memory = min(
        float(free_gpu_memory),
        max(0.0, float(total_gpu_memory) - used_memory),
    )
except Exception:
    # 对于不支持查询空闲内存的设备 / 驱动，使用回退方式
    used_memory = float(torch.xpu.memory_allocated(device_id))
    total_gpu_memory = float(
        torch.xpu.get_device_properties(device_id).total_memory
    )

```

```
free_gpu_memory = max(0.0, total_gpu_memory - used_memory)
```

python/sglang/srt/server_args.py

配置层核心：在 `_handle_xpu_backends` 中新增逻辑，默认禁用 decode graph，实现 opt-in 行为。

```
def _handle_xpu_backends(self):
    if self.device == "xpu":
        # Decode graph is opt-in on XPU: unless the user explicitly set
        # --cuda-graph-backend-decode (or --cuda-graph-config), keep it
        # disabled so the default startup doesn't require graph capture.
        if (Phase.DECODE, "backend") not in self._cuda_graph_config_locked:
            self.cuda_graph_config.decode.backend = Backend.DISABLED
        elif self.cuda_graph_config.decode.backend not in (
            Backend.DISABLED,
            Backend.FULL,
        ):
            logger.warning(
                "XPU platform only supports decode backend 'full'; "
                "disabling unsupported decode backend '%s'.",
                self.cuda_graph_config.decode.backend,
            )
            self.cuda_graph_config.decode.backend = Backend.DISABLED

        # Prefill 端仅支持 tc_pieewise 后端
        if self.cuda_graph_config.prefill.backend not in (
            Backend.DISABLED,
            Backend.TC_PIECEWISE,
        ):
            logger.warning(
                "XPU platform currently only supports prefill tc_pieewise CUDA graph; "
                "disabling unsupported prefill backend."
            )
            self.cuda_graph_config.prefill.backend = Backend.DISABLED
```

评论区精华

- 内存查询回退必要性：gemini-code-assist[bot] 指出部分 Intel XPU 设备（如 Arc 或集显）不支持 `torch.xpu.mem_get_info`，会直接抛出 `RuntimeError`，建议增加 try-except 回退。PR 作者 CaoE 采纳并修复，提升了跨设备兼容性。
- 内存查询兼容性问题：mem_get_info 需要 try-except 回退 (correctness)：作者采纳建议，添加 try-except 回退到 get_device_properties + memory_allocated 的计算方式。

风险与影响

- 风险：
 - 性能回退：默认禁用 decode graph 后，未显式启用的 XPU 用户将无法获得 CUDA Graph 带来的性能提升，特别是长序列 decode 场景。需通过文档引导用户手动开启。

- `mem_get_info` 兼容性: `try-except` 捕获的异常可能掩盖真实问题（如驱动错误），但回退逻辑合理，风险较低。
- 配置锁定冲突: 若用户通过 `--cuda-graph-config` 明确指定 `decode backend`，但指定了 XPU 不支持的 `backend` 值（非 `FULL` 或 `DISABLED`），PR 中会发出警告并强制设为 `DISABLED`；但用户期望可能被忽略，需文档说明。
- 影响:
 - 对 XPU 用户: 默认启动不再执行 Graph 捕获，启动时间缩短且避免崩溃；但需要手动添加 `--cuda-graph-backend-decode full` 以获得最佳性能。影响面中等。
 - 对系统维护: 移除了 106 行重复代码，`resolve_decode_backend` 路由更加清晰。测试文件中的 `import` 和对 `backend` 的引用同步更新。
 - 对其他硬件: 无影响，只修改了 XPU 相关分支和通用工具函数（`common.py`）。
 - 风险标记: 默认性能回退，`mem_get_info` 兼容性

关联脉络

- PR #23180 Unknown (conflict base): 本 PR 旨在解决与 #23180 的合并冲突，并调整默认行为。