

PR #51259 完整报告

vllm-project/vllm

[Bugfix] Import each packed IPC export once on the consumer side

合并时间: 2026-08-11 10:05

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51259>

执行摘要

- 一句话: 单槽导入缓存修复 packed IPC 多 chunk 传输的 producer 侧显存永久泄漏
- 推荐动作: 值得精读。这是对 torch 跨进程 IPC refcount 语义的一次深入剖析, 展示了 '一次导出只能配对一次释放' 的系统约束, 以及用单槽缓存 + 精确释放时机解决协议失配的简洁设计。建议关注两个可改进点: 把 fresh-export 假设固化为仓库内测试 (当前只有临时 harness), 以及在 PackedBufferImporter 中补充 key 不匹配时的显式替换语义注释, 防止后续调用方误用导致双进程计数组失衡。

功能与动机

issue #51258 定位了该泄漏: torch 的跨进程 IPC refcount 把一次导出与一次消费端 rebuild-release 配对, 而 packed_ipc_consumer 对每个 chunk 都重建并释放一次映射, 'A multi-chunk transfer therefore drives the counter negative (1 - n_chunks), and the staging buffer becomes permanently unreclaimable once the producer drops it'. 在 RL 训练循环中, 'Every packed weight transfer strands packed_buffer_size_bytes on the producer side... until the colocated inference engine OOMs on wake'. 该 per-chunk 释放模式自 #37476 引入 chunked packed 传输起就存在, 影响所有后续 release, 因此需要在消费端把 '一次导出' 与 '一次释放' 重新配对。

实现拆解

1. 根因确认 (issue 阶段完成): 通过直接读取共享内存 refcount slot 验证计数轨迹 (导出 $\rightarrow 1$, 三次消费后 $\rightarrow -2$), 并用真实 packed_ipc_producer / packed_ipc_consumer 双进程跑 4 次传输 $\times$ 3 chunks (每次 64 MB 新 buffer), producer 侧 memory_allocated 按 0 / 64 / 128 / 192 MB 线性增长。由此确定修复必须满足约束: 一次 reduce_tensor 导出只能对应一次消费端释放。
2. packed_tensor.py 新增 PackedBufferImporter: 单槽缓存, 以 tuple(list_args) 为 key 保存已打开的映射; rebuild() 命中 key 时直接返回缓存张量 (不再触发每 chunk 的 cudaIpcOpenMemHandle / close 周期), 未命中才调用 torch.multiprocessing.reductions.rebuild_cuda_tensor 并缓存; 替换条目或 close() 即丢弃旧张量引用, 恰好完成一次释放。由于映射别名 producer 内存, 同一导出的后续 chunk 总能读到当前 chunk 的字节, 原有 per-tensor clone 保证返回数据与 producer 后续写入隔离。

3. `packed_ipc_consumer` 接入 `importer`: 新增可选参数 `importer`, `None` 时创建临时 `PackedBufferImporter`, 保持单 `chunk` 调用方 (现有单元测试) 的旧行为; 函数体内原先的 `rebuild_cuda_tensor` 调用改为 `importer.rebuild(list_args)`, 同时删除了原先每 `chunk` 重建时的冗余句柄开合开销。
4. `ipc_engine.py` 持有与释放: `IPCWeightTransferEngine.__init__` 创建 `self._packed_importer`; `receive_weights` 将其传入 `packed_ipc_consumer`; `finish_weight_update` 在 `finalize_layerwise_reload` 之后调用 `close()`, 让 `trainer` 在两次更新之间回收 `buffer`; `shutdown` 也调用 `close()`。该设计依赖 ' 每次 `reduce_tensor` 都是带独立 `refcount` 槽的新导出 ' 这一经验结论, 从而保证每次更新的释放恰好平衡 `fresh-buffer` 与 `persistent-buffer` 两类 `producer`。
5. 测试与验证配套: 本 PR 未向仓库提交测试文件 (2 个变更文件均为源码); 验证依靠 `issue` 中的最小复现 / `refcount` 读取脚本与作者自建的临时双进程 `harness` (4 transfers × 3 chunks, 64 MB `fresh buffer`, H100 + torch 2.11)。aoshen02 在本地确认 `tests/distributed/test_packed_tensor.py` 通过, 并将 CI 失败归因于无关的 `test_compile_correctness` (已由 #51271 修复), `rebase` 后 CI #83144 转绿。

关键文件:

- `vllm/distributed/weight_transfer/packed_tensor.py` (模块 权重传输; 类别 `source`; 类型 `core-logic`; 符号 `PackedBufferImporter`, `PackedBufferImporter.init`, `PackedBufferImporter.rebuild`, `PackedBufferImporter.close`): 核心修复文件: 新增 `PackedBufferImporter` 单槽导入缓存, `packed_ipc_consumer` 增加 `importer` 参数, 使一次 `reduce_tensor` 导出只对应一次消费端重建 - 释放, 解决多 `chunk` 传输把跨进程 `refcount` 压成负数导致的 `producer` 侧显存永久泄漏; 同时删除每 `chunk` 的 `cudaIpcOpenMemHandle` / `close` 冗余开销。
- `vllm/distributed/weight_transfer/ipc_engine.py` (模块 权重传输; 类别 `source`; 类型 `core-logic`; 符号 `IPCWeightTransferEngine.init`, `IPCWeightTransferEngine.receive_weights`, `IPCWeightTransferEngine.finish_weight_update`, `IPCWeightTransferEngine.shutdown`): 集成与释放时机: `IPCWeightTransferEngine` 持有唯一的 `PackedBufferImporter` 实例, 经 `receive_weights` 传给 `packed_ipc_consumer`, 并在 `finish_weight_update` (`finalize_layerwise_reload` 之后) 与 `shutdown` 时 `close()` 释放, 是 ' 每次 `update` 恰好释放一次 ' 的落点。

关键符号: `PackedBufferImporter.init`, `PackedBufferImporter.rebuild`, `PackedBufferImporter.close`, `packed_ipc_consumer`, `IPCWeightTransferEngine.init`, `IPCWeightTransferEngine.receive_weights`, `IPCWeightTransferEngine.finish_weight_update`, `IPCWeightTransferEngine.shutdown`

关键源码片段

`vllm/distributed/weight_transfer/packed_tensor.py`

核心修复文件: 新增 `PackedBufferImporter` 单槽导入缓存, `packed_ipc_consumer` 增加 `importer` 参数, 使一次 `reduce_tensor` 导出只对应一次消费端重建 - 释放, 解决多 `chunk` 传输把跨进程 `refcount` 压成负数导致的 `producer` 侧显存永久泄漏; 同时删除每 `chunk` 的

cudaIpcOpenMemHandle / close 冗余开销。

vllm/distributed/weight_transfer/packed_tensor.py (整理后)

```
class PackedBufferImporter:
```

```
    """单槽导入缓存：让一次 reduce_tensor 导出只对应一次消费端重建释放周期。
```

torch 的跨进程 IPC refcount 把一次导出与一次消费端 rebuild-release 配对：
导出把共享计数置为 1，消费端释放 rebuilt 张量时减一；producer 只有在计数恰好为 0 时才能通过 torch.cuda.ipc_collect() 回收已丢弃的 buffer。

packed_ipc_producer 对同一 staging buffer 只导出一次，却把同一份 rebuild args 随每个 chunk 发出；若每个 chunk 都重建并释放映射，计数会被压成负数 (1 - n_chunks)，producer 侧的 buffer 将永久无法回收。

这里以 rebuild args 为 key 缓存已打开的映射：同一导出的多个 chunk 复用该映射（它别名 producer 内存，因此每个 chunk 都能读到当前 chunk 的字节），换入下一次导出或调用 close() 时才释放上一次映射——每次导出恰好释放一次。

```
    """
```

```
def __init__(self) -> None:
```

```
    # 缓存项为 (rebuild_args_tuple, packed_tensor), None 表示无缓存
    self._entry: tuple[tuple, torch.Tensor] | None = None
```

```
def rebuild(self, list_args: list) -> torch.Tensor:
```

```
    from torch.multiprocessing.reductions import rebuild_cuda_tensor
```

```
    key = tuple(list_args) # 同一导出的每个 chunk 携带完全相同的 args
```

```
    if self._entry is not None and self._entry[0] == key:
```

```
        # 命中缓存：直接复用已有映射，不再触发新一次打开与释放
```

```
        return self._entry[1]
```

```
    packed = rebuild_cuda_tensor(*list_args)
```

```
    # 替换旧条目即隐式释放上一个 mapping (恰好一次 release)
```

```
    self._entry = (key, packed)
```

```
    return packed
```

```
def close(self) -> None:
```

```
    # 置空条目即丢弃张量引用，释放时对应 producer 侧一次 refcount 减一
```

```
    self._entry = None
```

```
def packed_ipc_consumer(
```

```
    ipc_handle: dict[str, tuple],
```

```
    names: list[str],
```

```
    shapes: list[list[int]],
```

```
    dtype_names: list[str],
```

```
    tensor_sizes: list[int],
```

```
    device_index: int,
```

```
    importer: "PackedBufferImporter | None" = None,
```

```
) -> list[tuple[str, torch.Tensor]]:
```

"""把一个 packed IPC chunk 解包成命名张量。

importer 必须被同一导出的所有 chunk 共享；为 None 时创建临时实例，临时实例无法在 chunk 间复用，只对单 chunk 调用方（如原有单元测试）安全。

"""

```
props = torch.cuda.get_device_properties(device_index)
physical_gpu_id = str(props.uuid)
if physical_gpu_id not in ipc_handle:
    raise ValueError(
        f"IPC handle not found for GPU UUID {physical_gpu_id}. "
        f"Available UUIDs: {list(ipc_handle.keys())}"
    )

args = ipc_handle[physical_gpu_id]
list_args = list(args)
list_args[6] = device_index # 覆盖为接收端设备索引，也是缓存 key 的一部分

if importer is None:
    importer = PackedBufferImporter()
packed = importer.rebuild(list_args) # 命中缓存则不再打开新的 IPC 映射

content_size = sum(tensor_sizes)
packed = packed[:content_size] # 共享 staging buffer 只取本 chunk 的内容段

dtypes = [getattr(torch, dn) for dn in dtype_names]
# 每个张量必须 clone 成独立存储：映射别名 producer 内存，且 producer 在
# 生成器恢复后会写入下一个 chunk；layerwise reload 会缓冲这些结果供
# _layerwise_process 重放，若保留视图将读到下一个 chunk 的字节
return [
    (name, t.clone())
    for name, t in unpack_tensor(packed, names, shapes, dtypes, tensor_sizes)
]
```

vllm/distributed/weight_transfer/ipc_engine.py

集成与释放时机：IPCWeightTransferEngine 持有唯一的 PackedBufferImporter 实例，经 receive_weights 传给 packed_ipc_consumer，并在 finish_weight_update（finalize_layerwise_reload 之后）与 shutdown 时 close() 释放，是 ' 每次 update 恰好释放一次 ' 的落点。

vllm/distributed/weight_transfer/ipc_engine.py（整理后）

```
class IPCWeightTransferEngine(WeightTransferEngine[...]):
    def __init__(self, config, vllm_config, device, model) -> None:
        super().__init__(config, vllm_config, device, model)
        # 由训练端握手信息设置的 wire 参数，默认值仅用于先接收后初始化的场景
        self.packed = False
        # 每个引擎持有唯一的导入缓存，跨同一导出的所有 chunk 复用；
        # 这是 refcount 配对关系的持有者（见 PackedBufferImporter 文档）
        self._packed_importer = PackedBufferImporter()
```

```

def receive_weights(self, update_info: IPCWeightTransferUpdateInfo) -> None:
    # 以模型所在设备为准重建 IPC 张量，而非环境当前设备
    device_index = self.device.index
    if self.packed:
        if update_info.tensor_sizes is None:
            raise ValueError("`tensor_sizes` is required when packed=True")
        assert isinstance(update_info.ipc_handles, dict)
        weights = packed_ipc_consumer(
            ipc_handle=update_info.ipc_handles,
            names=update_info.names,
            shapes=update_info.shapes,
            dtype_names=update_info.dtype_names,
            tensor_sizes=update_info.tensor_sizes,
            device_index=device_index,
            importer=self._packed_importer, # 同一导出的所有 chunk 共享
        )
    else:
        # 逐张量 IPC 路径：每个 tensor 一次 rebuild，不经过 importer
        ...
    with disable_mtp_completeness_check():
        self.model.load_weights(weights)

def finish_weight_update(self) -> None:
    finalize_layerwise_reload(self.model, self.model_config)
    # 每次 reduce_tensor 都是带独立 refcount 槽的新导出，因此在每次更新
    # 结束时释放一次恰好平衡本次导出，trainer 可在两次更新之间回收其
    # staging buffer；跳过 finish 的调用方由 importer 的换入替换路径覆盖
    self._packed_importer.close()

def shutdown(self) -> None:
    self._packed_importer.close()

```

评论区精华

该 PR 的 review 评论为 0 条（fork 提交，claude[bot] 的自动 review 被禁用），核心讨论集中在 issue 评论与 PR body：

- aoshen02 追踪了 Buildkite #83061 两个失败 job（distributed-compile-plus-comm-4-gpus、distributed-compile-plus-rpc-tests-2-gpus），根因是 compile/fullgraph/test_basic_correctness.py::test_compile_correctness，已由 main 上的 #51271 修复：'The original packed IPC change from this PR is still valid; I rebased and signed it locally... could not push directly to acmore:fix/packed-ipc-import-once'。acmore rebase 到最新 main 后 CI #83144 转绿，aoshen02 给出 'LGTM'。
- 设计层面的关键论据在 PR body：'Every reduce_tensor call is a fresh export with its own refcount slot (verified empirically)'，因此每次 update 的 close() 必然平衡本次导出；importer=None 走临时实例以保持单 chunk 调用方兼容，并在 docstring 中明确警告 '只对单 chunk 消费安全'。

- CI 失败根因：无关的编译测试与 rebase 要求 (other): acmore 将分支 rebase 到最新 main (cf8f3a3bb) 后，重新触发的 CI #83144 转绿，aoshen02 给出 LGTM。
- fresh-export 假设：修复正确性的关键前提 (correctness): 该假设被接受为当前设计前提，但未固化为仓库内测试断言；后续若 producer 改为复用导出，需重新评估释放时机。

风险与影响

- 风险：
 - 正确性依赖经验假设：' 每次 reduce_tensor 都是带独立 refcount 槽的新导出 ' 是修复成立的前提。若未来某个 producer 改为复用同一导出跨多次 update（不重新 reduce_tensor），finish_weight_update 中的 close() 会释放一个仍被下一次 chunk 使用的映射，而 rebuild 又会按相同 key 重新打开，可能造成 refcount 与句柄数量失衡。该前提目前只有实证，未固化进代码断言或仓库测试。
 - 缺少仓库内回归测试：PR 只改源码、未新增测试文件，双进程 harness 是临时验证手段；后续改动 importer 或 consumer 时 CI 无法拦住同类回归。
 - 内存别名语义脆弱：缓存映射直接别名 producer 内存，正确性依赖每个 chunk 都执行 per-tensor clone。packed_tensor.py 的 docstring 明确警告 layerwise reload 会缓冲 bound_args 供 _layerwise_process 重放，若未来省略 clone 或把 importer 传给其他调用方，会读到下一个 chunk 的字节造成静默权重损坏。
 - 释放时序：close() 发生在 finalize_layerwise_reload 之后，依赖 load_weights 已同步完成；若将来出现异步加载路径，过早释放会导致映射失效。
- 影响：
 - 用户侧：修复 RL / 连续权重同步场景下 producer (trainer) 侧随每步权重同步累积的显存泄漏，消除同驻推理引擎因 staging buffer 持续滞留而 OOM 的风险；每次更新不再泄漏默认 64 MB (packed_buffer_size_bytes)。
 - 系统侧：仅影响 vllm/distributed/weight_transfer 的 packed IPC 路径 (IPCWeightTransferEngine)，非 packed 的逐张量 IPC 路径不变；packed_ipc_consumer 签名向后兼容，importer 为可选参数，现有单测不受影响。
 - 性能侧：顺带消除每 chunk 重复执行 cudaIpcOpenMemHandle / close 的开销，一次导出只打开一次映射。
 - 工程侧：改动集中在 2 个文件 59 行，以参数注入方式扩展，未引入新依赖，无 wire 协议变更。
 - 风险标记：核心分布式路径变更，正确性依赖经验假设，缺少仓库内回归测试，内存别名语义依赖 clone

关联脉络

- PR #49519 [Bugfix][Model Loader] Defer post-load attention weight processing: 同属模型加载与重载生命周期链路：packed IPC 消费发生在 layerwise reload 中，packed_tensor.py 的 docstring 明确提到 layerwise reload 会缓冲 bound_args 供 _layerwise_process 重放，本 PR 的映射复用与 clone 语义直接服务于该路径。

- PR #51271 [CI] Remove compile correctness test from 2/4-GPU suites: 评论区说明该 PR 修复了本 PR CI 失败的根因（移除了 2/4-GPU 套件中的 test_compile_correctness），本 PR 依赖它的合入才能让 CI 转绿；精确标题未在材料中给出，此处为按评论内容推断，仅属 CI 依赖。
- PR #37476 Chunked packed weight transfer: issue #51258 指出 chunked packed 传输由 #37476 引入，per-chunk 释放的消费端模式自此存在，是本次泄漏的起源；不在近期历史 PR 列表中，标题为按 issue 描述的推断。