

PR #49341 完整报告

vllm-project/vllm

[Rust Frontend] Send multimodal tensors in auxiliary frames

合并时间: 2026-07-29 20:21

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49341>

Rust Frontend: 多模态张量通过辅助帧零拷贝发送

执行摘要

本 PR 改造 Rust 前端 `EngineCoreRequest::Add` 的多模态张量发送路径, 利用 ZMQ multipart 辅助帧将大张量移出主 msgpack 负载, 消除两次全量拷贝, 使大张量编码时间减少约 3x, 峰值 RSS 降低 33%-50%。核心设计包括使用 `Bytes::from_owner` 的零拷贝包装 `PodVec`、递归提取方法以及阈值可配置支持。

功能与动机

现有实现中, 多模态张量 (图像 / 视频) 在 Rust 前端经过“typed buffer -> raw bytes -> msgpack 序列化 -> ZMQ 发送”的多次全量复制。对于大图或多图请求, 这导致显著的延迟和内存峰值。虽然底层 wire 协议早已支持 `AuxIndex` 标记 (以索引引用辅助帧), 但 Rust 出站路径从未真正生成过辅助帧。本 PR 完成该缺失的路径, 本质上是“追赶上 Python 端已有的辅助帧支持”。

实现拆解

1. 零拷贝张量构造(protocol/tensor.rs): 新增 `PodVec<T>` 包装器, 实现 `AsRef<[u8]>`, 通过 `cast_slice` 零成本转换。`bytes_from_pod_vec` 函数使用 `Bytes::from_owner` 将 `typed Vec` 直接转换为 `Bytes`, 避免中间复制。所有 `typed` 构造器 (`from_f32`、`from_f16`、`from_i64`、`from_u32`、`from_bf16`) 改为接收 `Vec<T>` 所有权, 并统一委托给 `from_raw_bytes`。
2. 张量提取与替换(protocol/multimodal.rs): 为 `MmFeatureSpec`、`MmFieldElem`、`MmKwargValue` 实现递归 `extract_aux_frames` 方法, 遍历字段树, 对大小 $\geq$ 阈值 (默认 256) 的张量: a) 将其数据字段替换为 `WireArrayData::AuxIndex(index)`; b) 将原始 `Bytes` 追加到输出向量。同时处理 `mm_position.is_embed` 等嵌套场景。
3. 专用发送路径(client/imp.rs): 新增 `send_request_to_engine` 方法, 先调用 `extract_aux_frames` 提取, 再编码残余主负载, 最后调用 `send_encoded_to_engine` 以 multipart 形式发送 (首帧为请求类型 + 主负载, 后续帧为辅助张量数据)。旧 `send_to_engine` 保留用于非请求消息。
4. 阈值可配置: 通过环境变量 `VLLM_MSGPACK_ZERO_COPY_THRESHOLD` 控制, 默认 256, 与 Python 端保持相同默认值。
5. 测试覆盖: 在 `protocol/request.rs` 中添加单元测试 `engine_core_request_extracts_large_nested_tensors_in_wire_order`, 验证内联 / 辅助帧

边界、AuxIndex 索引顺序及缓冲区零拷贝一致性。在 tests/client.rs 中添加集成测试 client_sends_large_multimodal_tensor_as_aux_frame，模拟引擎接收 multipart 消息，验证张量数据正确分离。

rust/src/engine-core-client/src/protocol/tensor.rs

core logic: 新增 PodVec 包装器和 bytes_from_pod_vec 实现零拷贝张量字节转换，重构所有 from_* 构造器以直接获取所有权并调用 from_raw_bytes，是消除副本的核心基础。

```
// SPDX-License-Identifier: Apache-2.0
// SPDX-FileCopyrightText: Copyright contributors to the vLLM project

use bytemuck::{Pod, cast_slice};
use bytes::Bytes;

/// Wrapper that lets a `Vec<T>` be used as `AsRef<[u8]>` via zero-copy cast.
/// 利用 bytemuck 的 `cast_slice` 将 typed 数组直接转为字节切片，无需复制。
struct PodVec<T: Pod> (Vec<T>);

impl<T: Pod> AsRef<[u8]> for PodVec<T> {
    fn as_ref(&self) -> &[u8] {
        cast_slice(&self.0)
    }
}

/// Convert a typed `Vec<T>` into `Bytes` without copying the underlying buffer.
/// 所有权由 `PodVec` 包裹，`Bytes::from_owner` 负责内存管理，发送完成后释放。
fn bytes_from_pod_vec<T>(data: Vec<T>) -> Bytes
where
    T: Pod + Send + 'static,
{
    Bytes::from_owner(PodVec(data))
}

// Example: from_f32 now takes ownership and uses bytes_from_pod_vec
impl WireNdArray {
    /// Build a float32 tensor backed by native-endian raw-view bytes.
    /// Takes ownership of the backing buffer without copying its data.
    pub fn from_f32(shape: Vec<usize>, data: Vec<f32>) -> Result<Self, String> {
        validate_element_count(&shape, data.len())?;
        Ok(Self::from_raw_bytes(
            "float32",
            shape,
            bytes_from_pod_vec(data),
        ))
    }
}
// 类似地，from_f16, from_bf16, from_i64, from_u32 均改为相同模式。
}
```

rust/src/engine-core-client/src/client/imp.rs

core logic: 新增 msgpack_zero_copy_threshold 配置和 send_request_to_engine 方法, 分离普通消息与请求消息的发送路径; 实现 multipart ZMQ 发送。

```
// 阈值配置: 环境变量 VLLM_MSGPACK_ZERO_COPY_THRESHOLD, 默认 256 字节
const MSGPACK_ZERO_COPY_THRESHOLD_ENV: &str = "VLLM_MSGPACK_ZERO_COPY_THRESHOLD";
const DEFAULT_MSGPACK_ZERO_COPY_THRESHOLD: usize = 256;

fn msgpack_zero_copy_threshold() -> usize {
    std::env::var(MSGPACK_ZERO_COPY_THRESHOLD_ENV)
        .ok()
        .and_then(|value| value.parse().ok())
        .unwrap_or(DEFAULT_MSGPACK_ZERO_COPY_THRESHOLD)
}

impl ClientInner {
    /// 发送 Add 请求, 将大张量移出主负载作为辅助帧。
    pub async fn send_request_to_engine(
        &self,
        engine_id: &EngineId,
        mut payload: EngineCoreRequest,
    ) -> Result<> {
        // 1. 提取超过阈值的张量 buffers, 替换为 AuxIndex
        let aux_frames = payload.extract_aux_frames(self.msgpack_zero_copy_threshold);
        // 2. 对剩下的请求编码 (小张量保持内联)
        let payload = Bytes::from(encode_msgpack(&payload)?);
        // 3. 发送 multipart: 主帧 + 辅助帧
        self.send_encoded_to_engine(engine_id, EngineCoreRequestType::Add, payload, aux_frames)
            .await
    }

    /// 底层 multipart 发送: ZMQ 分帧发送主负载和辅助帧。
    async fn send_encoded_to_engine(
        &self,
        engine_id: &EngineId,
        request_type: EngineCoreRequestType,
        payload: Bytes,
        aux_frames: Vec<Bytes>,
    ) -> Result<> {
        let mut input_send = self.input_send.clone();
        let engine_id = engine_id.clone();
        self.handle.spawn(async move {
            transport::send_multipart_message(
                &mut input_send,
                &engine_id,
                request_type.to_frame(),
                payload,
                aux_frames,
            )
        })
    }
}
```

```

        )
        .await
    }).await.map_err(|_| Error::ClientClosed)?;
    Ok(())
}
}

```

rust/src/engine-core-client/src/protocol/request.rs

core logic: 新增 extract_aux_frames 方法，遍历 mm_features 递归提取大张量；同时添加单元测试验证提取逻辑正确。

```

impl EngineCoreRequest {
    /// Extract large request tensors into ordered auxiliary frames.
    /// 遍历 mm_features，将超过阈值的张量移出主负载，返回辅助帧向量。
    pub(crate) fn extract_aux_frames(&mut self, threshold: usize) -> Vec<Bytes> {
        let mut aux_frames = Vec::new();
        if let Some(features) = &mut self.mm_features {
            for feature in features {
                feature.extract_aux_frames(&mut aux_frames, threshold);
            }
        }
        aux_frames
    }
}

#[cfg(test)]
mod tests {
    #[test]
    fn engine_core_request_extracts_large_nested_tensors_in_wire_order() {
        let inline = vec![1_u8; AUX_FRAME_THRESHOLD - 1]; // 小于阈值，保持内联
        let first_aux = vec![2_u8; AUX_FRAME_THRESHOLD]; // 等于阈值，移出
        let second_aux = vec![3_u8; AUX_FRAME_THRESHOLD + 1]; // 大于阈值，移出
        let first_aux_ptr = first_aux.as_ptr();
        let second_aux_ptr = second_aux.as_ptr();
        let mut request = EngineCoreRequest {
            mm_features: Some(vec![MmFeatureSpec {
                data: Some(BTreeMap::from([
                    ("inline".to_string(), MmFieldElem {
                        data: Some(MmKwargValue::Tensor(WireTensor::from_raw(
                            "uint8", vec![inline.len()], inline,
                        ))),
                        field: MmField::Batched(MmBatchedField { keep_on_cpu: false }),
                    }),
                    ("nested".to_string(), MmFieldElem {
                        data: Some(MmKwargValue::List(vec![
                            MmKwargValue::Int(7),
                            MmKwargValue::Tensor(WireTensor::from_raw(
                                "uint8", vec![first_aux.len()], first_aux,
                            )),
                        ])),
                    })
                ])),
            }])
        }
    }
}

```

```

    )),
    field: MmField::Batched(MmBatchedField { keep_on_cpu: false }),
  )),
  modality: "image".to_string(),
  identifier: "id".to_string(),
  mm_position: PlaceholderRange { /* ... */ },
  mm_hash: None,
  )),
  ..EngineCoreRequest::default()
};

let aux_frames = request.extract_aux_frames(AUX_FRAME_THRESHOLD);

assert_eq!(aux_frames.len(), 2);
// 验证辅助帧与原缓冲区地址一致（零拷贝）
assert_eq!(aux_frames[0].as_ptr(), first_aux_ptr);
assert_eq!(aux_frames[1].as_ptr(), second_aux_ptr);
// 验证内联张量仍是 RawView 而非 AuxIndex
let inline_tensor = /* 获取 inline tensor */;
assert!(matches!(inline_tensor.data, WireArrayData::RawView(_)));
// 验证提取后的张量变为 AuxIndex(1) 和 AuxIndex(2)
let nested_tensor = /* 获取 nested 中 tensor */;
assert_eq!(nested_tensor.data, WireArrayData::AuxIndex(1));
let is_embed = /* 获取 is_embed tensor */;
assert_eq!(is_embed.data, WireArrayData::AuxIndex(2));
}
}

```

评论区精华

- BugenZhao在批准前进行了两处调整：代码风格局部重构，以及将硬编码阈值改为从环境变量 VLLM_MSGPACK_ZERO_COPY_THRESHOLD 加载。这体现了对可配置性的重视，确保用户可根据实际多媒体负载调整阈值，与 Python 端保持一致。
- 无其他实质性讨论或争议。

风险与影响

维度	分析
兼容性	接收端必须支持 <code>AuxIndex</code> 和 multipart ZMQ。Python 引擎已支持，风险低。阈值默认 256 与 Python 端一致。
性能	大张量拷贝时间减少约 3x，内存占用降低 33%-50%。小张量始终内联，无性能退化。
回归	非请求消息仍走原 <code>send_to_engine</code> 路径，不受影响。
安全	<code>Bytes::from_owner</code> 确保 <code>Pod + Send + 'static</code> ，所有权安全。

影响范围：所有使用 Rust 前端的多模态推理请求。团队可通过调整

`VLLM_MSGPACK_ZERO_COPY_THRESHOLD` 控制零拷贝阈值，默认为 256 字节，适合大多数场景。

关联脉络

- 与 #49604 (Add --limit-mm-per-prompt support) 同属 Rust 前端多模态功能线，配合实现对多模态输入的全面性能和功能覆盖。
- 与 #48145 (prefill token reuse) 协同，整体提升 Rust 前端在分离式部署场景下的效率。
- 本 PR 完成了“Python- 侧已有但 Rust 侧缺失”的辅助帧路径，是 Rust 前端日趋成熟的关键一步。