

PR #49585 完整报告

vllm-project/vllm

[EC Connector] Added Build Connector Worker Meta for EC Connector

合并时间: 2026-08-16 06:40

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/49585>

执行摘要

- 一句话: 为 EC Connector 打通 worker 到调度器的元数据通道
- 推荐动作: 值得精读。这是 vLLM EC 分布式缓存架构的关键基础设施 PR, 虽然不直接面向用户, 但其设计决策有普遍借鉴意义: ①“聚合器只做通用归并、屏障语义下放给 connector 实现”的职责边界划分; ② shared 单例的 copy-on-write 防御模式; ③显式拒绝不支持的后端组合而非静默丢数据。建议与 #38390 (V2 EC connector 基础)、#49994 (CUDA events 拷贝完成检测)、#47941 (NIXL 共享) 串读, 理解完整演进链路。

功能与动机

PR #38390 已为 V2 model runner 接入 EC Connector, 但 EC 仍缺少 KV connector 早已具备的 worker->scheduler 元数据通道三件套: `build_connector_worker_meta()`、输出字段、输出聚合器。PR body 明确指出: 没有这条通道, worker 侧 `ECConnector` 无法把每步状态 (如“这个 mm_hash 已持久化”) 回报给 scheduler 侧 connector, 直接阻塞基于 CUDA events 的 CPU offload EC connector 与 NIXL P2P EC 共享。作者原话: “We are retiring the StepTracker in the CPU EC Connector, because using CUDA events is the right way to handle 'copy done' notifications”。评论中还说明这是 #49994 (CUDA events 检查 GPU->CPU 拷贝完成) 与 #47941 (NIXL 基于 EC 共享) 的前置条件。

实现拆解

实现按 6 步拆解:

1. 输出契约扩展 (`vllm/v1/outputs.py`): `ECConnectorOutput` 新增 `ec_connector_worker_meta` 字段与 `is_empty()`; `ModelRunnerOutput` 新增 `with_ec_conn_output()` / `with_ec_conn_output_only()` 两个 copy-on-write 静态方法——目标是共享的 `EMPTY_MODEL_RUNNER_OUTPUT` 单例时先浅拷贝再挂载, 避免把本步状态写进全局单例 (对应 review 中 P1 问题)。
2. 元数据基类 (`vllm/distributed/ec_transfer/ec_connector/base.py`): 新增 `ECConnectorWorkerMetadata` 抽象基类, 声明 `aggregate()` 契约; `ECConnectorBase.build_connector_worker_meta()` 默认返回 `None`, 因此所有现有 connector 行为不变, 只有主动覆盖的 connector (如后续 CPU offload、NIXL 实现) 才走新通道。

3. 聚合器与执行器接线 (`ec_connector/utils.py`、`executor/abstract.py`、`executor/multiproc_executor.py`) : `ECOutputAggregator` 镜像 `KVOutputAggregator`, 把每个 worker 的 `finished_sending` / `finished_recving` 取并集、worker meta 逐个调用抽象 `aggregate()` 归并, 最终 fold 到 `output_rank` 的输出上;
`multiproc_executor.collective_rpc` 的 `_aggregate` 把 KV 与 EC 两个聚合器链式作用于同一份 outputs, 各自字段互不覆盖; `executor/abstract.py` 提供 `init_ec_output_aggregator()` 初始化入口。
4. runner 出口修复 (`vllm/v1/worker/gpu/model_runner.py`、`ec_connector.py`) : MRv2 的 `execute_model` 在 `total_num_scheduled_tokens == 0` 与 `batch_desc.num_tokens == 0` 两个零工作分支调用新增的 `_merge_ec_connector_no_forward`, 让 EC connector 在无 forward 步也执行收发并上报 meta; `sample_tokens()` / `pool()` / encoder-only 分支改用 `with_ec_conn_output`, 修复此前 `ec_connector_output` 计算后被丢弃的问题;
`ExecuteModelState` 新增 `ec_connector_output` 字段在 `execute` 与 `sample` 之间传递。
`ActiveECConnector.no_forward()` 复用 `maybe_get_output` 上下文管理器的收尾逻辑; 同时把 `save_new_caches` 从 `is_ec_producer_only` 修正为 `ec_connector.is_producer` (`ec_both` 节点也要保存, 属附带 bugfix)。
5. 后端约束 (`vllm/v1/executor/ray_executor.py`) : legacy Ray executor + EC connector + `world_size > 1` 时抛 `NotImplementedError`, 明确提示改用 `RayExecutorV2` (`VLLM_USE_RAY_V2_EXECUTOR_BACKEND=1`) 或多进程 executor, 避免静默丢弃其他 worker 的 EC 状态。
6. 测试与 CI 配套: 新增 `test_ec_output_aggregator.py` 覆盖聚合折叠、空上报、单例不写穿、与 KV 聚合器链式组合四类语义; 新增 `test_worker_ec_connector.py` 覆盖 producer 保存、上下文退出上报 meta、no-forward 上报; `test_outputs.py` 补 copy-on-write 测试; `.buildkite/test_areas/misc.yaml` 把新测试挂入 CPU 测试区域。

关键文件:

- `vllm/distributed/ec_transfer/ec_connector/utils.py` (模块 聚合器; 类别 source; 类型 core-logic; 符号 `ECOutputAggregator`, `ECOutputAggregator.aggregate`) : 新增 `ECOutputAggregator`, 是本 PR 的核心聚合逻辑: 把任意 rank 上运行的 EC connector 输出 folded 到 `output_rank` 的输出上, 并处理共享单例的 copy-on-write 防护。
- `vllm/v1/outputs.py` (模块 输出契约; 类别 source; 类型 data-contract; 符号 `ECCConnectorOutput`, `ECCConnectorOutput.is_empty`, `ModelRunnerOutput.with_ec_conn_output`, `ModelRunnerOutput.with_ec_conn_output_only`) : 输出数据契约扩展: `ECCConnectorOutput` 新增 worker meta 字段与 `is_empty()`, `ModelRunnerOutput` 新增 `with_ec_conn_output` 系列 copy-on-write 助手, 是全链路数据传递的基石。
- `vllm/distributed/ec_transfer/ec_connector/base.py` (模块 连接器基类; 类别 source; 类型 core-logic; 符号 `ECCConnectorWorkerMetadata`, `ECCConnectorWorkerMetadata.aggregate`, `ECCConnectorBase.build_connector_worker_meta`) : 新增 `ECCConnectorWorkerMetadata` 抽象基类与 `build_connector_worker_meta()` 默认实现, 定义 worker -> scheduler 元数据通道的契约, 现有 connector 行为不变。

- vllm/v1/worker/gpu/model_runner.py (模块 模型执行器; 类别 source; 类型 data-contract; 符号 GPUModelRunner._merge_ec_connector_no_forward, GPUModelRunner.execute_model, GPUModelRunner.sample_tokens, GPUModelRunner.pool) : MRv2 核心接线点: 零 forward 步骤合并 EC 输出、encoder-only 与 sample/pool 路径修复 ec_connector_output 被丢弃的问题, ExecuteModelState 数据契约新增字段。
- vllm/v1/worker/gpu/ec_connector.py (模块 工作端连接器; 类别 source; 类型 core-logic; 符号 ActiveECCConnector, ActiveECCConnector.no_forward, ActiveECCConnector.maybe_get_output, ECCConnector.no_forward) : worker 侧收出口: ActiveECCConnector.no_forward 让无 forward 步骤也能收发与上报 meta; 修复 ec_both 节点 save_new_caches 判定。
- vllm/v1/executor/multiproc_executor.py (模块 多进程执行器; 类别 source; 类型 core-logic; 符号 MultiprocExecutor.collective_rpc, MultiprocExecutor._aggregate) : collective_rpc 的 _aggregate 把 KV 与 EC 两个聚合器链式作用于同一份 outputs, 保证输出字段互不覆盖, 是多 rank 场景下元数据送达 scheduler 的关键。
- vllm/v1/executor/ray_executor.py (模块 后端执行器; 类别 source; 类型 core-logic; 符号 RayExecutor._init_executor) : 显式拒绝 legacy Ray executor + EC connector + world_size > 1 的组合, 避免其他 worker 的 EC 状态被静默丢弃, 是 review P2 的决策落地。
- vllm/v1/executor/abstract.py (模块 执行器基类; 类别 source; 类型 core-logic; 符号 Executor.init_ec_output_aggregator) : 提供 init_ec_output_aggregator() 与属性声明, 是执行器侧聚合器生命周期的统一入口。
- tests/v1/ec_connector/unit/test_ec_output_aggregator.py (模块 聚合器测试; 类别 test; 类型 test-coverage; 符号 FakeWorkerMeta, test_aggregate_folds_every_rank_onto_output_rank, test_aggregate_leaves_no_ec_output_when_no_worker_reported, test_aggregate_does_not_write_through_the_shared_empty_output) : 覆盖聚合器四类核心语义: 跨 rank 折叠、空上报、共享单例不写穿、与 KV 聚合器链式组合, 是 P1 问题回归防线。
- tests/v1/ec_connector/unit/test_worker_ec_connector.py (模块 连接器测试; 类别 test; 类型 test-coverage; 符号 test_saves_newly_added_caches_for_every_producer, test_worker_meta_is_reported_on_context_exit, test_no_forward_reports_without_running_the_model) : 覆盖 worker 侧行为: ec_both 节点作为 producer 保存、上下文退出后上报 worker meta、no-forward 步仍上报, 直接验证本 PR 的核心新行为。
- tests/v1/test_outputs.py (模块 输出测试; 类别 test; 类型 test-coverage; 符号 test_with_ec_conn_output_copies_shared_empty_output) : 补充 with_ec_conn_output 对共享单例的 copy-on-write 行为测试, 直接对应 review P1 修复。

关键符号: ECTOutputAggregator.aggregate, ECCConnectorBase.build_connector_worker_meta, ECCConnectorWorkerMetadata.aggregate, ModelRunnerOutput.with_ec_conn_output, ModelRunnerOutput.with_ec_conn_output_only, ECCConnectorOutput.is_empty, ActiveECCConnector.no_forward,

ActiveECCConnector.maybe_get_output, GPUModelRunner._merge_ec_connector_no_forward, MultiprocExecutor.collective_rpc, Executor.init_ec_output_aggregator

关键源码片段

vllm/distributed/ec_transfer/ec_connector/utils.py

新增 EOutputAggregator, 是本 PR 的核心聚合逻辑: 把任意 rank 上运行的 EC connector 输出 folded 到 output_rank 的输出上, 并处理共享单例的 copy-on-write 防护。

```
# SPDX-License-Identifier: Apache-2.0
# SPDX-FileCopyrightText: Copyright contributors to the vLLM project
"""EC connector helper utilities."""

from vllm.v1.outputs import ECCConnectorOutput, ModelRunnerOutput

class EOutputAggregator:
    """ 把每个 worker 的 EC connector 输出合并到唯一到达 scheduler
    的 ModelRunnerOutput 上。

    镜像 KVOutputAggregator: 只有 output_rank 的输出会返回给 scheduler,
    但 EC connector 可能在任意 rank 上运行。
    """

    def aggregate(
        self, outputs: list[ModelRunnerOutput | None], output_rank: int = 0
    ) -> ModelRunnerOutput | None:
        output = outputs[output_rank]
        if not output:
            # output_rank 本步没有输出, 整体聚合结果为空
            return None

        finished_sending = set[str]()
        finished_recving = set[str]()
        worker_meta = None
        for model_runner_output in outputs:
            assert model_runner_output is not None
            ec_output = model_runner_output.ec_connector_output
            if not ec_output:
                continue

            # 并集语义: 任何 rank 上报的完成集合都要保留,
            # 供 scheduler 侧 connector 判断 mm_hash 收发是否结束
            finished_sending |= ec_output.finished_sending or set()
            finished_recving |= ec_output.finished_recving or set()

            if meta := ec_output.ec_connector_worker_meta:
                # 首个非空 meta 作为种子, 后续用 worker 自定义的
                # aggregate() 归并 (返回新对象, 不原地修改)
```

```

        worker_meta = (
            meta if worker_meta is None else worker_meta.aggregate(meta)
        )

    aggregated = ECConnectorOutput(
        finished_sending=finished_sending or None,
        finished_recving=finished_recving or None,
        ec_connector_worker_meta=worker_meta,
    )
    if aggregated.is_empty():
        # 没有任何有效上报时，不要把空对象继续传给 scheduler
        output.ec_connector_output = None
        return output

    # output 可能是共享空输出单例（output_rank 本步无工作），
    # 必须用 copy-on-write 助手挂载，避免污染全局 EMPTY_MODEL_RUNNER_OUTPUT
    return ModelRunnerOutput.with_ec_conn_output(output, aggregated)

```

vllm/v1/outputs.py

输出数据契约扩展：ECConnectorOutput 新增 worker meta 字段与 is_empty(), ModelRunnerOutput 新增 with_ec_conn_output 系列 copy-on-write 助手，是全链路数据传递的基石。

```

@dataclass
class ECConnectorOutput:
    # [mm_hash]
    finished_sending: set[str] | None = None
    finished_recving: set[str] | None = None
    # worker 侧 connector 按步构建、随输出回传 scheduler 的元数据
    ec_connector_worker_meta: ECConnectorWorkerMetadata | None = None

    def is_empty(self):
        return (
            not self.finished_sending
            and not self.finished_recving
            and not self.ec_connector_worker_meta
        )

class ModelRunnerOutput:
    # ... 其余字段省略，仅展示 EC 相关静态方法 ...

    @staticmethod
    def with_ec_conn_output_only(
        ec_connector_output: ECConnectorOutput | None,
    ) -> "ModelRunnerOutput":
        """返回一个只携带 ec_connector_output 的空输出。"""
        return ModelRunnerOutput.with_ec_conn_output(
            EMPTY_MODEL_RUNNER_OUTPUT, ec_connector_output

```

```

)

@staticmethod
def with_ec_conn_output(
    output: "ModelRunnerOutput",
    ec_connector_output: ECCConnectorOutput | None,
) -> "ModelRunnerOutput":
    """返回携带 ec_connector_output 的 output。

    共享的空输出会被复制而不是原地写入，因此调用方必须使用返回值。
    """
    if ec_connector_output is None or ec_connector_output.is_empty():
        return output
    if output is EMPTY_MODEL_RUNNER_OUTPUT:
        # copy-on-write: 绝不把本步特有的 metadata 写进全局单例
        output = copy(EMPTY_MODEL_RUNNER_OUTPUT)
    output.ec_connector_output = ec_connector_output
    return output

```

vllm/v1/worker/gpu/ec_connector.py

worker 侧收发出口：ActiveECCConnector.no_forward 让无 forward 步骤也能收发与上报 meta；修复 ec_both 节点 save_new_caches 判定。

```

class ActiveECCConnector(ECCConnector):
    """包装真实 EC connector，负责在 step 边界收集 worker 侧状态。"""

    def __init__(self, vllm_config: VllmConfig, encoder_cache: dict) -> None:
        self.encoder_cache = encoder_cache
        self.ec_connector = get_ec_transfer()
        assert isinstance(self.ec_connector, ECCConnectorBase)
        # 每个 producer 都要卸载刚算出的 encoder 输出，包括既保存
        # 又加载的 ec_both 节点（此前仅认 is_ec_producer_only，会漏存）
        self.save_new_caches = self.ec_connector.is_producer

    @contextmanager
    def maybe_get_output(
        self, scheduler_output: "SchedulerOutput"
    ) -> Generator[ECCConnectorOutput | None, None, None]:
        # ... 进入时初始化 output，退出（finally）时统一收尾：
        # 取回本步完成收发集合、构建 worker 元数据并清空 connector 状态
        output.finished_sending, output.finished_recving = (
            ec_connector.get_finished(scheduler_output.finished_req_ids)
        )
        output.ec_connector_worker_meta = (
            ec_connector.build_connector_worker_meta()
        )
        ec_connector.clear_connector_metadata()
        # yield ...

```



```
def no_forward(self, scheduler_output: "SchedulerOutput") -> ModelRunnerOutput:
    # 即使本步没有模型 forward, 也让 EC connector 继续收发,
    # 并把 worker 元数据带回 scheduler 侧
    with self.maybe_get_output(scheduler_output) as ec_connector_output:
        pass
    return ModelRunnerOutput.with_ec_conn_output_only(ec_connector_output)
```

NO_OP_EC_CONNECTOR = ECConnector()

评论区精华

评审中 yewentao256 最初两次质疑必要性（要求给出 main 不支持的场景、希望 diff 缩到 200 LOC 内），作者在 PR body 补充动机后得到认可，最终 APPROVED（“LGTM, thanks for the iterations!”）。

核心交锋有四处：

- P1 全局单例污染：yewentao256 指出直接给共享的 EMPTY_MODEL_RUNNER_OUTPUT 赋 ec_connector_output 会永久污染单例，后续步骤可能重放过期状态；depthfirst-app[bot] 也以 LOW 严重度报告同一问题。作者改为 with_ec_conn_output 的 copy-on-write 语义并补测试。
- P2 legacy Ray 执行器：yewentao256 指出 legacy Ray 只取 refs[0]，EC+KV 组合时只跑 KV 聚合，EC 元数据会被静默丢弃。作者选择显式拒绝该组合，Ray 完整支持留作 future work。
- 聚合语义是否过早报完成：gty111 对比 KVOutputAggregator 的 expected_finished_count 屏障，质疑 union 语义在多 rank 下会过早报告完成。作者回应：PP > 1 时仅首个 rank 持有 encoder（其余为 no-op connector 不上报）；TP 下 encoder_cache[mm_hash] 在全部 rank 上是一致的完整张量（weights 模式 all_reduce、data 模式 all_gather）；避免重复上报由 connector 实现负责。
- 消费者加载屏障：gty111 追问调度器是否应等待所有参与 rank 完成加载再恢复请求。作者认为屏障应在 connector 实现中做——没有 encoder cache 的 rank 用 NO_OP_EC_CONNECTOR 永不上报，若采用 KV 式全 rank 屏障会永久挂起。
 - 另外 yewentao256 对 test_with_ec_conn_output_sets_field_in_place 评价“Not that meaningful test”，作者随后移除了该测试。
- EC 聚合并集语义是否会过早报告完成 (correctness): omerpaz95 回应：PP > 1 时仅首个 rank 持有 encoder（其余为 no-op connector 不上报）；TP 下 encoder_cache[mm_hash] 在所有 rank 上是完整一致的张量（weights 模式 all_reduce、data 模式 all_gather）；避免重复上报由 connector 实现负责。
- 共享 EMPTY_MODEL_RUNNER_OUTPUT 全局单例污染 (correctness): 作者新增 with_ec_conn_output copy-on-write 助手，并新增 test_aggregate_does_not_write_through_the_shared_empty_output 与 test_with_ec_conn_output_copies_shared_empty_output 防护。
- 消费者异步加载是否需要完成屏障 (correctness): omerpaz95 认为屏障应由 connector 实现；无 encoder cache 的 rank 使用 NO_OP_EC_CONNECTOR 永不上报，采用

KVOutputAggregator 式全 rank 屏障会永久挂起。

- legacy Ray executor 静默丢弃 EC 元数据 (design): 作者选择显式拒绝该组合 (world_size > 1 时抛 NotImplementedError, 提示改用 RayExecutorV2 或多进程 executor), Ray 完整支持留作后续工作。
- 38390 落地后本 PR 的必要性 (question): 作者在 PR body 补充动机 (worker -> scheduler 元数据通道缺失, 阻塞 CUDA events CPU offload 与 NIXL P2P), 并通过多轮提交压缩 diff、回退 MRv1 改动, 最终获得 APPROVED。
- 低价值测试的取舍 (testing): 作者在后续提交中移除了该测试, 保留关注单例写穿的测试。

风险与影响

• 风险:

1. 聚合器并集语义缺少全 rank 屏障: ECTOutputAggregator.aggregate 对 finished_sending / finished_recving 取并集, 不像 KVOutputAggregator 那样等待 expected_finished_count 个 worker; 若某个 connector 实现未自行保证“所有参与 rank 完成”, 调度器可能过早恢复对 embedding 的读取 / 写入, 提前解除请求挂起。
2. 共享单例写穿回归风险: EMPTY_MODEL_RUNNER_OUTPUT 是模块级共享对象, with_ec_conn_output 的 copy-on-write 已防御, 但 model_runner.py 这类核心路径若未来有代码直接赋值 output.ec_connector_output = ... 或忽略返回值, 仍会重新引入污染。
3. legacy Ray 组合行为断裂: ray_executor.py 对 world_size > 1 + EC connector 的组合直接抛 NotImplementedError, 属于启动期行为变更, 存量用户需切换 RayExecutorV2 或多进程 executor。
4. MRv2 核心路径回归: execute_model / sample_tokens / pool 与 ExecuteModelState 数据契约均有改动, 零 forward 步骤与 encoder-only 步骤是回归高发区; 提交历史中的“Fix eplb CI”说明本 PR 曾引发 EPLB 相关测试失败。
5. 跨进序列序列化契约未覆盖: ECConnectorWorkerMetadata 实现需随 ModelRunnerOutput 跨进程传递, 但 PR 未提供该抽象类序列化契约的测试; 自定义 meta 若含不可序列化对象会在运行时失败。- 影响: 对默认配置 (无 EC connector) 行为完全不变: build_connector_worker_meta() 默认返回 None, 聚合器仅在配置 EC connector 时创建。对启用 EC connector 的分布式部署, 这是首个 worker -> scheduler 的 EC 状态通道, 使 CPU offload (CUDA events 方案, 替代 StepTracker) 与 NIXL P2P 两类 EC connector 可以落地。对团队而言, EC 补齐了与 KV 对齐的三件套 (build_connector_worker_meta / 输出字段 / 聚合器), 后续 connector 只需实现抽象类即可接入调度器。影响面覆盖 MRv1/MRv2 两个 runner、multiproc 与 Ray executor、v1 输出数据契约, 属于分布式 EC 缓存架构的关键地基。- 风险标记: MRv2 核心路径变更, 聚合语义缺少全 rank 屏障, 共享单例写穿风险 (已防御), legacy Ray 组合被显式拒绝, 元数据序列化契约未覆盖

关联脉络

- PR #38390 : V2 model runner 的 EC Connector 基础实现; review 中 yewentao256 要求等其落地后再评审本 PR, PR body 也以“Why this is still needed after #38390”开篇,

本 PR 在其基础上补齐元数据通道。

- PR #49994 : 基于 CUDA events 检查 GPU->CPU 拷贝完成、以替代 StepTracker 的 CPU offload EC connector; 作者在评论中明确这是本 PR 的消费者场景与前置依赖。
- PR #47941 : NIXL 基于 EC 的共享方案; 生产者需借助本 PR 的机制正确标记 CPU 上的 EC 条目就绪, 消费者才能安全 NIXL READ。