

PR #29799 完整报告

sgl-project/sglang

support rust sglang server

合并时间: 2026-08-01 02:56

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29799>

执行摘要

- 一句话: 引入嵌入式 Rust 服务器替换 Python 前端, 提升吞吐与扩展性
- 推荐动作: 值得精读。该 PR 是 SGLang 前端架构演进的关键一步, 展示了如何通过嵌入式异构语言线程池降低 GIL 瓶颈, 以及列式 wire protocol 的设计权衡。建议重点关注: rust_server.py 的列式 egress 实现、flatten.py 的缓冲区契约、tokenizer.rs 的 stop 语义修正。同时注意 review 中 P0/P1 问题的修复方式 (如背压策略) 与后续 PR 的衔接。

功能与动机

PR body 引用 RFC #23206, 动机包括: 迁移 pre-scheduler python server 到 rust; 为 DP attn 相关改动做准备; 支持 IO 密集的 thread-per-core 运行时; 支持 standalone (或外部) api-server 部署。作者在 issue 评论中明确目标是将 tokenize/detokenize 等 CPU 密集工作从 Python scheduler 的 critical path 中移出, 以降低 GIL 占用。

实现拆解

该 PR 以环境变量 `SGLANG_RUST_SERVER` 为开关, 将 Rust 服务器嵌入 scheduler 进程。

1. Rust 服务器核心 (新增 / 扩展 `rust/sglang-server/`): 包含 api_server (axum/tokio, 支持 /generate unary + SSE、/server_info、/health 等端点)、tokenizer 线程池、detokenizer 分片、TokenizerManager 的 ingress/egress 双 FSM 线程。所有 CPU 密集阶段均使用 Runnable 特质 + CPU 亲和性 pinning, 避免与 scheduler 争抢 GIL。
2. Python 侧嵌入生命周期 (新增 `python/sglang/srt/managers/rust_server.py`): RustServer.launch() 启动嵌入式 Server, _partition_cores() 划分 CPU 核心, drain() 从 ingress ring 取请求, push_generation() 将整批输出推入 egress ring, push_control_output() 处理控制请求响应。关键设计是 列式传输: ingress 中 input_ids 作为连续 int64 buffer 直接传, egress 中 logprobs/hidden states 通过 flatten.py 的 FlatPairColumns/RaggedPairColumns/NestedRowColumns 类展平为 f32/i32 buffer + 长度向量, 避免 per-request msgpack。
3. scheduler 集成 (`python/sglang/srt/managers/scheduler.py`、`idle_sleeper.py`、`output_streamer.py`): maybe_init_rust_server() 在 rank 0 启动 RustServer 并将 recv_from_tokenizer 指向它; RustServerIdleSleeper 替代 zmq poller 直接在 ingress ring 上休眠, 避免空转 100% CPU; _GenerationStreamAccumulator 在 rust mode 下跳过 detok 相关累积工作。

4. 协议与错误处理配套 (`managers/utils.py`) : 新增 `msgpack_decode_explained()` (将 `msgspec` 校验错误定位到具体字段并以 400 返回给客户端) 和 `compute_num_reserved_tokens()` (供 Rust 侧做 total-token 预算检查, 兼顾 eagle 草稿 token) 。
5. 测试配套: `test/registered/core/test_srt_endpoint.py` 新增 `TestRustServerEndpoint` (继承整个 Python endpoint 套件并跳过未实现项)、`test_cargo_workspace.py` 跑 Rust 单元测试、`simple_eval_common.py` 新增 `GenerateSampler` 支持 gsm8k 评测、`test_modelopt_fp8.py` 增加 Rust 模式覆盖。

关键文件:

- `python/sglang/srt/managers/rust_server.py` (模块 嵌入式服务; 类别 source; 类型 core-logic; 符号 `RustServer`, `launch`, `wait_ingress`, `drain`) : Python 侧嵌入式 `RustServer` 生命周期封装, 定义了列式 ingress/egress 传输的核心逻辑 (`drain/push_generation`) , 是 Python 与 Rust 边界的枢纽。
- `python/sglang/srt/utils/flatten.py` (模块 传输协议; 类别 source; 类型 core-logic; 符号 `flatten_ragged`, `flatten_hidden`, `_flatten_floats`, `FlatPairColumns`) : 定义列式 egress 的展平契约 (`FlatPairColumns/RaggedPairColumns/NestedRowColumns`) , 包含浓郁的防御性断言, 保证 Python 与 Rust 的 buffer 布局一致。
- `python/sglang/srt/managers/scheduler_components/idle_sleeper.py` (模块 调度器; 类别 source; 类型 core-logic; 符号 `RustServerIdleSleeper`, `maybe_sleep`) : 新增 `RustServerIdleSleeper`, 将 scheduler 空闲等待从 `zmq poller` 改为在 ingress ring 上阻塞, 消除 100% CPU 空转。
- `python/sglang/srt/managers/utils.py` (模块 协议解析; 类别 source; 类型 dependency-wiring; 符号 `MsgpackDecodeError`, `msgpack_decode_explained`, `compute_num_reserved_tokens`) : 新增 `msgpack_decode_explained` 与 `compute_num_reserved_tokens`, 前者把 `msgspec` 校验错误转为可回传客户端的字段级错误, 后者同步 Rust 侧 token 预算。
- `python/sglang/srt/managers/scheduler.py` (模块 调度器; 类别 source; 类型 core-logic ; 符号 `maybe_init_rust_server`) : 调度器核心集成点: `maybe_init_rust_server` 决定是否替换 request receiver、idle sleeper 与 output streamer 的 egress 目标。
- `rust/sglang-server/src/tokenizer.rs` (模块 Tokenizer; 类别 source; 类型 core-logic; 符号 `TokenizerWorker`, `tokenizing_replaces_the_byte_window_with_a_token_count`) : 修复 `stop_str_max_len` 从字节数改为 token 数, 配套单元测试验证, 是语义正确性的关键修正。
- `python/sglang/srt/entrypoints/engine.py` (模块 引擎入口; 类别 source; 类型 dependency-wiring; 符号 `SchedulerInitResult.block_until_scheduler_exits`) : Engine 离线 API 显式拒绝 Rust 模式, 并重命名 `wait_for_completion` 语义, 避免 offline 与 server 模式混用。
- `python/sglang/srt/managers/scheduler_components/output_streamer.py` (模块 输出流 ; 类别 source; 类型 core-logic; 符号 `SchedulerOutputStreamer._stream_output_generation`, `_GenerationStreamAccumulator`) : rust mode 下跳过增量 `detok` 相关累积 (`decode_ids/read_offset` 等) , 避免 dead work, 同时将 egress 改为 `push_generation`。

- test/registered/core/test_srt_endpoint.py (模块 端点测试; 类别 test; 类型 test-coverage; 符号 TestRustServerEndpoint, test_greedy_token_equals_top1) : 新增 TestRustServerEndpoint 复用整套 endpoint 测试, 并补充 logprob 列对齐的贪婪测试, 验证列式传输正确性。
- rust/sglang-server/src/tokenizer_manager/ingress.rs (模块 Ingress; 类别 source; 类型 core-logic; 符号 Ingress) : ingress FSM 核心, review 中披露了注册泄露与 RID 前缀匹配问题, 最终修复。学习其单消费者状态机设计。
- python/sglang/srt/managers/scheduler_components/request_receiver.py (模块 请求接收; 类别 source; 类型 dependency-wiring; 符号 SchedulerRequestReceiver) : 请求接收方根据 rust_server_mode 选择从 Rust ring 还是 zmq 取请求。
- test/registered/rust/test_cargo_workspace.py (模块 测试; 类别 test; 类型 test-coverage; 符号 TestCargoWorkspace) : 新增 Rust workspace 的 CI 单测入口, 确保 cargo test 在注册测试中运行。
- python/sglang/test/simple_eval_common.py (模块 评测工具; 类别 test; 类型 test-coverage; 符号 GenerateSampler) : 为 gsm8k 准确率评测提供原生 /generate 采样器 GenerateSampler。

关键符号: RustServer.launch, RustServer.drain, RustServer.push_generation, RustServer.push_control_output, flatten_ragged, flatten_hidden, FlatPairColumns.accept, RaggedPairColumns.accept, NestedRowColumns.accept, msgpack_decode_explained, compute_num_reserved_tokens, RustServerIdleSleeper.maybe_sleep, TokenizerWorker.run, Scheduler.maybe_init_rust_server, SchedulerOutputStreamer._stream_output_generation

关键源码片段

python/sglang/srt/utils/flatten.py

定义列式 egress 的展平契约 (FlatPairColumns/RaggedPairColumns/NestedRowColumns), 包含浓郁的防御性断言, 保证 Python 与 Rust 的 buffer 布局一致。

python/sglang/srt/utils/flatten.py (核心摘录)

def flatten_ragged(per_pos_val, per_pos_idx):

"""展平每位置 `list[Optional[list]]` val/idx 对为 flat buffer + 共享 lens。

idx 必须与 val 完全镜像 (断言保证) : 线上只用一条 lens 向量配对两个 buffer, 一旦发散就会把后续列的偏移整体推移, 导致下游 token id 移位或丢失。

"""

flat_val = []

flat_idx = []

lens = []

if not per_pos_val:

assert not per_pos_idx, (

f"ragged idx 列有 {len(per_pos_idx)} 个位置但 val 列为空")

return flat_val, flat_idx, lens

assert per_pos_idx is not None and len(per_pos_idx) == len(per_pos_val)

```

for p, pv in enumerate(per_pos_val):
    pi = per_pos_idx[p]
    if pv: # truthy 位置只含真实 logprobs; None/ 空位置走 else 分支
        assert pi is not None and len(pi) == len(pv)
        flat_val.extend(pv)
        flat_idx.extend(pi)
        lens.append(len(pv))
    else:
        assert not pi, f"position {p}: idx 有 {len(pi)} 项但 val 为空"
        lens.append(0)
return flat_val, flat_idx, lens

```

```

class FlatPairColumns:

```

```

    """一列平坦 val/idx 对 (如 per-token logprob 值与 token id) 。

```

```

    header 里放每请求元素数, data 里放拼接的 f32 + i32 buffer。

```

```

    `first_none_to_nan` 把输入 logprob 首 token 的 None 哨兵编码为 NaN。
    """

```

```

    def accept(self, j):

```

```

        vv = (self.vals[j] if self.vals else None) or []
        ii = (self.idxs[j] if self.idxs else None) or []
        # 防御性断言: lens 只记录 len(vv), 但两个 buffer 都会扩展,
        # 若 idx 更长会静默把后续列的偏移推偏, 解码端无法察觉。
        assert len(ii) == len(vv), (
            f"{self.name}: request {j} 有 {len(ii)} 个 idx 但 {len(vv)} 个 val")
        if self.first_none_to_nan and vv and vv[0] is None:
            self.v.append(float("nan"))
            self.v.extend(vv[1:])
        else:
            self.v.extend(vv)
        self.i.extend(ii)
        self.lens.append(len(vv))

```

python/sclang/srt/managers/utils.py

新增 msgpack_decode_explained 与 compute_num_reserved_tokens, 前者把 msgspec 校验错误转为可回传客户端的字段级错误, 后者同步 Rust 侧 token 预算。

```

# python/sclang/srt/managers/utils.py (核心摘录)

```

```

class MsgpackDecodeError(ValueError):

```

```

    """msgpack 帧被 typed decoder 拒绝, 且带解释: rid + 人类可读 reason。"""

```

```

    def __init__(self, rid: Optional[str], reason: str):
        super().__init__(reason)
        self.rid = rid
        self.reason = reason

```

```

def msgpack_decode_explained(data: bytes) -> Any:

```

"""`io\_struct.msgpack\_decode` 的增强版：被拒帧抛 `MsgpackDecodeError`，携带 rid（通过无类型重新解码 tagged array 恢复）和带字段名的 reason。

供需要把失败回传给客户端（如 rust ingress）而不是直接崩溃的调用方使用。

"""

```
try:
    return io_struct.msgpack_decode(data)
except Exception as e:
    msg = str(e)
    try:
        arr = msgspec.msgpack.decode(data)
    except Exception:
        arr = None
    if not (isinstance(arr, (list, tuple)) and arr):
        raise MsgpackDecodeError(None, msg) from e
    # tagged array_like 布局为 [tag, *fields]; rid 是所有 BaseReq 的首字段
    rid = str(arr[1]) if len(arr) > 1 and arr[1] is not None else None
    tag_to_fields = {
        cls.__struct_config__.tag: cls.__struct_fields__
        for cls in io_struct._all_types
        if isinstance(cls, type) and issubclass(cls, msgspec.Struct)
    }
    fields = tag_to_fields.get(arr[0])
    if fields is not None:
        # 把 msgspec ValidationError 路径里的 `${<n>}` 下标映射为字段名
        m = re.search(r"\${(\d+)\}", msg)
        if m is not None:
            idx = int(m.group(1))
            if 1 <= idx <= len(fields):
                msg = f"{msg[:m.start()]}${fields[idx - 1]}{msg[m.end():]}"
        raise MsgpackDecodeError(rid, msg) from e
```

rust/sclang-server/src/tokenizer.rs

修复 stop_str_max_len 从字节数改为 token 数，配套单元测试验证，是语义正确性的关键修正。

// rust/sclang-server/src/tokenizer.rs (核心摘录)

```
impl Runnable for TokenizerWorker {
    fn run(self) {
        while let Ok(mut req) = self.rx.recv() {
            let event = {
                let RequestKind::Generate(g) = &mut req.kind else {
                    tracing::error!("tokenizer pool received a non-generate request");
                    continue;
                };
                // 把 scheduler 的 stop 匹配窗口按 TOKEN 数设置，与 Python
                // `normalize(tokenizer)` 的行为对齐。
                // 某个 stop 编码失败时回退到其字节长度：仍为过估计，绝不低估，
                // 因此 scheduler 不可能漏掉该 stop。
            }
        }
    }
}
```

```

        let stop_tokens = g
            .sampling_params
            .stop_strs
            .iter()
            .map(|sl| self.tokenizer.encode(s).map_or(s.len(), |lidl| ids.len()))
            .max();
        if let Some(n) = stop_tokens {
            g.sampling_params.stop_str_max_len = n;
        }
        match self.tokenizer.encode(g.text.as_deref().unwrap_or("")) {
            Ok(ids) => {
                g.input_ids = Some(ids);
                Event::TokenizeDone
            }
            Err(err) => Event::Error(err),
        }
    };
    let _ = req.state.apply(event);
    if self.tm.send(TmEvent::Tokenized(req)).is_err() {
        tracing::error!("tm inbox closed; dropping request");
        break;
    }
}
}
}

#[cfg(test)]
mod tests {
    // WordTokenizer: 每个空格分隔的单词一个 token，使 stop 的 token 数与字节数可区分。
    #[test]
    fn tokenizing_replaces_the_byte_window_with_a_token_count() {
        // 8 字节 vs 3 token——单位可区分。
        // `Normalizing` 阶段留下的是字节数（安全过估计），但会导致 scheduler
        // 每步解码更长的尾部（典型 stop 集合 14 token vs 6 token）。
        // 此阶段拥有 tokenizer，因此在这里解析出精确计数。
        assert_eq!(g.sampling_params.stop_str_max_len, 3);
    }
}
}

```

评论区精华

P0 安全边界问题 (ishandhanani)：Rust router 未安装 API-key 中间件，`/server_info` 泄露全局 `ServerArgs` 中的 `credential` 字段。作者回应：authz 留作 follow-up，当前通过 `INTERNAL_STATE_ALLOWLIST` 过滤敏感字段。

P0 默认路径回归 (ishandhanani)：`maybe_init_rust_server()` 在 `SGLANG_RUST_SERVER=0` 时把 `recv_from_tokenizer` 置为 `None`，导致默认 zmq 路径启动失败。作者确认已修复。

P1 背压丢帧与泄漏 (ishandhanani、alexnails) : egress ring 满时 `push_batch` 非阻塞返回 `False`, 整批输出被丢弃, 且 `detok shard` 的 `per-rid` 状态永不清理, 造成 SSE 悬挂和内存泄漏; 作者表示已修复, 转而采用 `push` 侧阻塞或重试策略。

P1 stop 语义与字符计数上界 (ishandhanani) : `stop_str_max_len` 用字符数作上界不安全 (如 ② 1 字符编码 3 token), 会漏匹配 stop; 作者最终在 `tokenizer.rs` 中改用 `tokenizer` 编码后统计 token 数 (带测试验证)。

性能讨论 (alexnails) : `frame_egress_batch` 在持有 GIL 时做整批 `memcpy`; `OutputAccumulator.snapshot()` 每帧全量 clone 导致 $O(T^2)$ 拷贝; egress 单 dispatcher 用阻塞 `send` 造成 HOL。作者对 HOL 回应: 当前阶段避免过度设计, 若未来成为瓶颈再对齐 `detok worker` 数。

设计意见 (sagearc) : 协议契约分散 (`dynamo_protocols` vs 手写 JSON vs `msgpack`), 建议集中化。作者答复: `openai.rs` 只是基于 `dynamo_protocol` 的 HTTP 包装。

代码结构 (merrymercy) : 要求用 `foo.rs` 替代 `foo/mod.rs` 现代模块布局, 作者已整改。

- API key 鉴权缺失与敏感信息泄露 (security): 作者: `authz` 作为 follow-up; 当前用 `INTERNAL_STATE_ALLOWLIST` 过滤敏感字段, 并加 TODO。
- 默认路径 `SGLANG_RUST_SERVER=0` 启动崩溃 (correctness): 作者确认修复: 非 Rust 模式不动 canonical receiver。
- egress ring 满时丢帧与 `detok` 状态泄漏 (correctness): 作者声称已修复 (改为阻塞 / 重试策略)。
- `stop_str_max_len` 字符数上界不安全 (correctness): 作者最终在 `tokenizer.rs` 中用 `tokenizer` 实际编码得到 token 数, 并加差分测试。
- GIL 下整批 egress `memcpy` (performance): 作者表示已修复, 但最终实现还可能保留部分复制; 该评论无后续回复。
- egress 单 dispatcher HOL 阻塞 (performance): 作者承认潜在问题, 但认为当前阶段避免过度设计, 已在代码中注释 TODO, 未来按需对齐 dispatcher 到 `detok worker`。
- 协议契约统一 (design): 作者回答: `dynamo_protocols` 支持 reasoning/function call, `openai.rs` 仅是 HTTP 包装。
- 停止 `regex` 的有限边界 (performance): 作者在后续提交中改用了有限边界计算 (或回退到 Python 侧规范化)。
- 模块布局 `mod.rs` vs `foo.rs` (style): 作者已整改。

风险与影响

- 风险:
 1. 默认路径回归风险: `SGLANG_RUST_SERVER=0` 时 `recv_from_tokenizer` 置 `None` 的问题虽已修复, 但 `maybe_init_rust_server()` 改动位于 `scheduler` 初始化主路径, 任何环境变量解析或 `rank` 判断失误都可能影响所有常规部署。
 2. 协议漂移风险: Python 与 Rust 之间通过手写列式布局 (`header_cols/data_cols` 顺序) 耦合, `flatten.py` 中大量断言 (如 `RaggedPairColumns` 的 `val/idx` 长度一致) 仅防 Python 侧错误, 若 Rust 侧 `for_each_chunk` 读取顺序不同会静默错位。

msgpack_decode_explained 依赖 io_struct._all_types 的内部结构。

3. CPU 亲和性与线程数: `_partition_cores()` 与 `plan_cores()` 的 fallback 逻辑曾导致全部线程继承 scheduler 的窄 mask (P1 已修)。若机器核心数不足, 所有 Rust 线程可能被 pin 到单核。
 4. 功能差异: PR body 明确列出众多未支持项 (`custom_logit_processor`、`cache_tokens`、`parallel sampling`、`preferred_sampling_params`、`embedding`、多节点 `dp_size>1` 等), 用户若开启 Rust 模式可能静默丢失这些能力 (部分已通过拒绝启动或 400 响应防护, 但仍有 `extra` 字段被接受却不生效的情况)。
 5. 构建与依赖风险: 新增 rust workspace 与 Cargo.lock (锁 `dynamo-tokenizers 1.5.3`), 若未正确构建扩展, `is_rust_server_built()` 会 skip 测试; 但生产环境若遗漏编译会启动失败。
 - 影响: 性能影响: 将 `tokenize/detokenize`、OpenAI API、采样归一化从 Python 移出, 显著降低 scheduler 线程的 GIL 占用; 空闲时从 100% CPU 降为 ring 阻塞等待。但 PR 中披露 scheduler 本身 (`accept/process_batch_result`) 在大 batch 下仍是吞吐天花板, Rust 化只解决前端部分。架构影响: 为后续 DP attn、`thread-per-core`、`standalone api-server` 打下基础; 新增 rust/sglang-server crate 与 Python 边界 API (`Server.recv_requests/push_batch/push_result`) 成为新契约。团队影响: 需要维护 Python/Rust 双实现直至 Python 路径退役, 增加跨语言协议同步成本; 测试矩阵新增 Rust workspace 单测与 Rust 模式 e2e 测试。用户影响: 默认不受影响 (需显式设置 `SGLANG_RUST_SERVER=1`), 但该模式目前仅支持 `/generate` 等少量端点, OpenAI 兼容 API 不可用。
- 风险标记: 核心路径变更, 跨语言协议漂移风险, 安全边界缺失, 功能不完整, 性能瓶颈未完全消除, 构建依赖新增

关联脉络

- PR #32014 splittable PR (DONE 系列): PR body 列出的 Done 系列拆分之一, 作为本 PR 的组成部分先期合入。
- PR #32872 tokenizer.rs, detokenizer.rs, tokenizer_manager/egress.rs 976: 本 PR 的 Rust 代码被拆分为独立 PR 合入, 这些 PR 构成 Rust 服务器的主体。
- PR #32873 ingress.rs (implementation), api_server/common.rs, log.rs, openai.rs 888: 同上, ingress 与 api_server 实现拆分。
- PR #32877 wiring: lib.rs, runtime.rs, message.rs, tokenizer_manager.rs, environ.rs, error.rs 735: Rust 服务器核心 wiring 拆分 PR。
- PR #33850 [diffusion] retire released warmup and decoder flags: 同为 sglang runtime 前端的演进, 但属于 diffusion 领域, 与本 PR 无直接关系 (不列为关联)。