

PR #36252 完整报告

sgl-project/sglang

config: stop handing the record to code that does not read it

合并时间: 2026-08-26 20:02

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/36252>

执行摘要

- 一句话: 删除未读 `server_args` 参数, 权重加载工厂改名改读 `published config`
- 推荐动作: 值得精读, 尤其是关注配置架构演进与大规模机械重构方法的读者。要点: ①判断「未读参数即负债」的标准 (函数体内是否已全部改读 `published config`) ; ②工厂命名按实际读取来源而非字面参数 (`from_published_config`) ; ③ratchet 测试对「重新引入 `record`」的约束作用。若只关心功能行为, 本 PR 无用户可见变化, 可跳过。

功能与动机

配置重构系列的关键收尾。PR body 指出: 一旦函数内部已经通过 `runtime_context` 读取 `published config`, 继续为它传递的 `server_args` 参数就是死重 (dead weight) ——'Leaving it in place is not free: it keeps a record reference alive in the 48 production files this PR drops it from, it invites the next contributor to read a field off it, and it makes the call graph look like configuration flows by argument when it no longer does.' 此外, `StartupWeightLoadOptions.from_server_args` 接收 `record` 却从 `bags` 读取十六个配置叶子, 命名与行为相反, 需要在名字上修正为 `from_published_config`。

实现拆解

本 PR 的变更按四条线展开:

1. 删除未读参数 (主战场): 在 48 个生产文件中删除函数签名里从未被读取的 `server_args` 参数, 并同步更新所有调用点。典型代表包括: `python/sglang/srt/utils/common.py` 的 `require_mlp_tp_gather`、`require_attn_tp_gather`、`require_gathered_buffer`、`require_mlp_sync`、`get_cuda_graph_batch_size_alignment`、`get_cuda_graph_max_batch_size`、`get_eager_max_batch_size` ; `python/sglang/srt/managers/disagg_service.py` 的 `start_disagg_service` 与 `maybe_create_ascend_config_store` ; `python/sglang/srt/eplb/expert_location.py` 的 `init_trivial` / `init_by_mapping` / `init_by_eplb` / `_init_common` ; `python/sglang/srt/managers/scheduler_components/batch_result_processor.py` 的 `_get_prefill_hidden_capture_mode`。这些函数体内早已通过 `get_exec()` / `get_parallel()` 等 `runtime_context` 入口读取 `published config`, 参数纯属死重; 删除后调用图不再错误暗示配置按参数流动。同步清理 `TYPE_CHECKING` 下的 `from sglang.srt.server_args import ServerArgs` (`common.py` 中该块变为 `pass`) 。

2. 重命名权重加载工厂，让名字反映事实：python/sglang/srt/model_executor/model_runner_components/startup_weight_load.py 中 StartupWeightLoadOptions.from_server_args 改为 from_published_config, StartupWeightLoadManager.create_from_server_args 改为 create_from_published_config, 签名删除 server_args, 仅保留 is_draft_worker (该参数描述 runner 角色而非进程配置，作为唯一例外继续传参)。调用点 load_model_utils.py 的 load_model_with_memory_saver 与 model_runner.py 同步更新。
3. 三个 Manager 改读 bags: EPLBManager 原先拿 ServerArgs 读取九个 exec.moe / parallel 叶子，改为直接读 bags 后不再接收记录；LoRAManager 与多模态 BaseMultimodalProcessor 同样去掉 handed record, 改为读 published config。对应地从 supplied-instance exposure ratchet 测试的允许暴露实例列表中删除这些条目，防止未来重新引入 record 依赖。
4. 测试配套：更新 test_supplied_instance_exposure_ratchet.py、test_startup_weight_load.py、test_kimi_k25.py 等至少 10 个测试文件，与源码模块改动联动；PR 提供三条验证命令覆盖 ratchet、启动权重加载与 Kimi K2.5 模型测试。纯重构声明，无行为变化；CI 中 PR Test (Base) 与 AMD ROCm 运行显示失败，PR Test (Extra) 通过，失败原因未在材料中说明，需在合并前确认。

关键文件：

- python/sglang/srt/model_executor/model_runner_components/startup_weight_load.py (模块 权重加载；类别 source；类型 data-contract；符号 from_server_args, from_published_config, create_from_server_args, create_from_published_config)：数据契约核心变更：StartupWeightLoadOptions / StartupWeightLoadManager 的工厂从 from_server_args 重命名为 from_published_config, 删除 server_args 参数，全部叶子改从 runtime_context 的 config bags 读取，是本 PR 命名原则的代表作。
- python/sglang/srt/utils/common.py (模块 通用工具；类别 source；类型 core-logic；符号 require_mlp_tp_gather, require_attn_tp_gather, require_gathered_buffer, require_mlp_sync)：通用工具中 MoE/DP 判定链与 cuda graph 桶大小计算函数全部去掉 server_args 参数；函数体早已只读 get_exec() / get_parallel(), 参数是死重，且 TYPE_CHECKING 导入清理为 pass。
- python/sglang/srt/eplb/expert_location.py (模块 专家路由；类别 source；类型 core-logic；符号 init_trivial, init_by_mapping, init_by_eplb, _init_common)：EPLB 专家位置元数据的全部构造入口 (init_trivial / init_by_mapping / init_by_eplb / _init_common) 删除 server_args, 配置读取统一走 runtime_context; broadcast_global_expert_location_metadata 也不再通过 get_server_args() 取记录。
- python/sglang/srt/managers/disagg_service.py (模块 分离服务；类别 source；类型 core-logic；符号 start_disagg_service, maybe_create_ascend_config_store)：start_disagg_service 与 maybe_create_ascend_config_store 删除 server_args, Ascend config store 的创建只依赖 runtime_context 与 ASCEND_MF_STORE_URL 环境变量。
- python/sglang/srt/model_executor/model_runner_components/load_model_utils.py (模块 模型加载；类别 source；类型 data-contract；符号 maybe_downgrade_dtype_for_legacy_gpu, report_online_quantization,

maybe_enable_ipc_weight_cache, load_model_with_memory_saver) :
maybe_downgrade_dtype_for_legacy_gpu / report_online_quantization
/ maybe_enable_ipc_weight_cache 等删除 server_args;
load_model_with_memory_saver 改用 create_from_published_config 构建启动权重加载。

- python/sglang/srt/managers/scheduler_components/batch_result_processor.py (模块调度器; 类别 source; 类型 core-logic; 符号 _get_prefill_hidden_capture_mode) : SchedulerBatchResultProcessor 删除 server_args 字段与 _get_prefill_hidden_capture_mode 的参数, 批处理结果路径不再持有记录。
- python/sglang/srt/eplb/eplb_manager.py (模块 专家路由; 类别 source; 类型 dependency-wiring) : EPLBManager 原先拿 ServerArgs 读九个 exec.moe / parallel 叶子, 改为直接读 bags, 不再接收记录。

关键符号: from_published_config, create_from_published_config,
require_mlp_tp_gather, require_attn_tp_gather, require_gathered_buffer,
require_mlp_sync, get_cuda_graph_batch_size_alignment,
get_cuda_graph_max_batch_size, get_eager_max_batch_size, init_trivial,
init_by_mapping, init_by_eplb, _init_common, start_disagg_service,
maybe_create_ascend_config_store, maybe_downgrade_dtype_for_legacy_gpu,
report_online_quantization, _get_prefill_hidden_capture_mode, init_soft_watchdog,
maybe_init_shared_mooncake_transfer_engine

关键源码片段

python/sglang/srt/model_executor/model_runner_components/startup_weight_load.py

数据契约核心变更: StartupWeightLoadOptions / StartupWeightLoadManager 的工厂从 from_server_args 重命名为 from_published_config, 删除 server_args 参数, 全部叶子改从 runtime_context 的 config bags 读取, 是本 PR 命名原则的代表作。

```
# StartupWeightLoadOptions: 启动权重加载选项, 全部来自 published config 的叶子。  
# 原先工厂接收 ServerArgs, 但函数体内早已通过 runtime_context 的  
# get_exec() / get_parallel() / get_lora() / get_model() 等读取配置,  
# 参数纯属死重, 本 PR 将其删除并重命名工厂。
```

```
@dataclasses.dataclass(frozen=True, slots=True, kw_only=True)  
class StartupWeightLoadOptions:  
    device: str  
    is_cuda_platform: bool  
    cuda_graph_enabled: bool  
    prefill_cuda_graph_backend: Backend  
    is_draft_worker: bool  
    speculative_algorithm: Optional[str]  
    tp_size: int  
    attn_cp_size: int  
    dcp_size: int  
    pp_size: int
```

```
dp_size: int
ep_size: int
cpu_offload_gb: int
offload_group_size: int
enable_memory_saver: bool
enable_weights_cpu_backup: bool
enable_lora: bool
has_lora_paths: bool
weight_loader_disable_mmap: bool
weight_loader_drop_cache_after_load: bool
has_custom_weight_loader: bool
enable_torch_compile: bool
prefetch_num_threads: int
```

```
@classmethod
```

```
def from_published_config(
```

```
    cls,
```

```
    *,
```

```
    is_draft_worker: bool,
```

```
) -> StartupWeightLoadOptions:
```

```
    # 一切所需都是 published leaf, 只有 is_draft_worker 是例外:
```

```
    # 它描述当前 runner 的角色, 而非进程的配置, 故继续以参数传递。
```

```
    cuda_graph_config = get_exec().graph.cuda_graph_config
```

```
    cuda_graph_enabled = any(
```

```
        getattr(cuda_graph_config, phase).backend != Backend.DISABLED
```

```
        for phase in Phase.ALL
```

```
    )
```

```
    return cls(
```

```
        device=get_device().device,
```

```
        is_cuda_platform=current_platform.is_cuda(),
```

```
        cuda_graph_enabled=cuda_graph_enabled,
```

```
        prefill_cuda_graph_backend=cuda_graph_config.prefill.backend,
```

```
        is_draft_worker=is_draft_worker,
```

```
        speculative_algorithm=get_spec().speculative_algorithm,
```

```
        tp_size=get_parallel().config.tp_size,
```

```
        attn_cp_size=get_parallel().config.attn_cp_size,
```

```
        dcp_size=get_parallel().config.dcp_size,
```

```
        pp_size=get_parallel().config.pp_size,
```

```
        dp_size=get_parallel().config.dp_size,
```

```
        ep_size=get_parallel().config.ep_size,
```

```
        cpu_offload_gb=get_exec().offload.cpu_offload_gb,
```

```
        offload_group_size=get_exec().offload.offload_group_size,
```

```
        enable_memory_saver=get_exec().features.enable_memory_saver,
```

```
        enable_weights_cpu_backup=get_exec().features.enable_weights_cpu_backup,
```

```
        enable_lora=get_lora().enable_lora,
```

```
        has_lora_paths=bool(get_lora().lora_paths),
```

```
        weight_loader_disable_mmap=get_model().weight_loader_disable_mmap,
```

```
        weight_loader_drop_cache_after_load=(
```

```
            get_model().weight_loader_drop_cache_after_load
```

```

    ),
    has_custom_weight_loader=bool(get_model().custom_weight_loader),
    enable_torch_compile=get_exec().graph.enable_torch_compile,
    prefetch_num_threads=get_model().weight_loader_prefetch_num_threads,
)

```

python/sglang/srt/utils/common.py

通用工具中 MoE/DP 判定链与 cuda graph 桶大小计算函数全部去掉 server_args 参数；函数体早已只读 get_exec()/get_parallel()，参数是死重，且 TYPE_CHECKING 导入清理为 pass。

以下函数原先都接收 server_args: ServerArgs 参数，但函数体内从未读过它；
实际决策全部来自 published config (get_exec() / get_parallel() 等) ，
参数纯属死重，本 PR 将其删除，调用图不再假配置按参数流动。

```

def require_mlp_tp_gather():
    'Check if the input of MLP is obtained by all-gather rather than all-reduce.'
    from sglang.srt.runtime_context import get_exec, get_parallel

    # elastic-EP 扩容会在 published config 上改写 dp_size
    if get_parallel().config.enable_dp_attention:
        assert get_parallel().config.dp_size > 1, 'dp_size must be greater than 1'
        if get_exec().moe.elastic_ep_backend is not None:
            from sglang.srt.elastic_ep.elastic_ep import (
                elastic_expanded_world_enabled,
            )

            if elastic_expanded_world_enabled():
                return True
        # 其余分支基于 moe_dense_tp_size、enable_dp_lm_head、
        # get_moe_a2a_backend() 等 published config 判断，主体因篇幅省略；
        # 所有分支都只读配置、不再接收 server_args。
    return False

```

```

def require_gathered_buffer():
    # 是否需要 gathered buffer: 由 MLP 与 attention 两侧的判定共同决定；
    # require_attn_tp_gather 同样去掉了 server_args 参数。
    return require_mlp_tp_gather() or require_attn_tp_gather()

```

```

def require_mlp_sync():
    # MLP 同步需求: DP attention 开启，或需要 gathered buffer。
    from sglang.srt.runtime_context import get_parallel

    return get_parallel().config.enable_dp_attention or require_gathered_buffer()

```

```

def get_cuda_graph_batch_size_alignment() -> int:
    # cuda graph 批次对齐要求: 与 overlap、gathered buffer、attn cp 相关。

```

```

alignment = 1
if get_exec().overlap.enable_two_batch_overlap:
    alignment *= 2
if require_gathered_buffer():
    alignment *= get_parallel().attn_tp_size
if alignment % get_parallel().attn_cp_size != 0:
    alignment *= get_parallel().attn_cp_size
return alignment

```

```

def get_cuda_graph_max_batch_size(max_batch_size: int) -> int:
    # 将用户给定 max_batch_size 对齐到 cuda graph 桶大小。
    return ceil_align(max_batch_size, get_cuda_graph_batch_size_alignment())

```

python/sglang/srt/eplb/expert_location.py

EPLB 专家位置元数据的全部构造入口 (init_trivial / init_by_mapping / init_by_eplb / _init_common) 删除 server_args, 配置读取统一走 runtime_context; broadcast_global_expert_location_metadata 也不再通过 get_server_args() 取记录。

```

# ExpertLocationMetadata: 专家位置元数据, EPLB (Expert Parallel Load Balancer)
# 需要知道物理专家到逻辑专家的映射以及各 rank 的本地专家数量。
# 本 PR 前, 构造入口接收 server_args 仅用于在 _init_common 中读取
# exec.moe / parallel 等已发布配置; 现在直接读 bags, 参数全部删除。

```

```

class ExpertLocationMetadata:
    @staticmethod
    def init_trivial(model_config: ModelConfig, moe_ep_rank: int):
        # Trivial location: 逻辑专家 i 对应物理专家 i。
        common = ExpertLocationMetadata._init_common(model_config)

        if common is None:
            return None

        num_physical_experts = common['num_physical_experts']
        model_config_for_expert_location = common['model_config_for_expert_location']
        num_layers = model_config_for_expert_location.num_layers
        num_logical_experts = model_config_for_expert_location.num_logical_experts

        base_num_physical_experts = common['base_num_physical_experts']
        physical_to_logical_map = (
            torch.arange(0, base_num_physical_experts).repeat(num_layers, 1)
            % num_logical_experts
        )
        physical_to_logical_map = append_trivial_expert_slots(
            physical_to_logical_map,
            num_physical_experts - base_num_physical_experts,
            num_logical_experts,
        )

```

```
return ExpertLocationMetadata.init_by_mapping(  
    model_config,  
    physical_to_logical_map=physical_to_logical_map,  
    moe_ep_rank=moe_ep_rank,  
)
```

评论区精华

本 PR 的 review 评论区在提供材料中为空（统计显示 2 条 inline 评论，但内容未包含在给定上下文内），因此无法还原具体讨论。核心设计思路集中在 PR body 与三个 commit message 中：①未读参数是负债，会让 48 个文件继续持有 record 引用并诱使后续贡献者读取字段；②工厂命名应与实际读取来源一致，from_published_config 取代 from_server_args；③EPLBManager、LoRAManager、多模态 BaseMultimodalProcessor 改读 bags 后不再需要 handed record。这些决策共同服务于同一目标：让 ServerArgs 回归原始输入，运行时配置统一从 published config 读取。

- 暂无高价值评论线程

风险与影响

- 风险：
 - 回归风险（高）：60 个文件、48 个生产文件的参数删除依赖调用点逐一同步，任何遗漏或间接调用仍按旧签名传参都会在启动或运行期抛 TypeError；utils/common.py 的 TYPE_CHECKING 导入被清成 pass，若存在运行时才触发的 ServerArgs 引用会直接 NameError。
 - 时序风险（中）：EPLBManager、LoRAManager、多模态 BaseMultimodalProcessor 改为直接读 published config bags，必须保证调用时机在配置发布之后；历史上 PR#36254 就是为修复 Ray 启动时 readers 读 bags 的崩溃，说明该类改动对初始化顺序敏感。
 - 测试覆盖风险（中）：supplied-instance exposure ratchet 删除条目后，对未来重新引入 record 依赖的拦截面变小；PR CI 中 PR Test (Base) 与 AMD ROCm 运行显示失败，失败原因未在材料中说明，合并前需确认与本 PR 无关。
 - 兼容性：纯重构声明无行为变化，但无新增测试覆盖所有 48 个文件的调用路径，行为不变依赖人工验证。
- 影响：
 - 影响范围：60 个文件（+168/-280），覆盖模型加载启动路径（startup_weight_load.py / load_model_utils.py / model_runner.py）、MoE 与 EPLB（expert_location.py / eplb_manager.py）、调度器（scheduler.py / batch_result_processor.py）、多模态处理器、分布式通信（mooncake_transfer_engine.py）与分离服务（disagg_service.py）。
 - 用户影响：无，纯重构；对行为的影响仅在理论上存在回归可能。
 - 团队影响：配置体系向「函数从 bags 读叶子、ServerArgs 只做原始输入」收敛，降低参数误用与调用图误导，并为后续 PR#36255 等收尾工作铺路。

- 系统影响：删除 280 行、减少 record 引用存活，轻微降低调用链持引用成本（非性能动机）。
- 风险标记：60 文件跨模块改动，调用点遗漏回归风险，bag 发布时序敏感，Base/ROCM CI 失败待确认，ratchet 拦截面缩小

关联脉络

- PR #36253 config: resolution reads the declarations, not the fields: 同系列配置重构：resolution 链改读声明，ServerArgs 开始回归原始输入；本 PR 是其后续清理。
- PR #36254 config: the runtime readers take the published bags: 运行时读者改读 bags，本 PR 删除这些读者的未读参数并重命名工厂。
- PR #36255 config: ServerArgs holds the raw input: 同系列收尾：ServerArgs 只存原始输入，本 PR 减少了 48 个文件对 record 的持有。
- PR #36250 config: spell the parallel config tier at the call site: parallel 配置层显式化，与本 PR 同属 config 重构系列。