

PR #29232 完整报告

sgl-project/sglang

[Spec] Replace shared-infra dflash special-cases with capabilities (WAR barrier + seq_lens_cpu)

合并时间: 2026-06-28 15:02

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/29232>

执行摘要

- 一句话: 替换 DFLASH 硬编码检查为统一能力模型
- 推荐动作: 该 PR 的设计方向 (能力抽象) 值得学习, 但实现尚不完整 (needs_war_barrier 未实际引入, 仍保留 is_dflash() 判断)。建议在后续 PR 中完成能力属性注册, 并评估性能回归后优化 publish 路径或选择性恢复直接携带。

功能与动机

消除 scheduler 中对特定 speculative 算法的硬编码依赖, 使 WAR 屏障和 seq_lens 同步机制更具可扩展性, 方便后续支持更多 speculative 算法。

实现拆解

1. 调度器 WAR 屏障抽象 (scheduler.py): 将 _war_barrier_enabled 从硬编码的 not spec_algorithm.is_dflash() 改为通过 worker 能力属性 needs_war_barrier 控制, DFLASH 自身会覆写为 False, 其他算法默认启用。
2. 移除 DFLASH 直接携带路径 (overlap_utils.py): 删除 _resolve_spec_extras 和 resolve_seq_lens_cpu 中针对 DFLASH direct_carry_valid 的早期返回, 使 DFLASH 与其他 spec-v2 算法一样通过 FutureMap (publish_ready 事件 + 异步拷贝) 获取 seq_lens。
3. 精简 DFLASH 状态字段 (dflash_info_v2.py): 移除 cur_allocated_seq_lens_cpu、planning_seq_lens_cpu、planning_seq_lens_sum、direct_carry_valid 以及与之配套的 _prepare_committed_kv_lens_cpu_buf、_prepare_planning_kv_lens_cpu_buf 等字段, 简化了 buffer 分配逻辑 (needs_cpu_alloc 不再检查 is_pinned())。
4. 清理 worker 接口 (dflash_worker_v2.py): 从 _make_next_draft_input_prefill / _make_next_draft_input_decode 中删除 cur_allocated_seq_lens_cpu 参数, 并统一方法参数命名 (batch 替代 model_worker_batch)。

关键文件:

- python/sglang/srt/speculative/dflash_info_v2.py (模块 推测解码; 类别 source; 类型 core-logic): 核心提交流: 移除了 DFLASH 特定的 seq_lens 携带字段, 统一到通用 FutureMap 路径, 并简化了 buffer 分配逻辑。
- python/sglang/srt/managers/overlap_utils.py (模块 调度器; 类别 source; 类型 core-logic; 符号 _resolve_spec_extras, resolve_seq_lens_cpu): 移除了 DFLASH 的 direct_carry_valid 早期返回和 resolve_seq_lens_cpu 中的 DFLASH 特殊分支, 使其与通

用 spec-v2 路径一致。

- python/sglang/srt/speculative/dflash_worker_v2.py (模块 推测解码; 类别 source; 类型 core-logic; 符号 _validate_phase1_sampling_support) : 清理 worker 接口, 删除 cur_allocated_seq_lens_cpu 参数, 统一方法参数命名。
- python/sglang/srt/managers/scheduler.py (模块 调度器; 类别 source; 类型 core-logic) : 更新 WAR 屏障注释, 移除硬编码 TODO, 引入更通用的描述。

关键符号: _validate_phase1_sampling_support, _ensure_prepare_length_buffers, resolve_seq_lens_cpu, _resolve_spec_extras, forward_batch_generation, _make_next_draft_input_prefill, _make_next_draft_input_decode

关键源码片段

python/sglang/srt/speculative/dflash_info_v2.py

核心提交流: 移除了 DFLASH 特定的 seq_lens 携带字段, 统一到通用 FutureMap 路径, 并简化了 buffer 分配逻辑。

```
@dataclass
class DFlashDraftInputV2(SpecInput):
    # Legacy Eagle-shaped fields; DFLASH relays via FutureMap so these are unused.
    topk_p: torch.Tensor
    topk_index: torch.Tensor
    bonus_tokens: torch.Tensor
    new_seq_lens: torch.Tensor
    hidden_states: torch.Tensor
    verify_done: Optional[torch.cuda.Event] = None
    max_top_k: int = 1
    uniform_top_k_value: Optional[int] = None
    # 以下字段被删除:
    # cur_allocated_seq_lens_cpu, planning_seq_lens_cpu, planning_seq_lens_sum,
    # direct_carry_valid, _prepare_committed_kv_lens_cpu_buf, _prepare_planning_kv_lens_cpu_
    # buf
    reserved_seq_lens_cpu: Optional[torch.Tensor] = None
    reserved_seq_lens_sum: Optional[int] = None
    _prepare_batch_seq_lens_cpu_buf: Optional[torch.Tensor] = None
    _prepare_cur_kv_lens_cpu_buf: Optional[torch.Tensor] = None
    _prepare_nxt_kv_lens_cpu_buf: Optional[torch.Tensor] = None
    _prepare_cur_kv_lens_gpu_buf: Optional[torch.Tensor] = None
    _prepare_nxt_kv_lens_gpu_buf: Optional[torch.Tensor] = None

    # 由调度器在调度后填充
    future_indices: Optional[torch.Tensor] = None

    def _ensure_prepare_length_buffers(self, bs: int, device: torch.device | str) -> None:
        pin_memory = is_pin_memory_available(device)
        # needs_cpu_alloc 不再检查 is_pinned()
        def needs_cpu_alloc(buf: Optional[torch.Tensor]) -> bool:
            return buf is None or buf.numel() < bs
```

```

def needs_gpu_alloc(buf: Optional[torch.Tensor]) -> bool:
    return buf is None or buf.numel() < bs or str(buf.device) != str(device)

# 三个 CPU 缓存一起增长；容量是唯一的不变量
if needs_cpu_alloc(self._prepare_batch_seq_lens_cpu_buf):
    capacity = grown_capacity(self._prepare_batch_seq_lens_cpu_buf)
    self._prepare_batch_seq_lens_cpu_buf = torch.empty(
        (capacity,), dtype=torch.int64, device="cpu"
    )
    self._prepare_cur_kv_lens_cpu_buf = torch.empty(
        (capacity,), dtype=torch.int32, device="cpu", pin_memory=pin_memory
    )
    self._prepare_nxt_kv_lens_cpu_buf = torch.empty(
        (capacity,), dtype=torch.int32, device="cpu", pin_memory=pin_memory
    )
# GPU 缓存分配不变

```

python/sclang/srt/managers/overlap_utils.py

移除了 DFLASH 的 `direct_carry_valid` 早期返回和 `resolve_seq_lens_cpu` 中的 DFLASH 特殊分支，使其与通用 `spec-v2` 路径一致。

```

def resolve_seq_lens_cpu(self, batch: ScheduleBatch) -> None:
    # CPU mirror 由 needs_cpu_seq_lens 控制；选择退出的后端走下面的 GPU-only 路径。
    draft_input = batch.spec_info
    if draft_input is None:
        return

    fi = draft_input.future_indices
    if fi is None:
        return

    if self.publish_ready is not None:
        if _is_hip:
            # 临时应对：AMD MI355 上 Event.wait() 会降低 TPOT
            self.publish_ready.synchronize()
        else:
            self.publish_ready.wait()
    batch.seq_lens = self.new_seq_lens_buf[fi]

    if not self.needs_cpu_seq_lens:
        # GPU gather 保留 (SB.seq_lens 每轮 verify 必须推进)；
        # 跳过 .cpu() D2H。下游只使用 GPU 路径。
        batch.seq_lens_cpu = None
        batch.seq_lens_sum = None
        return

    if self.fwd_prepare_d2h_stream is None or self.publish_ready is None:
        batch.seq_lens_cpu = batch.seq_lens.cpu() # 引导 / 非 CUDA
        batch.seq_lens_sum = int(batch.seq_lens_cpu.sum())
        return

```

```
# 不在调度流上同步；在发布事件上阻塞一个私有流并拷贝到静态固定缓存。
self.fwd_prepare_d2h_stream.wait_event(self.publish_ready)
with torch.get_device_module(self.device).stream(self.fwd_prepare_d2h_stream):
    self.new_seq_lens_cpu_pinned.copy_(self.new_seq_lens_buf, non_blocking=True)
self.fwd_prepare_d2h_stream.synchronize()

batch.seq_lens_cpu = self.new_seq_lens_cpu_pinned[batch.req_pool_indices_cpu]
batch.seq_lens_sum = int(batch.seq_lens_cpu.sum())
```

评论区精华

测试者 dcw02 在 PR 合并后报告了性能回归：

- 使用 `bench_dflash_gsm8k_sweep.py` 基准测试，关闭 `direct_carry_valid` 后，`conc=1` 时吞吐从 1164.20 tok/s 下降到 1159.32 tok/s（小幅下降），`conc=32` 时从 15181.06 tok/s 下降到 14625.93 tok/s（约 3.6% 下降）。
- 尝试修复（通过 `scheduler WAR fallback` 替代）导致 `conc=32` 时进一步下降到 14625.93 tok/s（与 baseline 相比下降明显）。
- 作者未在后续提交中解决该回归，PR 已合并。该讨论表明统一路径引入了 `publish` 同步开销，可能需要在后续优化中恢复特定路径。
- 性能回归报告 (performance): PR 已合并，性能回归未被修复；建议后续优化 `publish` 路径或部分恢复 DFLASH 的特殊处理。

风险与影响

- 风险：性能回归风险（高）：移除 `direct_carry_valid` 后 DFLASH 必须等待 `FutureMap` 同步，导致 `conc` 较高时吞吐下降约 3.6% ~ 8%。内存分配风险（低）：`needs_cpu_alloc` 不再检查 `is_pinned()`，可能在特定配置下重复分配或使用非 `pin_memory` buffer，但影响可控。兼容性风险（中）：依赖 `FutureMap` 的路径要求 `publish_ready` 事件正确设置，若其他 `spec-v2` 算法未正确初始化可能导致死锁或错误。
- 影响：用户影响：DFLASH 用户在高并发场景可能观察到吞吐下降，低并发场景影响较小。系统影响：统一了 `spec-v2` 的同步路径，为后续算法接入提供了更一致的接口，降低了维护成本。团队影响：需要监控性能指标，确认是否需恢复部分 DFLASH 优化路径或优化 `FutureMap` 实现。
- 风险标记：核心路径变更，性能回归风险，未解决讨论

关联脉络

- PR #29464 Fix EAGLE draft hidden dim extraction and centralize spec helpers: 同为 speculative decoding 模块的重构，集中化了 `spec` 工具函数，与本 PR 的抽象方向一致。
- PR #29223 (perf): Shard Kimi-K2.5 Eagle3 draft fc + symm-mem AG: 也涉及 speculative decoding 的性能优化，与本 PR 的同步机制变化有潜在交互。