

PR #33926 完整报告

sgl-project/sglang

[DCP] Support decode context parallelism on the trtllm_mla decode path

合并时间: 2026-09-01 10:19

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/33926>

执行摘要

- 一句话: trtllm_mla 解码路径支持 DCP, 长上下文吞吐提升达 68%
- 推荐动作: 值得精读。重点看三处设计决策: 一是 `q_len == 1` 免全局因果边界的论证与其在 `_run_decode_kernel` 的门控实现 (理解 DCP 内核契约的关键); 二是 LSE 基数问题从 backend 内注册表方案演化到采纳 #34240 merge-site 约定的过程, 展示了跨 PR 设计协同的正确姿势; 三是 `_apply_dcp_cuda_graph_metadata` 三支对 global / local 长度视图的维护, 是 CUDA graph 捕获路径下容易出错但必须一致的典型样例。

功能与动机

PR body 说明: `trtllm_mla` 是 DCP 路线图 #29736 中的未指派任务项; 其 `forward_decode` 用全局序列长度构建页表和 `seq_lens` 并返回裸张量, 而 DCP 下模型期望 rank-local 的 (`out`, `lse`) 对用于跨 rank 合并。该问题无需显式请求即可触达: `arg_groups/overrides.py` 在 sm100 上为 `DeepseekV3ForCausalLM` 自动选择 `trtllm_mla`, 且 `forward_mla.py` 在 `dcp_enabled` 时无条件把 decode 路由到 `attn_mqa_for_dcp_decode`, 使 `--dcp-size > 1` 在 Blackwell 上直接撞上没有实现的路径。作者还明确提出了核心假设: `q_len == 1` 的单 token 解码不需要全局因果边界, 因为循环所有者规则已精确选中因果可见前缀。

实现拆解

实现按 5 步推进:

1. DCP 元数据管道提升到基类 (`trtllm_mla_backend.py`, +247/-9)。新增 `_get_dcp_local_seq_lens`、`_get_dcp_local_max_seq_len`、`_fill_dcp_block_kv_indices`、`_apply_dcp_cuda_graph_metadata`, 同时 `cutedsml_mla_backend.py` (-203 行) 与 `tokenspeed_mla_backend.py` (-188 行) 删除各自重复实现, 收尾两个子类里 `collapse into the base` 的 TODO。这些 helper 只依赖 `dcp_size / dcp_rank / page_size / req_to_token`, 与具体内核无关; 所有分支在 `dcp_enabled=False` 时均为 no-op, 非 DCP 路径保持原样。
2. `_run_decode_kernel` 收窄拒绝条件并转发 `return_lse`。基类把一票否决的 `cp_world > 1 or return_lse` 改为按需拒绝: 显式传 `causal_seqs` (需要全局因果边界) 或 `return_lse=False` 的 spec 路径 (单 token draft-extend 跳过跨 rank 合并) 才抛 `NotImplementedError`, 保证 `q_len == 1` 的普通解码不被误伤。同时修复 `CuteDsIMLABackend` 委托基类时丢弃 `return_lse` 的陷阱 (`thanhhao98` review 发现)。

3. 新增 `_forward_decode_dcp` 解码路径。DCP 下查询 /KV 准备镜像既有 decode 逻辑（FP8 KV 含 `_fused_set_kv_concat_q_fp8` 路径），内核返回 rank-local (`out, lse`)，用 `fixup_zero_kv_rows` 把不拥有任何切片的 rank 中和为 (0, -inf) 中性态，交给 `forward_mla.py` 跨 rank merge。新增 `_dummy_dcp_decode_for_autotune` 短路 FlashInfer MoE autotune 模拟解码，修复单节点 tp8 `--dcp-size 2` 冷缓存预热 workspace 溢出崩溃。
4. 四个正确性修复：LSE 基数（实测 trtllm-gen 已输出 base-2、cute-dsl 输出自然对数，作者最初在 backend 内建注册表 rebase，后因 #34240 merge-site 约定落地而弃用自家实现）；DCP 下禁用融合 bf16 KV 写入（`set_mla_kv_concat_q` 无 DCP 参数，会写虚拟行，回退 pool 的 `set_mla_kv_buffer`）；新增 `num_decode_q_heads = num_q_heads * attn_dcp_size` 修正 multi-CTA 计数器尺寸；修复 `draft_extend_v2` 分支 global/local 长度视图交换。
5. 测试配套。新增 `test_trtllm_mla_family_dcp_metadata.py` (250 行)：三后端共享 DCP 元数据断言、`_run_decode_kernel` 5 个拒绝 / 放行场景、`TestDcpDecodeLayout` 分区不变量；删除 `test_tokenspeed_mla_dcp_metadata.py` 并入 family 测试（原测试 patch tokenspeed 模块的 `get_parallel`，提升后会绕过真实路径）。B200 上 `test_dcp_layout_unit.py` 回归通过。

关键文件：

- `python/sglang/srt/layers/attention/trtllm_mla_backend.py`（模块 注意力后端；类别 source；类型 core-logic；符号 `_get_dcp_local_seq_lens`, `_get_dcp_local_max_seq_len`, `_fill_dcp_block_kv_indices`, `_apply_dcp_cuda_graph_metadata`）：PR 主战场：DCP 元数据管道（`_get_dcp_local_seq_lens` / `_get_dcp_local_max_seq_len` / `_fill_dcp_block_kv_indices` / `_apply_dcp_cuda_graph_metadata`）从子类提升到基类；`_run_decode_kernel` 收窄拒绝条件并转发 `return_lse`；新增 `_forward_decode_dcp` 与 `_dummy_dcp_decode_for_autotune`；同时修复 multi-CTA 计数器尺寸、DCP 下禁用融合 KV 写入等正确性问题。
- `python/sglang/srt/layers/attention/cutedsl_mla_backend.py`（模块 注意力后端；类别 source；类型 refactor；符号 `_get_dcp_local_seq_lens`, `_create_block_kv_indices`, `init_forward_metadata`, `_run_decode_kernel`）：删除与 tokenspeed 重复的约 190 行 DCP 元数据实现（原有 TODO 的收尾），仅保留 cute-dsl 内核调用与 decode 前向；`_run_decode_kernel` 的 `cp_world <= 1` 委托补上 `return_lse` 转发，避免 #34240 工作踩坑。
- `python/sglang/srt/layers/attention/tokenspeed_mla_backend.py`（模块 注意力后端；类别 source；类型 refactor；符号 `_get_dcp_local_seq_lens`, `_create_block_kv_indices`, `init_forward_metadata`, `forward_decode`）：同样删除约 180 行重复 DCP 元数据代码，autotune dummy-run 跳过逻辑改为复用基类 `_dummy_dcp_decode_for_autotune`，是本次提升重构的另一半。
- `test/registered/dcp/test_trtllm_mla_family_dcp_metadata.py`（模块 DCP 测试；类别 test；类型 test-coverage；符号 `_DCPMetadataTests`, `TestTRTLLMMLADCPMetadata`, `TestTokenspeedMLADCPMetadata`, `TestCuteDsIMLADCPMetadata`）：新增测试：把原

来 tokenspeed 专属的 DCP 元数据测试重命名并扩到三个后端，另加_run_decode_kernel 拒绝路径测试与 TestDcpDecodeLayout 分区不变量测试，钉死q_len == 1 假设所依赖的 rank 分区数学。

- test/registered/dcp/test_tokenspeed_mla_dcp_metadata.py (模块 DCP 测试；类别 test；类型 deletion；符号 TestTokenspeedMLADCPMetadata)：删除：内容并入 test_trtllm_mla_family_dcp_metadata.py；原测试 patch 的是 tokenspeed 模块的 get_parallel，元数据提升到基类后会绕过真实路径，必须随重构迁移。

关键符号：_forward_decode_dcp, _apply_dcp_cuda_graph_metadata, _dummy_dcp_decode_for_autotune, _get_dcp_local_seq_lens, _get_dcp_local_max_seq_len, _fill_dcp_block_kv_indices, _run_decode_kernel

关键源码片段

python/sglang/srt/layers/attention/trtllm_mla_backend.py

PR 主战场：DCP 元数据管道 (`_get_dcp_local_seq_lens` / `_get_dcp_local_max_seq_len` / `_fill_dcp_block_kv_indices` / `_apply_dcp_cuda_graph_metadata`) 从子类提升到基类；`_run_decode_kernel` 收窄拒绝条件并转发 `return_lse`；新增 `_forward_decode_dcp` 与 `_dummy_dcp_decode_for_autotune`；同时修复 multi-CTA 计数器尺寸、DCP 下禁用融合 KV 写入等正确性问题。

```
# python/sglang/srt/layers/attention/trtllm_mla_backend.py
# DCP 元数据管道：从 cutedsl / tokenspeed 两个子类提升到 TRTLLMMLABackend 基类。
# 这些 helper 只依赖 dcp_size / dcp_rank / page_size / req_to_token,
# 与具体内核无关，因此整个 trtllm_mla 家族共享一份实现。
```

```
def _get_dcp_local_seq_lens(self, seq_lens: torch.Tensor) -> torch.Tensor:
    """把全局序列长度按循环分片切成 rank-local 长度。
```

DCP 关闭时原样返回，保证非 DCP 路径是彻底的 no-op。

```
"""
```

```
parallel = get_parallel()
if not parallel.dcp_enabled:
    return seq_lens
return get_dcp_lens(seq_lens, parallel.dcp_size, parallel.dcp_rank).to(
    torch.int32
)
```

```
def _get_dcp_local_max_seq_len(self, max_seq_len: int) -> int:
    """rank-local 调度上界；空行时也必须返回正数（内核约束）。"""
    parallel = get_parallel()
    if not parallel.dcp_enabled:
        return max_seq_len
    local_max = max_seq_len // parallel.dcp_size + int(
        parallel.dcp_rank < max_seq_len % parallel.dcp_size
    )
    return max(local_max, 1)
```

```

def _apply_dcp_cuda_graph_metadata(
    self,
    bs: int,
    req_pool_indices: torch.Tensor,
    seq_lens: torch.Tensor,
    forward_mode: ForwardMode,
    metadata: TRTLLMMLADecodeMetadata,
):
    """DCP 变体的 CUDA graph 捕获 + 重放主体。

    每步一次把全局长度与 rank-local 长度刷新进 capture-stable 缓冲区，
    再在本 rank 的循环切片上重建页表，而不是每 MLA 层重算。
    """
    if forward_mode.is_target_verify():
        # verify 的 KV 长度 = 前缀长度 + draft token 数
        torch.add(
            seq_lens[:bs], self.num_draft_tokens, out=metadata.global_seq_lens_k
        )
        metadata.seq_lens_k.copy_(
            self._get_dcp_local_seq_lens(metadata.global_seq_lens_k)
        )
        local_seq_lens = metadata.seq_lens_k
    elif forward_mode.is_draft_extend_v2():
        num_tokens_per_req = self.num_draft_tokens
        metadata.max_seq_len_q = num_tokens_per_req
        metadata.sum_seq_lens_q = num_tokens_per_req * bs
        seq_lens = seq_lens[:bs]
        # 修复点：此前这一分支把全局长度写进 seq_lens_k（所有读者都按
        # rank-local 解读）且从不写 global_seq_lens_k，导致内核拿着全局
        # 长度去读只覆盖本 rank 切片的页表。
        metadata.global_seq_lens_k.copy_(seq_lens)
        metadata.seq_lens_k.copy_(self._get_dcp_local_seq_lens(seq_lens))
        local_seq_lens = metadata.seq_lens_k
    else:
        seq_lens = seq_lens[:bs]
        metadata.global_seq_lens_k.copy_(seq_lens)
        metadata.seq_lens_k.copy_(self._get_dcp_local_seq_lens(seq_lens))
        local_seq_lens = metadata.seq_lens_k

    self._fill_dcp_block_kv_indices(
        metadata.block_kv_indices, req_pool_indices[:bs], local_seq_lens
    )

```

test/registered/dcp/test_trtllm_mla_family_dcp_metadata.py

新增测试：把原来 tokenspeed 专属的 DCP 元数据测试重命名并扩到三个后端，另加 `_run_decode_kernel` 拒绝路径测试与 `TestDcpDecodeLayout` 分区不变量测试，钉死 `q_len == 1` 假设所依赖的 rank 分区数学。

test/registered/dcp/test_trtllm_mla_family_dcp_metadata.py

本 PR 的核心假设: $q_len == 1$ 解码时, 循环所有者规则选中的
rank-local 切片恰好等于全局因果可见前缀。下面两个性质是
跨 rank 合并不重复、不遗漏 token 的前提, 用可移植的数学断言钉死。

```
class TestDcpDecodeLayout(CustomTestCase):
    """rank-local 长度数学——decode 页表正是建立在这之上。"""

    SIZES = [1, 2, 3, 4, 8]
    LENS = list(range(0, 41))

    def test_ranks_partition_the_global_length(self):
        # 各 rank 的局部长度之和必须精确等于全局长度, 否则跨 rank 合并
        # 会系统性丢失或重复 KV 覆盖范围。
        lens = torch.tensor(self.LENS, dtype=torch.int32)
        for n in self.SIZES:
            total = sum(
                get_dcp_lens(lens, n, rank).to(torch.int64) for rank in range(n)
            )
            self.assertTrue(
                torch.equal(total, lens.to(torch.int64)),
                f"per-rank lengths do not sum to the global length at n={n}",
            )

    def test_newest_token_is_owned_by_exactly_one_rank(self):
        # 解码步只追加一个 token; 除非恰好一个 rank 的长度增长,
        # 跨 rank 合并会重复计数或丢掉该 token。
        for n in self.SIZES:
            for global_len in self.LENS[1:]:
                prev = torch.tensor([global_len - 1], dtype=torch.int32)
                cur = torch.tensor([global_len], dtype=torch.int32)
                grew = [
                    int(get_dcp_lens(cur, n, rank).item())
                    - int(get_dcp_lens(prev, n, rank).item())
                    for rank in range(n)
                ]
                self.assertEqual(sum(grew), 1, f"n={n}, global_len={global_len}")
                self.assertEqual(grew[(global_len - 1) % n], 1)
```

评论区精华

Review 中最有价值的交锋集中在 4 个正确性缺陷与测试组织:

- [thanhhao98](#) 在 `trtllm_mla_backend.py:1397` 指出 `_forward_decode_dcp` 缺 autotune dummy-run 守卫:

"Not GB300-only. Single node, tp8 `--dcp-size 2, --moe-runner-backend flashinfer_trtllm`, cold autotune cache -> dies in warmup ... Buffer overflow when allocating `trtllm_gen_softmax_workspace`" [qtris123](#) 确认后用 `_dummy_dcp_decode_for_autotune` 修复, 并以 `get_in_autotune_dummy_run()` 限定范

围。

- thanhhao98 发现 `CuteDslMLABackend._run_decode_kernel` 的 `cp_world <= 1` 分支丢弃 `return_lse`:

"Unreachable today only because `forward_decode` shadows the DCP path - but it is a trap for the #34240 work this defers to." 作者在 46c321c 补齐转发。

- kpham-sgl 指出 `causal_seqs is not None` 代理 `q_len > 1` 会被 `draft_extend_v2` 破坏, 要求补拒绝测试; 作者在 d46cb16ad / e1e71ee43c 用 `return_lse` 补齐门控, 覆盖单 token draft-extend 漏网场景。
- kpham-sgl 指出 `draft_extend_v2` 分支长度视图与兄弟分支相反 (`seq_lens_k` 装全局、`global_seq_lens_k` 不填); b57f3b7 修复, 作者说明该分支只服务 draft token、不影响 DCP 性能。
- 测试组织上 kpham-sgl 建议把散落的 DCP 布局测试并入 family 元数据测试文件并删除独立 hooks 测试, 作者照办并顺带去掉与 `test_dcp_layout_unit.py` 重复的用例。另有 kpham-sgl 要求清理 AI 生成注释, 最终由 kpham-sgl 亲自提交两轮注释修剪 (81ed959e、76c2bbe8), 声明 AST 不变。
- DCP 解码缺 autotune dummy-run 守卫导致预热崩溃 (correctness): qtris123 引入 `_dummy_dcp_decode_for_autotune` 并限定在 `get_in_autotune_dummy_run()` 内, 真实解码与 CUDA graph 捕获不受影响, 作者声明不影响精度与性能。
- CuteDsl 委托基类时丢弃 `return_lse` (correctness): qtris123 在 46c321c 转发全部 hook kwargs (含 `return_lse` 等), 非 DCP 默认仍为 `False`, 既有 cute-dsl 结果不受影响。
- `causal_seqs` 代理 `q_len > 1` 在 `draft_extend_v2` 失效 (correctness): qtris123 在 d46cb16ad / e1e71ee43c 双重门控: `causal_seqs` 或 `not return_lse` 任一成立即拒绝, 覆盖多 token verify、显式因果边界、单 token draft-extend 三类 spec 路径; 配套 5 个拒绝 / 放行测试。
- draft-extend 分支全局 / 局部长度视图交换 (correctness): qtris123 在 b57f3b7 修复: 该分支现在同时维护 `global_seq_lens_k` 与 `rank-local seq_lens_k`, 与 `target-verify / decode` 分支一致; 作者说明该分支只服务 draft token, 不影响 DCP 性能。
- DCP 布局测试合并进 family 测试 (testing): 合并完成: `TestDcpDecodeLayout` 保留分区求和与单 rank 增长两条不变量, 全部 MLA DCP attn 布局测试集中在 family 文件。
- AI 生成注释清理 (style): 注释修剪完成, 声明 AST 不变 (剥离 docstring 后无代码差异)。

风险与影响

- 风险: 主要风险点:
- `_run_decode_kernel` 门控复杂: 拒绝条件依赖 `q_len / causal_seqs / return_lse` 三重信号的组合, 未来新增 spec 路径或 `forward_extend` 调用约定变化时容易漏配; 当前测试覆盖 5 个场景, 但 spec × DCP 组合在 flashinfer 上游本就被拒, 整体风险可控。
- 单模型验证: 精度与性能均只在 `DeepSeek-V2-Lite` 上实测, 作者自述该模型最小、解码步便宜、collective 占比更高, 结论可能低估 DCP 收益, 但生产级 MLA 模型未覆盖仍是盲区。

- 通信后端覆盖不全：仅测 `ag_rs`，`a2a` 未覆盖，跨 rank 合并的数值行为在 `a2a` 下有不同实现路径。
- CUDA graph 双阶段分支：捕获 (`_init_cuda_graph_metadata`) 与重放 (`_apply_dcp_cuda_graph_metadata`) 两个阶段都加了 DCP 分支，`init_forward_metadata` 的 eager 路径还有视图切换，三处必须保持一致，后续改动容易只改一处。
- 非 DCP 路径回归面：`_fused_set_kv_concat_q` 条件新增 `not get_parallel().dcp_enabled`、计数器分配改走 `num_decode_q_heads` (非 DCP 时等于 `num_q_heads`)，作者用控制实验证明门控零开销，但这两个位置仍是回归敏感点。
- 影响：影响范围与程度：
 - 用户侧：Blackwell (sm100 / B200) 上 MLA 模型用户从此可用 `--dcp-size > 1`，此前该组合直接不可用或静默产出错误结果；长上下文 ($\geq 128K$) 高并发 (≥ 16) 收益最大，256K 下即使并发 1-2 DCP 已领先 TP。
 - 系统侧：三个 MLA 后端 (`trtllm_mla` / `cutedsl_mla` / `tokenspeed_mla`) 消除约 380 行重复代码，后续 DCP 元数据演进只需改基类一处；`TRTLLMMLABackend` 成为家族共享的 DCP 元数据源。
 - 团队侧：与 #34240 的 LSE merge-site 约定完成对齐，避免两套独立 LSE 基数修正并存；该 PR 标记 `release-highlight`，是 DCP 路线图在 MLA 家族的关键落子。
 - 风险标记：核心解码路径变更，spec-decode 与 DCP 组合受限，仅单模型硬件验证，CUDA graph 双阶段分支，`a2a` 通信后端未覆盖

关联脉络

- PR #34240 [DCP] Move DCP LSE base to the merge site: review 讨论中明确提及：该 PR 把 DCP LSE 基数从 `backends` 移到 `merge` 站点 (`is_mla_dcp_lse_base_on_e` 按 `sglang attention-backend` 名解析)；作者 rebase 后弃用自家 `backend` 内注册表方案 (`f54f050`)，避免两套基数修正并存与二次 rebase 回归。
- PR #35245 DCP decode page-table refactor: 作者评论中提及该 PR 改动 DCP decode 页表，合并前做了 B200 前后 A/B 验证 (`pre d46cb16ad5` vs `post e1e71ee43c`)，确认 `kvcache` 重构不破坏本 PR。
- PR #29736 DCP roadmap: DCP 路线图 issue: PR body 开头即说明本 PR 是其中未指派的 `trtllm_mla` decode 路径任务项，同时相关 #26432。
- PR #37307 fix(unified-memory): forward the KV-index translator through every wrapper backend: 同属 attention 后端管道一致性修复：该 PR 修复 `wrapper` 后端未转发 `kv_index_translator` 导致 MLA 前缀缓存读错，与本 PR 的 `return_lse` 参数透传是同一类 '后端间管道透传漏项' 问题，说明 MLA 家族后端正在集中加固。