

PR #6469 完整报告

verl-project/verl

[fsdp, megatron, trainer] feat: add top-k distillation overlap metrics

合并时间: 2026-05-26 18:58

原文链接: <http://prhub.com.cn/verl-project/verl/pull/6469>

执行摘要

- 一句话: 添加 top-k 蒸馏重叠诊断指标
- 推荐动作: 值得精读, 特别是对 on-policy 蒸馏研究和调试感兴趣的同学。设计上保持了非侵入性 (不改变损失计算), 在 FSDP 和 Megatron 两套引擎中实现了相同的诊断逻辑, 是跨后端功能一致的范例。Megatron 中通过 all_gather 获取全局 student top-k 的方式值得注意, 其 world_size=1 的保护建议虽未采纳, 但可作为后续优化点。

功能与动机

参照 PR body: 旨在跟踪 on-policy top-k 蒸馏中 student 的 top-k 与 teacher 的 top-k 的重叠情况, 帮助研究人员分析 OPD 是否在高概率 token 上对齐 student 与 teacher。基于论文 'Rethinking On-Policy Distillation of Large Language Models' (arXiv:2604.13016)。

实现拆解

1. FSDP 路径 ([verl/trainer/distillation/fsdp/losses.py](#)) : 在 compute_forward_kl_topk 中新增 torch.topk 获取 student 的 top-k 索引, 通过广播比较计算 teacher 与 student 的 token 级重叠掩码, 进而得到 overlap_count 和 overlap_token_advantage。这些张量随原有 distillation_losses 等一并返回至 model_output。
2. Megatron 路径 ([verl/trainer/distillation/megatron/losses.py](#)) : 在 ForwardKLTOPKFunction.forward 中, 对不同 TP rank 的局部 student top-k 进行 all_gather 并拼接为全局 top-k, 再计算与 teacher 全局 top-k 的重叠指标。新增的 overlap_count 和 overlap_token_advantage 通过 ctx.mark_non_differentiable 标记为非可微。
3. 公共聚合层 ([verl/trainer/distillation/losses.py](#)) : 在 compute_forward_kl_topk 中增加了 overlap_count 和 overlap_token_advantage 的条件提取与 padding 对齐, 然后计算 distillation/overlap_ratio (平均重叠比例) 和 distillation/overlap_token_advantage (重叠 token 上的平均负 KL 贡献), 并合并到原有的 distillation_metrics 字典。
4. 测试配套: 在 tests/workers/test_distillation_topk_symmetry_on_cpu.py 中新增两个测试函数验证 FSDP 路径的重叠指标数值正确性和指标聚合逻辑; 在 tests/utils/test_special_megatron_kl_loss_tp.py 中扩展了 TP 正确性测试, 对比 Megatron 与 FSDP 输出的重叠指标。

5. 文档更新：在 docs/algos/opd.md 中增加了新指标的说明和解释，并更新了蒸馏流程图和引用文献。

关键文件：

- verl/trainer/distillation/fsdp/losses.py (模块 FSDP 蒸馏；类别 source；类型 core-logic；符号 compute_forward_kl_topk)：FSDP 蒸馏路径的核心逻辑，新增 student top-k 索引计算和重叠指标
- verl/trainer/distillation/megatron/losses.py (模块 Megatron 蒸馏；类别 source；类型 core-logic；符号 ForwardKLTOPKFunction.forward)：Megatron 蒸馏路径，通过 all_gather 实现跨 TP rank 的 student top-k 全局化并计算重叠指标
- verl/trainer/distillation/losses.py (模块 蒸馏聚合；类别 source；类型 core-logic；符号 compute_forward_kl_topk)：公共蒸馏损失聚合层，将底层引擎返回的重叠张量转换为最终指标
- tests/workers/test_distillation_topk_symmetry_on_cpu.py (模块 蒸馏测试；类别 test；类型 test-coverage；符号 _nested_from_rows, test_forward_kl_topk_emits_overlap_metrics, test_forward_kl_topk_metric_aggregation_for_overlap_outputs)：新增两个测试函数验证 FSDP 路径的重叠指标数值和聚合逻辑
- tests/utis/test_special_megatron_kl_loss_tp.py (模块 Megatron 测试；类别 test；类型 test-coverage)：扩展 Megatron TP 正确性测试，对比 FSDP 与 Megatron 的重叠指标
- docs/algos/opd.md (模块 蒸馏文档；类别 docs；类型 documentation)：更新 OPD 文档，说明新指标的含义和用途
- README.md (模块 项目文档；类别 docs；类型 documentation)：轻微修改，可能更新了蒸馏相关引用

关键符号：compute_forward_kl_topk (FSDP), compute_forward_kl_topk (common), ForwardKLTOPKFunction.forward, _nested_from_rows, test_forward_kl_topk_emits_overlap_metrics, test_forward_kl_topk_metric_aggregation_for_overlap_outputs

关键源码片段

verl/trainer/distillation/fsdp/losses.py

FSDP 蒸馏路径的核心逻辑，新增 student top-k 索引计算和重叠指标

```
def compute_forward_kl_topk(
    student_logits: torch.Tensor,
    teacher_topk_log_probs: torch.Tensor,
    teacher_topk_ids: torch.Tensor,
    config: DistillationConfig,
    data_format: str,
) -> dict:
    # ... 前置处理：断言嵌套张量、SP 拆分、log_softmax ...
    student_log_probs = F.log_softmax(student_logits, dim=-1)
    # 计算 student 自身的 top-k 索引（用于重叠分析）
    student_topk_ids = torch.topk(student_log_probs, k=teacher_topk_ids.shape[-1], dim=-1).
```

```

indices
# 按 teacher 的 top-k 位置收集 student log prob
student_topk_log_probs = torch.gather(student_log_probs, dim=-1, index=teacher_topk_ids)
student_mass = student_topk_log_probs.exp().sum(dim=-1)
teacher_mass = teacher_topk_log_probs.exp().sum(dim=-1)
# ... 可选 clamp ...
distillation_losses = kl_divergence(log_q=student_topk_log_probs, log_p=teacher_topk_log_probs)

# 诊断: 计算 teacher 与 student 的 top-k 重叠
# overlap_mask: teacher 的每个 top-k token 是否出现在 student 的 top-k 中
overlap_mask = (teacher_topk_ids.unsqueeze(-1) == student_topk_ids.unsqueeze(-2)).any(dim=-1)
overlap_count = overlap_mask.sum(dim=-1) # 每个位置的重叠 token 数
# 计算每个 teacher token 的 KL 贡献, 仅关注重叠部分的负值
token_kl = teacher_topk_log_probs.exp() * (teacher_topk_log_probs - student_topk_log_probs)
overlap_token_advantage_sum = (-token_kl * overlap_mask).sum(dim=-1)
overlap_token_advantage = overlap_token_advantage_sum / overlap_count.clamp_min(1)
overlap_token_advantage = torch.where(overlap_count > 0, overlap_token_advantage, torch.zeros_like(overlap_token_advantage))

return {
    "distillation_losses": distillation_losses,
    "student_mass": student_mass,
    "teacher_mass": teacher_mass,
    "overlap_count": overlap_count,
    "overlap_token_advantage": overlap_token_advantage,
}

```

verl/trainer/distillation/losses.py

公共蒸馏损失聚合层, 将底层引擎返回的重叠张量转换为最终指标

```

def compute_forward_kl_topk(config, distillation_config, model_output, data) -> tuple:
    # ... 提取基础张量 ...
    overlap_count = model_output.get("overlap_count")
    overlap_token_advantage = model_output.get("overlap_token_advantage")
    if overlap_count is not None and overlap_token_advantage is not None:
        overlap_count = no_padding_2_padding(overlap_count, data)
        overlap_token_advantage = no_padding_2_padding(overlap_token_advantage, data)

    # ... response_mask 处理 ...

    overlap_metrics = {}
    if overlap_count is not None and overlap_token_advantage is not None:
        assert overlap_count.shape == overlap_token_advantage.shape == response_mask_bool.shape
        valid_overlap_count = overlap_count[response_mask_bool]
        k = distillation_config.distillation_loss.topk
        assert k is not None

```

```

# overlap_ratio: 平均重叠比例 = 重叠数 / k
overlap_metrics["distillation/overlap_ratio"] = (valid_overlap_count.float().mean() / k).item()
overlap_position_mask = response_mask_bool & (overlap_count > 0)
if overlap_position_mask.any():
    overlap_metrics["distillation/overlap_token_advantage"] = (
        overlap_token_advantage[overlap_position_mask].mean().item()
    )
else:
    overlap_metrics["distillation/overlap_token_advantage"] = 0.0

# ... 原有的 mass 指标 ...
distillation_metrics = {
    "distillation/student_mass": ...,
    ...
    **overlap_metrics,
}
# ... clamp 和返回 ...

```

评论区精华

Megatron all_gather 的冗余问题:

- gemini-code-assist[bot] 在 review 中指出, 当 world_size=1 (即 TP=1) 时, 调用 all_gather 是不必要的, 可能导致额外通信开销甚至运行时错误, 建议用 if world_size > 1 保护。
- PR 提交者未在后续 commit 中修改该处, 合并者 wuxibin89 仍批准了 PR。可能认为 Megatron 通常运行在 TP>1 环境, 或该路径在 TP=1 时会被早期返回避免。
- Megatron all_gather 在 world_size=1 时的冗余 (performance): PR 未采纳该修改, 合并者 wuxibin89 仍批准了 PR, 推测 Megatron 通常运行在 TP>1 环境或该路径在 TP=1 时不会执行。

风险与影响

- 风险:
 1. Megatron all_gather 对 TP=1 的兼容性: 当 tensor model parallel size 为 1 时, all_gather 可能会产生空集合或未初始化通信后端问题, 但实际场景中 Megatron 很少使用 TP=1, 因此风险较低。
 2. 新增计算开销: 每个 token 需要计算 student 的 top-k 索引和重叠掩码, 但仅在前向中执行且不改变梯度, 对整体训练吞吐影响可忽略。
 3. 兼容性: 新增的 overlap_count 和 overlap_token_advantage 键仅在 distillation_loss.topk 启用时存在, 旧版 compute_forward_kl_topk 调用者通过 model_output.get 安全处理, 因此向后兼容。- 影响: 用户: 无 API 或配置变更, 现有启用 distillation_loss.topk 的蒸馏训练会自动记录两个新指标, 不启用则无影响。系统: 无性能或稳定性影响, 指标仅用于监控。团队: 新指标有助于诊断 OPD 过程中 student 与 teacher 的对齐程度, 为后续算法调优提供依据。
- 风险标记: Megatron all_gather 在 TP=1 时冗余, 新计算开销小但未评估

关联脉络

- 暂无明显关联 PR