

PR #51727 完整报告

vllm-project/vllm

[Bugfix] Fix DeepSeek V4/3.2 tokenizer vocab size overcount crashing guided decoding

合并时间: 2026-08-11 07:00

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/51727>

执行摘要

- 一句话: 修复 DSV4/3.2 词表重复计数引发的引导解码崩溃
- 推荐动作: 值得精读。这是“小改动、深根因”的典型 bugfix: tokenizer 包装类覆写 `__len__` 会静默破坏结构化输出的 grammar bitmask 宽度, 而上游库语义变化 (transformers 5.x) 让曾经的 workaround 从必要变为有害。建议关注两点: 一是作者选择整体删除而非修正公式, 避免后续再次失配; 二是测试被回退的取舍——若能保留一个极简回归测试 (构造 id 落在基础词表内的 added token) 会更有长期保障。对结构化输出或 DeepSeek 模型维护者, 本 PR 是必读的参考案例。

功能与动机

PR body 与两个关联 issue 给出了完整链条: DeepseekV4Tokenizer.`__len__` 返回 `tokenizer.vocab_size + len(get_added_vocab())`, 对 DeepSeek-V4-Flash-0731 得到 $128000 + 1283 = 129283$, 但其中 3 个 added token 的 id 低于 128000, 真实词表应为 129280, 恰好等于 `config.vocab_size`, 因此该覆盖重复计数; v3.2 同款模式源于早期 HF fast tokenizer 的 `len()` 不含 added tokens, 而 transformers 5.x TokenizersBackend 已包含, 覆盖“既冗余又错误”。issue #50924 描述首个 guided-decoding 请求即令 EngineCore 崩溃 (tensor 4040 vs 4041), 且无投机配置时引导解码正常, 说明过宽 bitmask 只在投机解码调度路径出现; issue #51467 在 8xH20 上复现了同一 tensor size mismatch。

实现拆解

1. 定位根因: 在 `vllm/tokenizers/deepseek_v4.py` 与 `vllm/tokenizers/deepseek_v32.py` 的包装类内发现自定义 `__len__`, 其实现为 `vocab_size + len(get_added_vocab())`。由于部分 added token 的 id 已落在基础词表区间内, 该公式重复计数; 对 DeepSeek-V4-Flash-0731 实测多算 1 个 token (129283 vs 129280)。
2. 删除覆盖实现: 两个文件均删除 `__len__` 方法, 并移除仅为它服务的 `added_vocab_size`、`tokenizer_vocab_size` 局部变量; `get_added_vocab()` 返回副本的逻辑保持不变。删除后 `len(tokenizer)` 回落到 TokenizersBackend 原生实现, 与 `config.vocab_size` 对齐。
3. 测试改动与回退: 早期提交曾在 `tests/tokenizers_/test_deepseek_v4.py` 增加参数化测试 `test_deepseek_tokenizer_preserves_vocab_size`, 并引入 `tokenizers`、`TokenizersBackend` 等 import; reviewer 认为该测试对本修复非必需且产生冗余 import, 作者最终回退全部测试改动, PR 收尾为 2 个源码文件的纯删除。

4. 验证：作者按 issue #50924 的复现命令（TP=4、fp8 KV cache、dspark、guidance 后端）验证引导解码不再崩溃，并确认 tool calling 正常。

关键文件：

- vllm/tokenizers/deepseek_v4.py（模块 分词器；类别 source；类型 core-logic；符号 len）：主修复文件：删除 `__len__` 覆盖后，DeepSeek-V4-Flash-0731 的 `len(tokenizer)` 从错误的 129283 修正为 129280，与 `config.vocab_size` 一致，grammar bitmask 宽度不再多 1，直接解决 issue #50924 / #51467 的 EngineCore 崩溃。
- vllm/tokenizers/deepseek_v32.py（模块 分词器；类别 source；类型 core-logic；符号 len）：同源修复：DeepSeek V3.2 tokenizer 继承同一过时模式（`vocab_size + len(added_vocab)`），且该覆盖在 transformers 5.x 下已冗余；一并删除避免 V3.2 模型出现同类引导解码故障。

关键符号：get_deepseek_v4_tokenizer, get_deepseek_v32_tokenizer, len

关键源码片段

vllm/tokenizers/deepseek_v4.py

主修复文件：删除 `__len__` 覆盖后，DeepSeek-V4-Flash-0731 的 `len(tokenizer)` 从错误的 129283 修正为 129280，与 `config.vocab_size` 一致，grammar bitmask 宽度不再多 1，直接解决 issue #50924 / #51467 的 EngineCore 崩溃。

```
def get_deepseek_v4_tokenizer(tokenizer: HfTokenizer) -> HfTokenizer:
    """
    Wraps a tokenizer to use the custom DeepSeek V4 chat template encoding.
    """
    dsv4_tokenizer = copy.copy(tokenizer)

    # 旧实现在此处还计算了 added_vocab_size = len(added_vocab) 和
    # tokenizer_vocab_size = tokenizer.vocab_size，它们只服务于
    # __len__ 覆盖；覆盖删除后这两个临时变量也一并移除。
    added_vocab = tokenizer.get_added_vocab()

    class _DeepseekV4Tokenizer(tokenizer.__class__): # type: ignore
        def apply_chat_template(
            self,
            messages: list["ChatCompletionMessageParam"],
            tools: list[dict[str, Any]] | None = None,
            **kwargs,
        ) -> str | list[int]:
            # 聊天模板编码逻辑（thinking / reasoning_effort 处理）未改动，
            # 本 PR 只涉及词表长度，保留原实现。
            thinking = kwargs.get("thinking")
            enable_thinking = kwargs.get("enable_thinking")
            thinking_enabled = bool(thinking) or bool(enable_thinking)
            if "thinking" not in kwargs and "enable_thinking" not in kwargs:
                thinking_enabled = True
            thinking_mode = "thinking" if thinking_enabled else "chat"
```

```

conversation = kwargs.get("conversation", messages)
messages = conversation.copy()
if tools is not None and len(tools) > 0:
    messages.insert(0, {"role": "system"})
    messages[0]["tools"] = tools # type: ignore[typeddict-unknown-key]

reasoning_effort = kwargs.get("reasoning_effort")
if not isinstance(reasoning_effort, str):
    reasoning_effort = "high" if thinking_enabled else None
elif reasoning_effort == "none":
    thinking_mode = "chat"
    reasoning_effort = None
elif reasoning_effort == "max":
    reasoning_effort = "max"
elif reasoning_effort in ("low", "minimal", "medium"):
    reasoning_effort = "low"
else:
    reasoning_effort = "high"

encode_config = dict(
    thinking_mode=thinking_mode,
    drop_thinking=kwargs.get("drop_thinking", True),
    reasoning_effort=reasoning_effort,
)

prompt_str = encode_messages(messages, **encode_config) # type: ignore

if kwargs.get("tokenize", True):
    tokenizer_kwargs = {
        k: kwargs[k] for k in ("truncation", "max_length") if k in kwargs
    }
    return self.encode(
        prompt_str,
        add_special_tokens=False,
        **tokenizer_kwargs,
    )

return prompt_str

def num_special_tokens_to_add(self) -> int:
    return len(self.encode(""))

# --- 已删除的 __len__ 覆盖 ---
# 旧逻辑: return tokenizer_vocab_size + added_vocab_size
# 问题: 部分 added token 的 id 已落在基础词表区间内
# (DeepSeek-V4-Flash-0731 的 1283 个 added token 中有 3 个
# id < 128000), 导致 128000 + 1283 = 129283, 而真实词表为
# 129280, 恰好等于 config.vocab_size。
# 结论: transformers 5.x 的 TokenizersBackend 中 len() 原生

```

```
# 已包含 added tokens, 因此该覆盖既冗余又错误; 删除后
# len(tokenizer) 与 config.vocab_size 对齐, grammar bitmask
# 宽度不再多 1。
```

```
def get_added_vocab(self) -> dict[str, int]:
    return added_vocab.copy()
```

```
def __reduce__(self):
    return get_deepseek_v4_tokenizer, (tokenizer,)
```

```
_DeepseekV4Tokenizer.__name__ = f"DSV4{tokenizer.__class__.__name__}"
dsv4_tokenizer.__class__ = _DeepseekV4Tokenizer
return dsv4_tokenizer
```

vllm/tokenizers/deepseek_v32.py

同源修复: DeepSeek V3.2 tokenizer 继承同一过时模式 (`vocab_size + len(added_vocab)`), 且该覆盖在 transformers 5.x 下已冗余; 一并删除避免 V3.2 模型出现同类引导解码故障。

```
def get_deepseek_v32_tokenizer(tokenizer: HfTokenizer) -> HfTokenizer:
    """
    Wraps a tokenizer to use the custom DeepSeek V3.2 chat template encoding.
    """
    dsv32_tokenizer = copy.copy(tokenizer)
```

```
# apply_chat_template 与 v4 版逻辑结构一致且未改动, 为聚焦本 PR
# 的改动不再重复展开。旧代码在这里计算 added_vocab_size 与
# tokenizer_vocab_size, 它们仅服务于被删除的 __len__ 覆盖。
added_vocab = tokenizer.get_added_vocab()
```

```
class _DeepseekV32Tokenizer(tokenizer.__class__): # type: ignore
    def num_special_tokens_to_add(self) -> int:
        return len(self.encode(""))
```

```
def get_added_vocab(self) -> dict[str, int]:
    return added_vocab.copy()
```

```
def __reduce__(self):
    return get_deepseek_v32_tokenizer, (tokenizer,)
```

```
# 已删除的 __len__ 覆盖 (旧注释: </think> 是 DeepseekV32 的
# added token)。原实现与 v4 相同: tokenizer_vocab_size +
# added_vocab_size, 同样会把基础词表内已存在的 added token
# 重复计数; 且 transformers 5.x TokenizersBackend 的 len()
# 已包含 added tokens, 所以该覆盖整体移除, 行为回落基类。
```

```
_DeepseekV32Tokenizer.__name__ = f"DSV32{tokenizer.__class__.__name__}"
dsv32_tokenizer.__class__ = _DeepseekV32Tokenizer
return dsv32_tokenizer
```

评论区精华

核心讨论围绕“新增单元测试是否必要”展开：

- reviewer yewentao256 在 tests/tokenizers_/test_deepseek_v4.py 的 diff 上留言：“I am thinking that we might not need this unit test for this specific update”，并另条指出“NIT: now becomes redundant import”；作者回复“Done, thanks~”，最终提交回退测试改动。
- reviewer chaunceyjiang 提问“Have you tested whether tool calling is working properly?”，作者答复“Yes the tool calling part works as expected.”
- 自动化 review 无阻塞项：Codex 给出“Didn't find any major issues”，Claude 因 fork 源 PR 自动跳过；yewentao256 最终 APPROVED。
- 新增单元测试是否必要 (testing): 作者回复“Done, thanks~”，最终提交回退了测试与相关 import，PR 收尾为纯源码删除。
- tool calling 是否正常 (question): sfeng33 答复：“Yes the tool calling part works as expected.”
- 自动化 review 结论 (other): 无阻塞项，PR 由作者合并。

风险与影响

- 风险：行为依赖底层实现：删除 `__len__` 后，`len()` 完全委托给 transformers `TokenizersBackend`；本修复的正确性建立在 transformers 5.x 的 `len()` 已包含 added tokens 这一语义上，若未来版本或个别 tokenizer 实现改变该行为，词表大小会被低估（与当前 overcount 方向相反）。缺少回归测试：最终版本删除了唯一的单元测试，同类 `__len__` 覆盖若被再次引入，CI 无法拦截；grammar bitmask、采样、投机草稿等多条路径都会静默受影响。跨模型影响面：`len(tokenizer)` 被结构化输出、推理 parser 等消费，所有 DeepSeek V3.2/V4 模型的该值都会从“多 1~N”修正为与 `config.vocab_size` 一致；方向正确，但依赖外部对旧行为的假设时可能出现边际兼容问题。
- 影响：对用户：修复 issue #50924 / #51467 中 dspark 投机解码 + guided decoding（JSON schema、tool_choice、reasoning parser）首请求即 EngineCore 崩溃的问题，不再出现整引擎 500 与分钟级重启。对系统：零运行时开销（纯删除），无新依赖、无配置变更；tokenizer 初始化少两次属性求值。对团队：为 DeepSeek 模型族 tokenizer 与 transformers 5.x 语义对齐提供范例，提示后续任何覆写 `__len__` / `vocab_size` 的包装类都需与 `TokenizersBackend` 行为核对。
- 风险标记：缺少测试覆盖，依赖 transformers 版本行为，影响 DeepSeek V3.2/V4 全系

关联脉络

- PR #51668 Bump Transformers version to 5.15.0: 本次修复的核心论据是 transformers 5.x `TokenizersBackend` 的 `len()` 已包含 added tokens；该升级使旧 `__len__` 覆盖从必要变为冗余，是此次删除的前提背景。
- PR #51768 [Bugfix] Guard DeepSeek V4 MRV1 piecewise CUDA graphs: 同属 DeepSeek V4 在 v1 引擎下的稳定性修复，与该 PR 共同完善 DSV4 部署可靠性。

- PR #51145 [Bugfix][ROCm] Fix DeepSeek V4 DSpark probabilistic startup: 同样修复 DeepSeek V4 + dspark 投机解码路径上的崩溃，复现环境与 issue #50924 一致。