

PR #23969 完整报告

sgl-project/sglang

[plugin][distributed] use active platform's backend in `get\_default\_distributed\_backend`

合并时间: 2026-06-12 11:05

原文链接: <http://prhub.com.cn/sgl-project/sglang/pull/23969>

执行摘要

- 一句话: `get_default_distributed_backend` 委托平台接口
- 推荐动作: 值得所有涉及分布式初始化和平台插件的开发者精读。本 PR 展示了如何将核心组件迁移到插件友好接口而不破坏向后兼容。重点关注委托决策与 fallback 机制。

功能与动机

`get_default_distributed_backend(device)` 目前使用硬编码 `_DEVICE_TO_DISTRIBUTED_BACKEND` 字典, out-of-tree 平台插件必须修补该字典。平台接口已声明 `get_torch_distributed_backend_str()` 作为真实来源, 此 PR 对齐 SRT 分布式初始化与平台接口, 同时保留字典作为安全 fallback。

实现拆解

1. 移动字典并修改默认接口: 将 `_DEVICE_TO_DISTRIBUTED_BACKEND` 从 `parallel_state.py` 移至 `device_mixin.py`, 并更改 `get_torch_distributed_backend_str` 的默认行为从 `raise NotImplementedError` 改为返回字典查找值 (`_DEVICE_TO_DISTRIBUTED_BACKEND.get(self.device_type, 'gloo')`)。子类可覆盖该方法提供自定义后端。
2. 重构 `get_default_distributed_backend`: 在 `parallel_state.py` 中, 移除本地字典定义, 导入共享字典; 修改函数: 当请求的 `device` 与 `platforms.current_platform.device_type` 相等时, 优先调用 `platforms.current_platform.get_torch_distributed_backend_str()`, 否则 fallback 到字典查找。跨设备查询 (如 CPU 辅助组) 路径不变。
3. 类型注解调整: 在 `platforms/__init__.py` 中添加 `current_platform: SRTPlatform` 类型声明, 并移除 `__getattr__` 的显式返回类型, 使 IDE 能正确推断。
4. 新增单元测试: 添加 `test_get_default_distributed_backend.py`, 通过替换 `platforms_mod.current_platform` 来模拟不同平台, 覆盖平台覆盖、跨设备不覆盖、默认平台、未知设备四种场景。测试注册在 `stage-a-test-cpu`。

关键文件:

- `test/registered/unit/distributed/test_get_default_distributed_backend.py` (模块测试覆盖; 类别 `test`; 类型 `test-coverage`; 符号 `_OverridingPlatform`, `_DefaultPlatform`, `TestGetDefaultDistributedBackend`, `setUp`): 新增测试文件, 完整覆盖 `get_default_distributed_backend` 的四种场景, 是验证重构正确性的关键。

- python/sglang/srt/platforms/device_mixin.py (模块 设备识别; 类别 source; 类型 dependency-wiring; 符号 get_torch_distributed_backend_str) : 将分布式后端字典迁移至此, 并修改 get_torch_distributed_backend_str 默认行为, 是重构的核心。
- python/sglang/srt/distributed/parallel_state.py (模块 分布式通信; 类别 source; 类型 dependency-wiring; 符号 get_default_distributed_backend) : 集中的分布式后端选择函数被重构为优先委托平台, 是变更的核心入口。
- python/sglang/srt/platforms/__init__.py (模块 平台入口; 类别 source; 类型 core-logic; 符号 getattr) : 增加类型注解, 改善了开发者体验。

关键符号: get_default_distributed_backend, get_torch_distributed_backend_str, getattr

关键源码片段

test/registered/unit/distributed/test_get_default_distributed_backend.py

新增测试文件, 完整覆盖 get_default_distributed_backend 的四种场景, 是验证重构正确性的关键。

```
"""Unit tests for sglang.srt.distributed.parallel_state — no server, no model loading."""
```

```
import unittest
```

```
import sglang.srt.platforms as platforms_mod
```

```
from sglang.srt.distributed.parallel_state import get_default_distributed_backend
```

```
from sglang.srt.platforms.interface import SRTPlatform
```

```
from sglang.test.ci.ci_register import register_cpu_ci
```

```
from sglang.test.test_utils import CustomTestCase
```

```
register_cpu_ci(est_time=2, suite="base-a-test-cpu")
```

```
# 模拟一个覆盖后端的平台子类 (out-of-tree 插件可参考此模式)
```

```
class _OverridingPlatform(SRTPlatform):
```

```
    device_type = "cuda"
```

```
    def get_torch_distributed_backend_str(self) -> str:
```

```
        return "fake_backend"
```

```
# 模拟当前所有 in-tree 平台: 不覆盖, 走 DeviceMixin 默认实现
```

```
class _DefaultPlatform(SRTPlatform):
```

```
    device_type = "cuda"
```

```
class TestGetDefaultDistributedBackend(CustomTestCase):
```

```
    def setUp(self):
```

```
        self._saved_platform = platforms_mod._current_platform
```

```
    def tearDown(self):
```

```

platforms_mod._current_platform = self._saved_platform

def _install(self, platform: SRTPlatform) -> None:
    platforms_mod._current_platform = platform

def test_overriding_platform_supplies_backend(self):
    # 当活跃平台覆盖了后端，应返回其值
    self._install(_OverridingPlatform())
    self.assertEqual(get_default_distributed_backend("cuda"), "fake_backend")

def test_overriding_platform_skipped_for_non_active_device(self):
    # 跨设备查询（如 CPU）不应受覆盖影响
    self._install(_OverridingPlatform())
    self.assertEqual(get_default_distributed_backend("cpu"), "gloo")

def test_default_platform_uses_device_mixin_table(self):
    self._install(_DefaultPlatform())
    self.assertEqual(get_default_distributed_backend("cuda"), "nccl")

def test_unknown_device_returns_gloo_default(self):
    self._install(_DefaultPlatform())
    self.assertEqual(get_default_distributed_backend("unobtainium"), "gloo")

```

python/sglang/srt/platforms/device_mixin.py

将分布式后端字典迁移至此，并修改 `get_torch_distributed_backend_str` 默认行为，是重构的核心。

```

# device_mixin.py（关键新增部分）
from sglang.srt.environ import envs

# 从 parallel_state 迁移而来的设备到分布式后端映射
_DEVICE_TO_DISTRIBUTED_BACKEND: dict[str, str] = {
    "cuda": "nccl",
    "xpu": "xccl",
    "hpu": "hccl",
    "cpu": "gloo",
    "npu": "hccl" if not envs.SGLANG_ZBAL_LOCAL_MEM_SIZE.get() > 0 else "zbal",
    "musa": "mccl",
}

class DeviceMixin:
    # ... 其他方法 ...

    def get_torch_distributed_backend_str(self) -> str:
        """Return the torch.distributed backend string (e.g. "nccl", "hccl").

        Default: lookup ``self.device_type`` in ``_DEVICE_TO_DISTRIBUTED_BACKEND``,
        falling back to ``"gloo"``. Subclasses override only when they need a
        non-default backend (e.g. mooncake, or a brand-new device).

```

```
"""
```

```
return _DEVICE_TO_DISTRIBUTED_BACKEND.get(self.device_type, "gloo")
```

python/slang/srt/distributed/parallel_state.py

集中的分布式后端选择函数被重构为优先委托平台，是变更的核心入口。

```
# parallel_state.py (关键变更片段)
from slang.srt import platforms
from slang.srt.platforms.device_mixin import _DEVICE_TO_DISTRIBUTED_BACKEND

# TODO: refactor in-tree platforms to get rid of this wrapper
def get_default_distributed_backend(device: str) -> str:
    # 通过 platforms.current_platform (而非直接导入 current_platform)
    # 使每次调用都经过 platforms 包的 __getattr__，从而能捕获 _current_platform 的
    # 运行时覆盖（例如在测试中）。
    if device == platforms.current_platform.device_type:
        # 委托给活跃平台的 get_torch_distributed_backend_str，
        # 若失败则 fallback（默认实现返回字典值）
        return platforms.current_platform.get_torch_distributed_backend_str()
    # 跨设备查询（如 CPU 辅助组）仍走字典
    return _DEVICE_TO_DISTRIBUTED_BACKEND.get(device, "gloo")
```

评论区精华

- alexnails: 指出函数中的局部导入 (from slang.srt.platforms import current_platform) 应改为全局导入，并建议在平台插件内部清理判断逻辑。最终代码改为 from slang.srt import platforms 后使用 platforms.current_platform，并精简了分支。
- ch-wan: 质疑在 `__init__.py` 中同时定义 `current_platform` 变量和 `_current_platform` 全局变量的做法，建议引入 `get_current_platform` 函数。AgainstEntropy 回应已有离线讨论，该单例模式可能在未来重构。
 - 函数内导入 vs 全局导入 (design): 作者将导入方式改为 from slang.srt import platforms 后使用 platforms.current_platform，并简化了分支判断，获得批准。
 - current_platform 变量与全局变量的设计 (design): AgainstEntropy 回复已在离线讨论，该单例模式可能在未来重构，目前保留现有方式。

风险与影响

- 风险：
 - 循环导入风险：通过 from slang.srt import platforms 而非直接导入 current_platform 来避免循环导入。
 - 平台异常处理：若平台 get_torch_distributed_backend_str 抛出 NotImplementedError 的异常，函数仅记录警告并 fallback 到字典，可能掩盖配置错误。
 - 跨设备路径不变：非活跃设备的查询仍通过字典，与之前一致，风险低。
 - 测试覆盖完整：新增测试覆盖所有分支，降低回归风险。
- 影响：

- 对用户：out-of-tree 平台开发者可直接覆盖 `get_torch_distributed_backend_str` 来自定义分布式后端，无需修补内部字典；现有 in-tree 平台因默认实现未变无感知。
- 对系统：统一了分布式 backend 选择入口，为后续平台插件化基础设施奠定基础。
- 对团队：减少硬编码映射维护成本，代码更加模块化。
- 风险标记：循环导入风险，平台异常处理，跨设备路径不变

关联脉络

- PR #21388 Unknown: 当前 PR 作为 #21388 的后续，对齐分布式后端选择与平台接口。