

PR #52016 完整报告

vllm-project/vllm

[Kernel] Add B12X dense linear backends

合并时间: 2026-08-14 10:48

原文链接: <http://prhub.com.cn/vllm-project/vllm/pull/52016>

执行摘要

- 一句话: 为 Blackwell SM12x 引入可选的 B12X 量化线性后端
- 推荐动作: 值得精读。该 PR 展示了如何把外部 kernel 库以最小侵入接入 vLLM 的线性层选择框架: `is_supported` / `can_implement` / `process_weights_after_loading` / `apply_*` 四段生命周期的职责切分、可用性统一在 `is_supported` 判定而 `apply` 路径只断言、以及 vLLM 刻意不复制 B12X 内部 policy 启发式的边界意识。review 中 mgoin 关于不做 backend 特判、恢复通用 fallback、按优先级列表正常排序的意见, 是通用框架维护的典型案列, 值得在后续集成类 PR 中复用这些评审标准。

功能与动机

PR body 说明目标是为 NVIDIA SM120/SM121 GPU 通过既有 vLLM linear backend 接口集成 B12X dense linear kernels。B12X 是纯 Python CuTe DSL 包, 固定 `b12x==1.2.4`, 不需要额外的 vLLM 构建步骤, 适合作为低集成成本的原生量化 GEMM 路径。其价值在于为 Blackwell 12x 用户提供 per-tensor FP8、block FP8、MXFP8、NVFP4、MXFP4 五种格式的额外性能选择; 同时本 PR 是 #51696 的 linear 与共享集成组件拆分, 明确 supersede 而非重复 #51696, 并指出 #41243、#47577 只覆盖 FlashInfer-embedded B12X 路径, 未提供独立可选后端。

实现拆解

变更入口是依赖与配置: `setup.py` 新增 `vllm[b12x]` extra (固定 `b12x==1.2.4`), `vllm/config/kernel.py` 为 `--linear-backend` 注册 `b12x` 取值。实现分以下步骤展开:

1. 共享懒加载与存储复用层: 新增 `vllm/utils/b12x.py`, 提供 `has_b12x`、`_get_submodule` 与 `get_b12x_blockscaled`、`get_b12x_intrinsics`、`get_b12x_mxfp8_linear`、`get_b12x_tensor_fp8_linear` 等懒加载 accessor, 避免在非 SM12x 平台无谓导入外部包; `reuse_packed_weight_storage` 递归比较 shape、stride、dtype、device 与 dataclass 字段, 允许权重重载时复用 packed 存储地址, 保障已捕获 CUDA Graph 仍有效。
2. 五条量化 kernel 实现:
 - per-tensor FP8: `scaled_mm/b12x_tensor.py` 的 `B12xTensorFP8ScaledMMLinearKernel`, 编译期将权重转置为 [K,N] 布局并 pack、合并 input/weight scale 为单一 output scale, 前向 `apply_weights` 用 `quant_fp8` 量化激活后调 `tensor_fp8.mm`, 并覆写 `apply_scaled_mm` 抛 `NotImplementedError` 防止误走

通用路径。

- block FP8: `scaled_mm/b12x_block.py` 的 `B12xFp8BlockScaledMMKernel`, `can_implement` 校验激活 `group shape` 须为 (1,128)、权重 `group shape` 须为 (128,128)、特征维 128 对齐; 128x128 UE8M0 `scale` 暂通过 `_upcast_e8m0_to_fp32` 转为 FP32 (代码内 TODO 注明待 B12X 支持后移除)。
 - MXFP8: `mxfp8/b12x.py` 的 `B12xMxfp8LinearKernel`, 采用与 tensor FP8 类似的 `pack` 模式, 原 `weight`、`weight_scale` 参数替换为空张量以释放显存。
 - NVFP4: `nvfp4/b12x.py` 的 `B12xNvFp4LinearKernel`, 复用既有 `scaled_fp4_quant` 量化激活 (`swizzled` 布局), 权重 `scale` 经 `swizzle_block_scale` 变换后走 `mm_nvfp4`。
 - MXFP4: `mxfp4/b12x.py` 的 `B12xMxFp4LinearKernel`, `can_implement` 只接受 `kMxfp4Dynamic` 动态激活, 激活量化走 `flashinfer_mxfp4_quantize(..., backend='cute-dsl')`, 再 `mm_mxfp4`。每条路径的 `is_supported` 统一集中校验 CUDA、SM12x、包安装与运行时支持 (review 后收敛), `can_implement` 只负责格式级约束。
3. 内核注册与选择优先级: `vllm/model_executor/kernels/linear/__init__.py` 将 5 个 kernel 类注册进 `_LINEAR_BACKEND_KERNEL_MAP['b12x']`, 并按格式放入 `_POSSIBLE_*_KERNELS[PlatformEnum.CUDA]` 的合适位置 (既有优化后端之后、模拟实现之前); 测试固定相对顺序防止回归。review 后恢复了 `_resolve_backend_kernels` 的通用 fallback 行为, 删除 b12x 特判。
4. Warmup 集成: 新增 `vllm/model_executor/warmup/b12x_warmup.py` 的 `b12x_warmup` dispatcher, 在 `kernel_warmup.py` 挂载; 每个 provider 先按 (`device`, `维度`, `dtype...`) 签名去重构建 `layer_map`, 无匹配层立即返回且不导入未用模块; `token` 集合由共享的 `b12x_warmup_token_counts` 生成 (1 + 全部 CUDA Graph 捕获尺寸 + `max_num_batched_tokens`), 日志用 `logger.info_once`。
5. 测试与文档配套: 新增 `test_b12x_mxfp8_linear.py` (935 行)、`test_b12x_nvfp4_linear.py`、`test_b12x_mxfp4_linear.py`、`test_b12x_warmup.py`, 覆盖 `backend` 映射、fallback 优先级、显式选择、`can_implement` 约束、`process_weights_after_loading/apply_weights` 行为 (含非连续输入回归测试)、warmup 签名去重与未用 provider 不导入; `tests/kernels/quantization/test_block_fp8.py` 增加 `test_w8a8_block_fp8_b12x_matmul`; `docs/features/quantization/b12x.md` 说明安装、自动 / 显式选择、支持格式与 W4A16 fallback 行为。

关键文件:

- `vllm/model_executor/kernels/linear/scaled_mm/b12x_tensor.py` (模块 FP8 内核; 类别 source; 类型 data-contract; 符号 `_apply_b12x_tensor_fp8_packed_linear`, `warmup_b12x_tensor_fp8_linear`, `B12xTensorFP8ScaledMMLinearKernel`, `is_supported`): per-tensor FP8 后端的核心实现, 完整展示了 `is_supported/can_implement/process_weights_after_loading/apply_weights` 生命周期与 `packed` 权重契约, 是理解本 PR 设计的入口。
- `vllm/model_executor/kernels/linear/scaled_mm/b12x_block.py` (模块 FP8 内核; 类别 source; 类型 data-contract; 符号 `_run_b12x_fp8_block_scaled_mm`, `warmup_b12x_block_fp8_linear`, `B12xFp8BlockScaledMMKernel`, `is_supported`): block FP8 内核, `can_implement` 对 `group shape` 与 128 对齐有严格约束, 且包含

UE8M0 scale upcast 的精度处理，是格式适配细节最多的文件。

- `vllm/model_executor/kernels/linear/mxftp8/b12x.py` (模块 MXFP8 内核; 类别 source; 类型 data-contract; 符号 `_b12x_mxftp8_expected_m`, `_apply_b12x_mxftp8_packed_linear`, `warmup_b12x_mxftp8_linear`, `B12xMxftp8LinearKernel`) : MXFP8 内核, 采用与 tensor FP8 类似的 pack 模式, `process_weights_after_loading` 将原参数替换为空张量, 展示另一条 packed 契约路径。
- `vllm/model_executor/kernels/linear/nvfp4/b12x.py` (模块 NVFP4 内核; 类别 source; 类型 data-contract; 符号 `_apply_b12x_nvfp4_linear`, `warmup_b12x_nvfp4_linear`, `B12xNvFp4LinearKernel`, `is_supported`) : NVFP4 内核是 review 讨论最集中的文件: `is_supported` 集中校验、移除显式 gate、去掉 contiguous 拷贝与 cutlass 硬编码, 最能体现评审落地过程。
- `vllm/model_executor/kernels/linear/mxftp4/b12x.py` (模块 MXFP4 内核; 类别 source; 类型 data-contract; 符号 `_apply_b12x_mxftp4_linear`, `warmup_b12x_mxftp4_linear`, `B12xMxFp4LinearKernel`, `is_supported`) : MXFP4 内核, 激活量化依赖 `flashinfer_mxftp4_quantize(backend='cute-dsl')`, 是唯一跨模块依赖点到外部 `flashinfer` 工具的文件, 后续 #52265 的懒加载修复正与此相关。
- `vllm/utils/b12x.py` (模块 工具层; 类别 source; 类型 dependency-wiring; 符号 `has_b12x`, `_get_submodule`, `get_b12x_blockscaled`, `get_b12x_intrinsics`) : 全 PR 的共享基础设施: 懒加载 accessor、warmup token 集合与 packed 存储复用, 跨 5 个 kernel 文件复用, 是避免污染常规导入路径的关键设计。
- `vllm/model_executor/warmup/b12x_warmup.py` (模块 预热模块; 类别 source; 类型 data-contract; 符号 `b12x_warmup`) : B12X 专属 warmup dispatcher, 统一调度 5 个 provider, 按选中层签名去重并支持未用 provider 零成本跳过, 是 warmup 集成的主入口。
- `vllm/model_executor/kernels/linear/__init__.py` (模块 内核注册; 类别 source; 类型 core-logic; 符号 `_resolve_backend_kernels`, `_LINEAR_BACKEND_KERNEL_MAP`, `_POSSIBLE_FP8_KERNELS`, `_POSSIBLE_FP8_BLOCK_KERNELS`) : 内核注册与选择优先级的关键修改点: 把 5 个 B12X kernel 挂进 backend map 和 CUDA 优先级列表 (既有优化后端之后、模拟实现之前), 并恢复通用 fallback 逻辑。
- `vllm/model_executor/warmup/kernel_warmup.py` (模块 预热挂载; 类别 source; 类型 core-logic; 符号 `warmup_kernels`) : 把 `b12x_warmup` 挂进通用 kernel warmup 流程, 决定引擎启动时是否调用 B12X 预热。
- `setup.py` (模块 依赖配置; 类别 source; 类型 configuration) : 新增 `vllm[b12x] extra` 并固定 `b12x==1.2.4`, 是本 PR 可选依赖的安装入口。
- `vllm/config/kernel.py` (模块 内核配置; 类别 source; 类型 configuration) : 为 `--linear-backend` 参数增加 `b12x` 合法取值, 与前端配置打通。
- `tests/model_executor/kernels/test_b12x_mxftp8_linear.py` (模块 内核测试; 类别 test; 类型 test-coverage; 符号 `test_b12x_backend_maps_mxftp8_kernel`, `test_b12x_backend_maps_tensor_fp8_kernel`, `test_b12x_fp8_fallback_priority`, `test_b12x_backend_does_not_intercept_unquantized_bf16`) : 最大的测试文件 (935 行), 覆盖 backend 映射、fallback 优先级固定、显式选择、`can_implement` 约束与 warmup 去重, 是选择逻辑回归的主要防线。

- tests/model_executor/test_b12x_warmup.py (模块 预热测试; 类别 test; 类型 test-coverage; 符号 test_b12x_linear_warmup_skips_unused_provider, fail_import, test_b12x_warmup_covers_linear_serving_shapes, linear_warmup) : 验证 warmup 对未使用 provider 的零成本跳过 (不导入 b12x 模块) 与 serving shape 覆盖, 是懒加载设计的关键回归测试。
- docs/features/quantization/b12x.md (模块 使用文档; 类别 docs; 类型 documentation) : 面向用户的安装、自动 / 显式选择、支持格式与 W4A16 fallback 行为说明, 是功能落地的文档配套。

关键符号: B12xTensorFP8ScaledMMLinearKernel, B12xFp8BlockScaledMMKernel, B12xMxfp8LinearKernel, B12xNvFp4LinearKernel, B12xMxFp4LinearKernel, b12x_warmup, warmup_b12x_tensor_fp8_linear, warmup_b12x_block_fp8_linear, warmup_b12x_mxfp8_linear, warmup_b12x_nvfp4_linear, warmup_b12x_mxfp4_linear, reuse_packed_weight_storage, b12x_warmup_token_counts, has_b12x, get_b12x_blockscaled

关键源码片段

vllm/model_executor/kernels/linear/nvfp4/b12x.py

NVFP4 内核是 review 讨论最集中的文件: is_supported 集中校验、移除显式 gate、去掉 contiguous 拷贝与 cutlass 硬编码, 最能体现评审落地过程。

```
# vllm/model_executor/kernels/linear/nvfp4/b12x.py
# NVFP4 路径复用 vLLM 既有 scaled_fp4_quant 做激活量化,
# GEMM 本体由 b12x.gemm.blockscaled.mm_nvfp4 提供。
```

```
def _apply_b12x_nvfp4_linear(
    x: torch.Tensor,
    weight: torch.Tensor,
    weight_scale_storage: torch.Tensor,
    input_global_scale_inv: torch.Tensor,
    alpha: torch.Tensor,
    bias: torch.Tensor | None,
) -> torch.Tensor:
    blockscaled = _import_b12x_blockscaled()
    assert blockscaled is not None # is_supported 已保证可用

    output_size = int(weight.shape[0])
    output_shape = [*x.shape[:-1], output_size]
    # 输入保持非连续视图直接进入量化, 不再做 contiguous 拷贝
    # (review 后移除, 避免多余显存搬运成本, 并有回归测试)
    x_2d = x.reshape(-1, x.shape[-1])
    x_packed, x_scale_swizzled = scaled_fp4_quant(
        x_2d,
        input_global_scale_inv,
        is_sf_swizzled_layout=True,
    )
```

```

output = blockscaled.mm_nvfp4(
    x_packed,
    x_scale_swizzled,
    weight,
    weight_scale_storage,
    alpha,
    out_dtype=x.dtype,
)
if bias is not None:
    output = output + bias
return output.view(*output_shape)

```

```

class B12xNvFp4LinearKernel(NvFp4LinearKernel):
    '''ModelOpt NVFP4 linear 通过 B12X SM120 dense GEMM。'''

    @classmethod
    def is_supported(cls, compute_capability=None) -> tuple[bool, str | None]:
        # 包、intrinsic、设备族、运行时支持全部在此集中校验（review 要求），
        # apply 路径只断言不变量，不再做运行时 ImportError 分支
        del compute_capability
        if not current_platform.is_cuda():
            return False, 'B12X NVFP4 kernels are only available on CUDA'
        if not current_platform.is_device_capability_family(120):
            return False, 'B12X NVFP4 kernels require a Blackwell 12x device'
        blockscaled = _import_b12x_blockscaled()
        if blockscaled is None or _import_b12x_intrinsic() is None:
            return False, 'Install the B12X backend with `pip install vllm[b12x]`'
        if not blockscaled.is_supported():
            return False, 'b12x native NVFP4 GEMM is not supported'
        return True, None

    @classmethod
    def can_implement(cls, config: NvFp4LinearLayerConfig) -> tuple[bool, str | None]:
        del config
        return True, None

    def process_weights_after_loading(self, layer: torch.nn.Module) -> None:
        # scale 先做 B12X 需要的 swizzle 布局变换，再原地替换参数；
        # swizzle 后的 tensor 继承原 Parameter 的 weight_loader
        intrinsic = _import_b12x_intrinsic()
        assert intrinsic is not None
        replace_parameter(
            layer,
            'weight_scale',
            intrinsic.swizzle_block_scale(layer.weight_scale.data),
        )
        layer.b12x_nvfp4_linear = True # warmup 按此标记识别选中层

```

```

def apply_weights(
    self,
    layer: torch.nn.Module,
    x: torch.Tensor,
    bias: torch.Tensor | None = None,
) -> torch.Tensor:
    return _apply_b12x_nvfp4_linear(
        x,
        layer.weight,
        layer.weight_scale,
        layer.input_global_scale_inv,
        layer.alpha,
        bias,
    )

```

vllm/utils/b12x.py

全 PR 的共享基础设施：懒加载 accessor、warmup token 集合与 packed 存储复用，跨 5 个 kernel 文件复用，是避免污染常规导入路径的关键设计。

```

# vllm/utils/b12x.py
# b12x 是可选纯 Python CuTe DSL 包，全部 accessor 懒加载：
# 没安装时返回 None，kernel 的 is_supported 会给出失败原因。

```

```

@functools.cache
def has_b12x() -> bool:
    '''Return whether the B12X package is installed.'''
    return importlib.util.find_spec('b12x') is not None

```

```

@functools.cache
def _get_submodule(module_name: str) -> ModuleType | None:
    if not has_b12x():
        return None
    try:
        return importlib.import_module(module_name)
    except (ImportError, ModuleNotFoundError):
        return None

```

```

def b12x_warmup_token_counts(
    *,
    max_tokens: int,
    cudagraph_capture_sizes: Iterable[int] = (),
) -> tuple[int, ...]:
    # B12X 内部会把选择同一 kernel policy 的 shape 去重；vLLM 刻意
    # 不去复刻 B12X 的选择启发式，只提供完整 serving shape 集合：
    # 1（模拟 decode）、全部 cuda graph 捕获尺寸、max_num_batched_tokens
    counts = {1}
    counts.update(int(size) for size in cudagraph_capture_sizes if int(size) > 0)

```

```
if int(max_tokens) > 0:
    counts.add(int(max_tokens))
return tuple(sorted(counts))
```

```
@torch.no_grad()
def reuse_packed_weight_storage(current: Any, replacement: Any) -> Any:
    '''Reuse packed tensor addresses when a compatible weight is reloaded.'''
    # 权重重载时, packed weight 的内存地址必须保持不变,
    # 否则已捕获的 CUDA Graph 会读到过期地址; shape、stride、
    # dtype、device 一致时原地 copy, 否则退回 replacement
    if current is None or not _same_packed_layout(current, replacement):
        return replacement
    _copy_packed_tensors(current, replacement)
    return current
```

评论区精华

mgoin 的 review 聚焦框架级设计原则, lukealongso 逐条落地并在测试中固化:

- 可用性校验应集中在 is_supported: mgoin 在 nvfp4/b12x.py 指出 "This should be resolved during the is_supported stage" 与 "ditto, check during is_supported"; lukealongso 将 package、intrinsic、device-family、runtime 支持全部收敛到 is_supported, apply 路径改为断言不变量。
- 删除 NVFP4 显式启用 gate: mgoin 质疑 "Why does it need to be explicitly enabled? ... I think we should just let it run if for some reason it gets triggered as a fallback"; lukealongso 移除 gate, 让 B12X 以 fallback 位置参与自动选择。
- 不要为单后端污染通用逻辑: mgoin 问 _resolve_backend_kernels 里的 "b12x special case"; lukealongso 承认 "You are right; there should not be a B12X-specific case here" 并恢复 upstream 行为。优先级列表位置同样从 "最前" 改为 "properly placed in the priority list"。
- 硬编码量化后端问题: mgoin 发现 NVFP4 的 scaled_fp4_quant 被硬编码 cutlass, 并指出 FlashInfer 路径里已有 b12x 量化, lukealongso 自嘲 "lol, the backend ouroboros", 改为让 scaled_fp4_quant 自行解析后端。
- warmup 覆盖策略: mgoin 确认保守覆盖可接受, 但指出 max_cudagraph_capture_size 与 max_num_batched_tokens 之间的 eager shape 无编译覆盖; lukealongso 保留该策略并说明 B12X 内部按 kernel policy 去重, vLLM 不应复制其启发式, 并在共享 helper 旁补充 rationale。
- warmup 缺 MXFP4/NVFP4: mgoin 问 "why not mxfp4 or nvfp4?"; lukealongso 补全 5 个 provider 并抽共享 dispatcher。
- 合并后 AMD CI 失败事件: AndreasKaratzas 质疑 "why was this force merged? There is a clear AMD failure here related to this PR."; mgoin 道歉 "I didn't look carefully enough at the AMD failure. The end looked like hardware procurement issues", 随后回归被快速修复。

- NVFP4 是否应显式启用 (design): lukealonso 移除显式 gate, B12X 以 fallback 位置参与 NVFP4 自动选择; --linear-backend b12x 保留作显式选择。
- 可用性校验应集中在 is_supported 阶段 (design): lukealonso 将 package、intrinsics、device-family、runtime 支持全部收敛到 is_supported, apply 路径改为 assert 不变量。
- scaled_fp4_quant 硬编码 cutlass (design): lukealonso 移除硬编码 backend 参数, 让 scaled_fp4_quant 自行解析, 并加回归测试确认无 backend 覆盖传入。
- NVFP4 输入是否需要 contiguous (correctness): lukealonso 移除 contiguous 拷贝, 并新增非连续输入回归测试验证同一存储直接到达 scaled_fp4_quant。
- warmup 是否覆盖每个 CUDA Graph 捕获尺寸 (performance): lukealonso 保留保守覆盖: B12X 内部按 kernel policy 去重, vLLM 不应复制其启发式; 在共享 helper 旁补充 rationale。
- warmup 缺少 MXFP4/NVFP4 provider (design): lukealonso 补全 5 个 provider, 抽共享 b12x_warmup_token_counts, MXFP4/NVFP4 获得与 FP8 相同的 gating、去重与 shape 覆盖。
- B12X 在优先级列表中的位置 (design): lukealonso 改为既有优化后端之后、模拟实现之前, 并用测试固定 FP8/block FP8/MXFP8/NVFP4/MXFP4 的相对顺序。
- _resolve_backend_kernels 的 b12x 特判 (design): lukealonso 承认不应有 B12X 特判, 恢复 upstream 通用 fallback 行为: 显式 backend 过滤支持族, 无 kernel 的族保留正常选择。
- 合并时 AMD CI 失败仍被 force merge (other): mgoin 承认未仔细查看 AMD 失败, 后续通过快速修复解决回归; 事件暴露了多平台 CI 信号把关不足。

风险与影响

- 风险: 本 PR 有实证与非实证两类风险:
- 非 CUDA 平台回归 (实证风险): 合并时 AMD CI 已有失败仍被 force merge (mgoin 事后承认未细看), 说明新增测试或导入路径在非 NVIDIA 平台存在回归窗口; 历史 PR #52265 随后修复 `vllm/utils/flashinfer.py` 的 `flashinfer_mxfp4_quantize` 懒加载与 b12x UT, 直接关联本 PR 的 `mxfp4/b12x.py` 与新增测试, 印证该风险。
- 自动选择默认路径变更: 安装 b12x 的 SM12x 用户在 auto 模式下会自动落到 B12X (排在 CUTLASS/FlashInfer 等之后), 基准只覆盖 block FP8 单请求 decode, 其他格式与 batch 形状的性能表现未充分验证。
- 数值精度变化: b12x_block.py 的 process_weights_after_loading 将 UE8M0 scale upcast 为 FP32 (带 TODO), 与 CUTLASS 结果可能有微小数值差异; b12x_tensor.py 与 mxfp8/b12x.py 把原 weight、weight_scale 参数替换为空张量, packed 权重存于 layer 属性, 任何绕过该属性的序列化、检查点或重载路径都会破坏运行。
- 启动时间与首请求延迟: warmup 覆盖全部 CUDA Graph 捕获尺寸加 max_num_batched_tokens, 中间 eager shape 的 JIT 编译在首个请求时发生, 作者说明 cutedsl 编译有磁盘缓存, 但冷启动场景仍需关注。
- 依赖 pinning: b12x==1.2.4 固定版本, 外部纯 Python DSL 包快速演进可能带来兼容压力。

- 影响：用户影响：RTX PRO 6000 Blackwell 及 SM120/SM121 设备用户可通过 `pip install vllm[b12x]` 获得新后端，block FP8 单请求解码吞吐提升约 4.33%；非 SM12x 或未安装 b12x 的用户完全无感（`is_supported` 返回失败原因）。系统影响：线性层 kernel 选择框架新增一个可选后端与注册表条目，warmup 流程新增一个分支，默认行为不变；B12X 未实现的格式（如 W4A16）保留原有 Marlin 等 fallback，混合格式模型可继续工作。团队影响：需要维护与外部 CuTe DSL 包的集成契约；这是由 AI（OpenAI Codex）辅助开发、作者声明逐行审查的贡献，review 重心应放在生命周期契约、异常路径与多平台测试覆盖。
- 风险标记：非 CUDA 平台回归，可选依赖引入，自动选择默认路径变更，中间 shape 首请求编译延迟，e8m0 scale 数值变化，权重契约变更

关联脉络

- PR #51696 B12X integration（本 PR 从中拆分的母 PR，被 supersede）：PR body 明确说明：This is the linear and shared-integration component split from #51696, which it supersedes rather than duplicates。本 PR 独立承载 linear 后端与共享集成逻辑。
- PR #41243 FlashInfer-embedded B12X path（PR body 提及）：PR body 指出 #41243 与 #47577 目标是 FlashInfer-embedded B12X 路径或更窄的集成，不提供本 PR 的 standalone 可选线性后端。
- PR #47577 Narrower B12X integration（PR body 提及）：同 #41243，属于 B12X 功能线的并行或前置探索，与本 PR 无代码重叠但功能脉络相关。
- PR #52265 [UT][XPU] fix b12x UT: 修复 flashinfer_mxfp4_quantize 懒加载与 b12x UT，涉及本 PR 新增的 `tests/model_executor/kernels/test_b12x_mxfp8_linear.py` 与 `mxfp4` 路径依赖的 flashinfer 工具，回应了合并时 AMD 平台测试失败。