

Application Note

在边缘人工智能应用中使用 MSPM0G5187



摘要

本应用手册介绍了在 MSPM0G5187 微控制器上实现边缘人工智能 (AI) 应用的方法。该微控制器集成了 TinyEngine™ 神经处理单元 (NPU)，这是一种专为资源受限嵌入式系统中高效神经网络推理而设计的专用硬件加速器。本文档首先概述 MSPM0G5187 硬件架构和 NPU 的关键特性，随后全面介绍 TI 的边缘 AI 软件生态系统，包括 TI Edge AI Studio、Tiny ML Tensorlab 和神经网络编译器 (NNC)。接着，通过两个代表性应用示例展示 MSPM0G5187 的边缘 AI 功能：基于 LeNet-5 CNN 模型的手写数字识别设计，以及基于时间序列 CNN 模型的波形分类器设计。文中提供了基于 NPU 与基于 CPU 推理实现的量化性能基准对比，充分展示了 NPU 硬件加速在推理延迟和能效方面的显著优势。这些结果为在嵌入式平台上进行边缘 AI 设计的工程师提供了实用指导和设计参考。

内容

1 简介.....	2
2 带 TinyEngine NPU 的 MSPM0G5187.....	3
3 边缘 AI 工具链.....	4
3.1 TI Edge AI Studio.....	4
3.2 TI Tiny ML Tensorlab.....	5
3.3 TI 神经网络编译器.....	5
4 边缘 AI 应用领域：数字识别.....	6
4.1 LeNet-5 变体 CNN 模型.....	6
4.2 NPU/CPU 性能比较.....	7
5 边缘 AI 应用领域：波形分类器.....	8
5.1 特征提取.....	9
5.2 时序分类模型.....	11
5.3 模型内存注意事项.....	12
5.4 NPU/CPU 性能比较.....	13
6 总结.....	14
7 参考资料.....	14

商标

所有商标均为其各自所有者的财产。

1 简介

边缘人工智能 (AI) 是指在设备本地、数据源附近或数据源处执行 AI 推理模型的技术。与云端 AI 相比，边缘 AI 具有显著优势，包括：通过降低延迟提升响应速度；通过减少带宽需求和降低云依赖提高运营效率；通过网络中断期间仍能持续运行增强系统可靠性；以及通过最小化数据传输和强化隐私保护来加强数据安全性。

MSPM0G5187 搭配 TinyEngine NPU，已助力将 AI 功能直接集成到消费电子、智能家居器件、工业及汽车系统等多个领域的终端器件中。专用的硬件 NPU 并行运行，使得 MCU 的推理延迟比无硬件加速的同类 MCU 最高可降低 90 倍，每次推理的能耗最高可降低 120 倍。

然而，边缘 AI 不仅关乎硬件，它需要一个完整的生态系统，才能真正加速用户的开发与部署。从通用分类、回归到预测任务，TI 全面的软件生态系统使工程师能够以极低的复杂度快速进行 AI 解决方案的原型设计、优化和部署。

- **TI Edge AI Studio**：提供直观的图形化环境，支持模型开发、训练、编译和部署
- **TI Tiny ML Tensorlab**：通过命令行工具提供覆盖模型训练、量化、编译和部署的端到端 workflow
- **TI Neural Network Compiler**：通过简单配置即可无缝部署至 NPU 或 CPU

MSPM0 SDK 进一步加速开发，提供丰富的即用型边缘 AI 示例项目，使开发者能够快速将参考代码部署到 MCU 上，导入预训练模型进行再训练，或将用户自定义开发的 AI 模型集成到应用中。

2 带 TinyEngine NPU 的 MSPM0G5187

MSPM0G5187 微控制器 (MCU) 属于 MSP 高度集成的超低功耗 32 位 MCU 系列，该 MCU 系列基于增强型 Arm® Cortex®-M0+ 32 位内核平台，工作频率最高可达 80MHz。这些 MCU 为需要采用小型封装或高引脚数封装（高达 64 个引脚）的最高 128KB 闪存存储器的应用同时提供了成本优化和设计灵活性。这些器件配备 USB2.0-FS 接口、数字音频接口和 TinyEngine NPU，可在整个工作温度范围内提供出色的低功耗性能。

硬件 NPU 是一款面向深度卷积神经网络 (CNN) 的高度优化内核，支持使用预先训练的模型进行机器学习推理。这与片上 CPU 结合使用，可为 CNN 推理提供更高的性能和更低的功耗。NPU 具备每秒 640 至 2560 百万次操作 (MOPS) 的性能，相比软件实现的延迟低 90 倍，且在待机模式下功耗低于 2 μ A。

图 2-1 展示了芯片内模块的顶层视图。

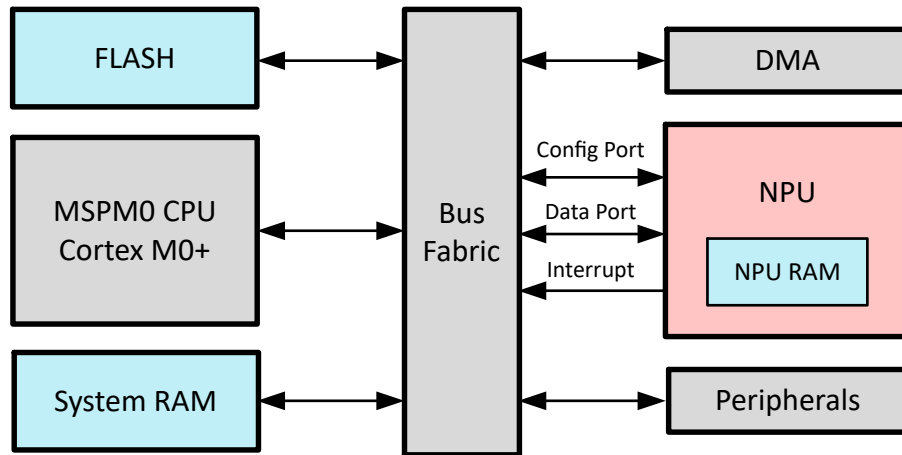


图 2-1. MSP CPU/NPU 系统结构

MSPM0 与 NPU 系统共享片上内存，包括非易失性闪存存储器和易失性 RAM 存储器。系统 RAM 支持 CPU 或 NPU 进行读写操作。此外，NPU 还拥有专用的 RAM，用于 CNN 推理。边缘 AI 模型必须适配可用闪存存储器的容量范围，且峰值 RAM 需求必须在可用系统 RAM 之内。

3 边缘 AI 工具链

3.1 TI Edge AI Studio

TI Edge AI Studio 是一套综合性的图形化及命令行工具集，旨在简化将 AI 模型开发、训练、编译和部署到德州仪器 (TI) 各类处理器、微控制器和雷达传感器上的流程。

Edge AI Studio 提供两种版本以适应不同的开发环境：云端版完全在线运行，无需任何本地环境配置，使开发者只需极少的配置即可快速入门；桌面版在用户本地机器上运行，确保所有数据和训练模型保留在本地，TI 无法访问，非常适合对数据保密性要求严格的项目。

Edge AI Studio 支持以下特性：

- 数据采集、注释和可视化
- 使用用户提供的数据集培训 TI 提供的模型（自带数据，BYOD）
- 具有自定义训练参数的训练模型
- 编译经过训练的模型
- 实时推理预览

Edge AI Studio 目前不支持以下特性：

- 训练用户定义的模型（训练自己的模型，TYOM）
- 修改 TI 模型架构
- 自定义特性提取流程或量化流程
- 直接编译预训练或客户提供的模型

3.2 TI Tiny ML Tensorlab

Tiny ML Tensorlab 是德州仪器 (TI) 为将 AI 引入微控制器而提供的完整设计方案，使用户能够：

- 训练用于时间序列分类和图像分类任务的机器学习模型
- 使用量化 (2 位、4 位、8 位) 优化模型，用于嵌入式部署
- 将模型编译为在 TI MCU 上高效运行，可选 NPU 加速
- 使用 Code Composer Studio (CCS) 部署模型

这支持以下机器学习 (ML) 任务类型：

表 3-1. 支持的任务类型

任务类型	说明
时序分类	将时间序列数据分类为离散类别 (例如，故障检测、活动识别)
时间序列回归	根据时间序列输入预测连续值 (例如，扭矩估算)
时间序列预测	根据历史模式预测未来数值 (例如，温度预测)
异常检测	使用基于自编码器的模型识别异常模式 (例如，设备监控)
图像分类	将图像分类到不同的类别 (例如，视觉检测、数字识别)

其他资源

- [Tiny ML Tensorlab GitHub 存储库](#)
- [Tiny ML Tensorlab 用户指南](#)

3.3 TI 神经网络编译器

TI 神经网络编译器 (NNC) 是 TI 为 MCU 器件开发的嵌入式 AI 编译工具链，基于 TVM 编译框架构建。TI NNC 通过简单的编译配置，可灵活地将 AI 模型部署到硬件 NPU 加速器或主机 CPU 上，无需对应用程序代码或模型端进行任何修改。推理功能始终通过统一接口调用，使开发者无需深入了解底层硬件架构，即可将 AI 推理集成到嵌入式应用中。

表 3-2. NPU/CPU 部署配置比较

目标部署	配置符号	初始化需求	调用推理
主机 CPU	TVMGEN_DEFAULT_TI_NPU_SOFT	是的，在推理之前初始化 NPU	tvmgen_default_run()
硬件 NPU	TVMGEN_DEFAULT_TI_NPU	否	

如需更多信息，请参阅[适用于 MCU 的 TI 神经网络编译器用户指南](#)。

4 边缘 AI 应用领域：数字识别

数字识别技术是一项基础功能，在嵌入式和个人消费市场中具有广泛的应用。然而，在嵌入式系统上部署可靠的识别功能面临重大挑战：传统模板匹配方法难以应对人类手写笔迹的自然变异性。边缘 AI 的出现提供了一种颇具吸引力的设计方案，能够直接在设备本身上实现准确、实时的推理，无需依赖云端连接。基于 CPU 的推理流水线需要大量的计算资源，超出了资源受限微控制器通常可提供的处理能力。

MSPM0G5187 微控制器通过集成专用的 TinyEngine NPU 来解决这一挑战，该 NPU 旨在以极低的功耗和内存开销加速机器学习工作负载。该设计演示了将单个数字字符图像分类为 10 个类别 (0-9) 之一的功能。应用通过 UART 反向通道接收 28×28 灰度图像作为像素数据，使用片上 NPU 进行硬件加速的模型推理，并将识别出的数字类别发送回主机 GUI 进行显示。

借助片上 NPU，该数字识别系统实现了约 99% 的分类准确率，推理延迟仅为 6.05ms，同时仅占用 73KB 的闪存和 10.9KB 的 RAM，如表 4-1 所示。这些结果证明了在资源受限的微控制器上进行器件端机器学习的可行性。

表 4-1. 数字识别边缘 AI 设计性能

Metric	值
精度	约 99%
闪存使用情况	73KB
RAM 使用情况	10.9KB
推理延迟 (NPU)	6.05ms
推理功耗 (AVG)	424.65uJ

表 4-2 显示了模式主要信息。

表 4-2. 数字识别 AI 模型信息

属性	值
模型架构	CNN
参数数量	60,000
输入形状	(1、28、28)
输出类	10
量化	INT8

4.1 LeNet-5 变体 CNN 模型

LeNet 是一系列开创性的卷积神经网络 (CNN) 架构，最初设计用于从小尺寸的灰度图像中识别手写数字和字符。数字识别边缘 AI 模型是通过 TI 神经网络编译器编译为 NPU 原生指令的 LeNet-5 变体 CNN 模型。

图 4-1 展示了其工作流程。该解决方案接受一个扁平化的 28×28 灰度图像作为 784 个元素的 float32 输入，并输出一个包含 10 个元素的 float32 值，表示每个类别的置信度分数。该网络由两个连续的卷积块组成，每个块包含一个带偏置的 Conv2D 层、ReLU 激活函数和 2×2 最大池化。这些模块在每个阶段逐步提取空间特性，同时将空间分辨率减半。生成的特征图随后被展平，并通过一系列全连接层，将高维特性映射到 10 类输出。由于所有层均使用 INT8 量化权重和激活操作来运行，在最后一个全连接层之后会执行去量化后处理步骤，将输出重新缩放为 float32 类型。预测的数字由主机端对 10 个输出分数取 argmax 来确定。

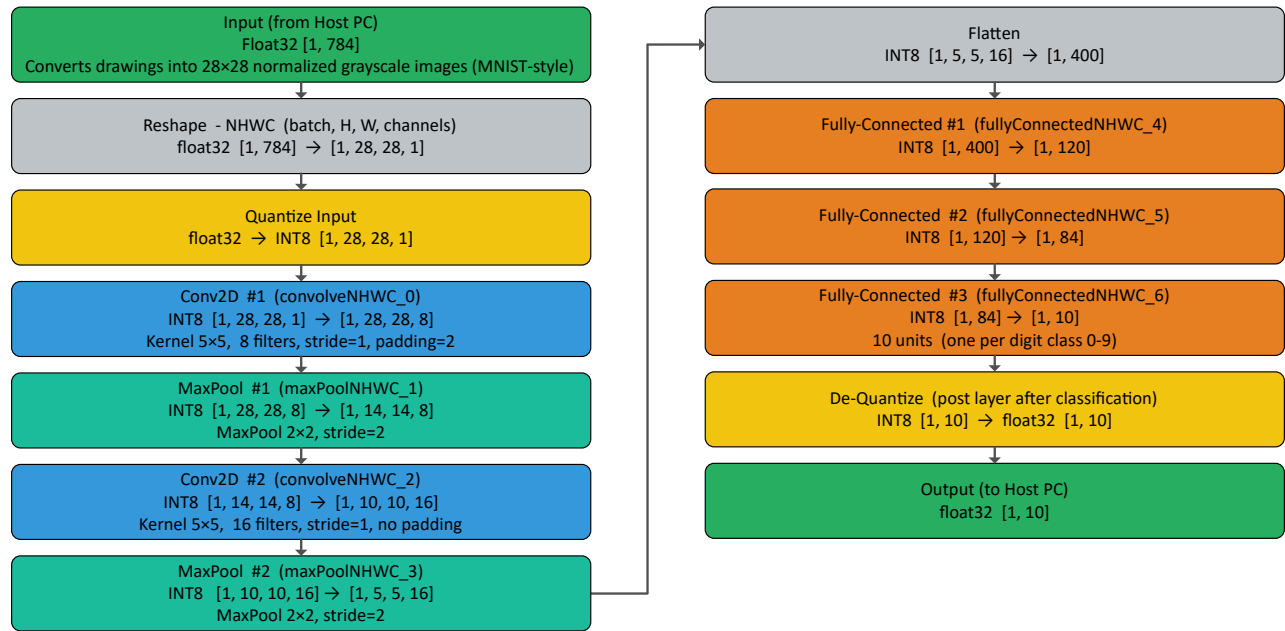


图 4-1. 数字识别 AI 模型流程图

4.2 NPU/CPU 性能比较

边缘 AI 模型可通过 TI 神经网络编译器部署到硬件上，目标可以是专用的硬件 NPU，也可以是主机 CPU。TinyEngine NPU 是一种专门针对神经网络计算进行优化的硬件加速器，能够在保证高效运算的同时，相比通用 CPU 大幅减少推理延迟并降低功耗。

为便于性能评估，MSPM0 SDK 同时提供了基于 NPU 和基于 CPU 的实现示例，供用户进行基准测试，参见：

- 基于 NPU 的数字识别边缘 AI 示例
- 基于 CPU 的数字识别边缘 AI 示例

表 4-3 总结了基于 NPU 和 CPU 的边缘 AI 设计在数字识别应用中的性能对比。

表 4-3. NPU/CPU 性能比较

性能指标	基于 NPU 的设计	基于 CPU 的设计
精度	~99%	~99%
闪存使用情况	73KB	68KB
RAM 使用情况	10.9KB	8.1KB
推理延迟	6.05ms	89.81ms
推理功耗 (AVG)	424.65 uJ	6,265.32 uJ

基于 NPU 的设计相比基于 CPU 的实现，推理延迟大约降低 14 倍，且每次推理的能耗有效减少约 93%，使其成为对功耗敏感的边缘应用场景中极具效率的选择。

5 边缘 AI 应用领域：波形分类器

MSPM0 提供了一种实验性波形分类器设计，能够对输入信号进行实时、连续的检测和分类，识别正弦波、三角波和锯齿波等常见波形类型。该设计基于一个 13K 参数的时间序列分类模型，具有较强的可扩展性和可移植性。

该波形分类器系统实现了 100% (R 平方) 的分类准确率，推理延迟仅为 5.84ms，同时仅占用 21.5KB 的闪存和 3.3KB 的 RAM，如表 5-1 所示。

表 5-1. 波形分类器边缘 AI 解决方案性能

Metric	值
精度	100%
闪存使用情况	21.5KB
RAM 使用情况	3.3KB
推理延迟 (NPU)	5.84ms
推理功耗 (AVG)	412.90uJ

表 5-2 显示了模式主要信息。

表 5-2. 波形分类器 AI 模型信息

属性	值
模型架构	CNN
参数数量	14,124
输入形状	张量 [(1, 1, 128, 1)]
输出类	3 (锯齿波、正弦波、方形波)
量化	INT8

5.1 特征提取

特性提取是常见 EdgeAI 应用的关键组成部分，尤其是在音频和时间序列信号处理中。特性提取模块提供针对嵌入式系统的高效信号处理功能，利用 ARM 库以实现最佳性能。可用的特性提取函数包括：

- 实数快速傅里叶变换 (RFFT)：将时间域信号转换为频域表示
- 幅度调节：应用缩放将信号幅度归一化到适当范围，并防止后续处理中出现上溢/下溢
- 复幅度计算：计算 FFT 所得复数的幅度
- 直流分量去除：消除信号直流偏置，以提高特征准确性
- 频带计算：对定义频带内的频率分量进行平均，以实现降维

图 5-1 展示了特性提取处理流水线。

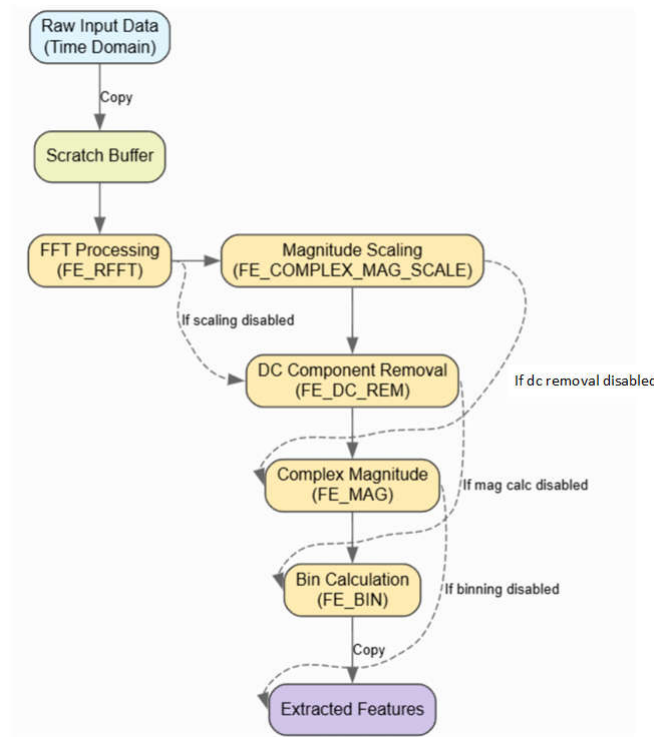


图 5-1. 特性提取处理流水线

表 5-3 展示了专为波形分类器设计的特性提取处理流水线配置。当相应的配置符号被定义时，其关联的任务将在处理流水线内执行。

表 5-3. 波形分类器的流水线配置

正在处理任务	配置符号	是否已定义
RFFT	FE_RFFT	Y
幅度缩放	FE_COMPLEX_MAG_SCALE	Y
直流分量去除	FE_DC_REM	Y
复幅度计算	FE_MAG	Y
频带计算	FE_BIN	Y

表 5-4 展示了专为波形分类器设计的特性提取配置参数。对于每个推理任务，将八个帧进行拼接，每个帧处理 256 个数据点以生成 16 个特征维度，总共产生 128 个特征维度输入到波形分类器 AI 模型中。

表 5-4. 波形分类器特性提取参数

Metric	配置符号	值
输入帧大小	FE_FRAME_SIZE	256
每帧特性大小	FE_FEATURE_SIZE_PER_FRAM	16
频率频带大小	FE_BIN_SIZE	8
归一化频带值	FE_BIN_NORMALIZE	1
要拼接的帧	FE_NUM_FRAME_CONCAT	8
输入特性大小	FE_STACKING_FRAME_WIDTH	128

备注

表 5-4 中配置的特性提取参数必须与模型训练期间使用的参数一致，以确保结果的可比性。

性能注意事项

- 经优化的计算：利用 ARM 优化的库进行高效处理
- 内存效率：缓冲区交换可最小化内存复制开销
- 降维：频带计算可降低特性维度，从而降低后续阶段的内存使用和计算量
- 参数调优：配置参数必须根据可用的内存和计算资源进行优化调整

5.2 时序分类模型

时序分类用于识别时序数据中的特定类别、运行状态或预定义标签，广泛适用于医疗健康、无线监控、智能安防和自动化测试等领域。

本波形分类器方案实现了一个轻量级、含 13K 参数的时序分类模型 (CLS_13k_NPU)，专为在嵌入式神经处理单元 (NPU) 上高效部署而设计。该设计采用 256 点实数快速傅里叶变换 (RFFT) 特性提取流水线，并结合每次 AI 推理任务中的 8 帧拼接，接受 128 元素的一维时序特性向量作为输入，输出 3 元素的输出向量，其中每个元素表示对应类别的置信度得分。

导出的模型文件从最后一个全连接层输出 INT8 (8 位量化) 张量。由于原始输出为量化形式的每类置信度得分，因此需要后处理步骤才能得到最终的分类结果。

波形分类器模型架构

模型网络结构为顺序式的特性变换流水线，其中每个阶段都直接构建在前一阶段的输出之上，逐步精炼并抽象输入表示，最终完成分类决策。图 5-2 展示了工作流程。

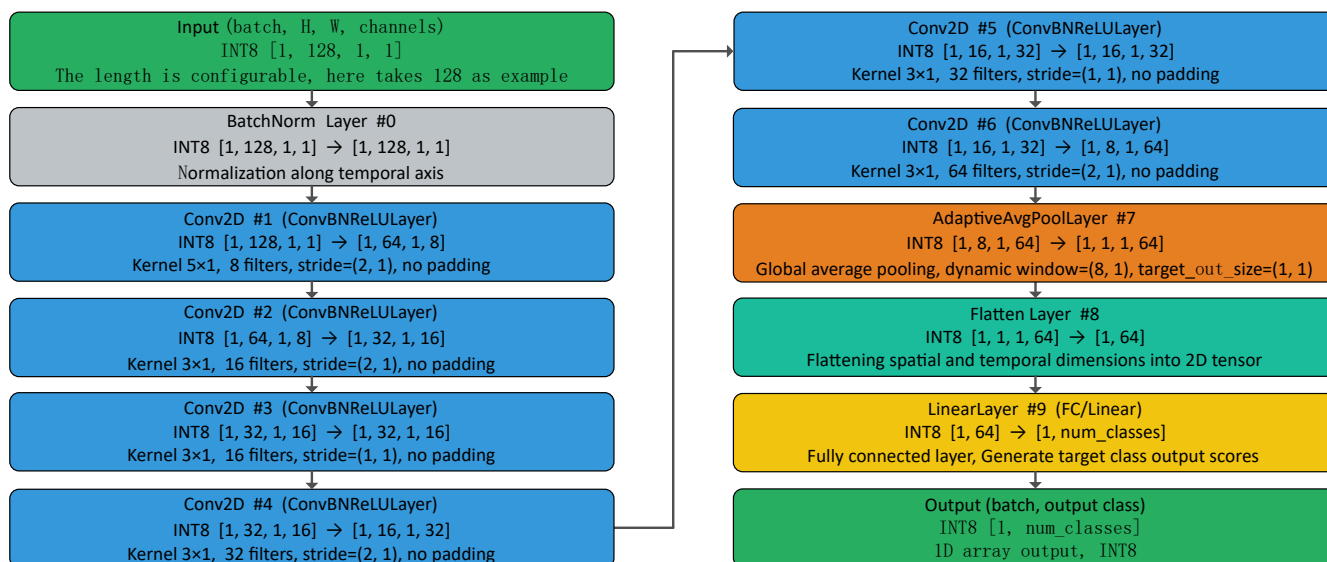


图 5-2. 时间序列分类模型 (CLS_13k) 流程图

流水线起始于**批量归一化层**，该层对原始输入序列进行实时标准化处理，确保在任何可学习的变换应用之前，输入信号幅值变化不会影响特性缩放的一致性。输入稳定后，网络通过六个**连续的卷积块**逐步加深特性表示。每个卷积块内部结构一致：一个带偏置的 **Conv2D** 层，后接**批量归一化**和 **ReLU** 激活，形成重复模式，从而在连续的阶段中捕捉越来越抽象的时序模式。每个卷积块内的**批量归一化**有助于在深度增加时维持训练稳定性，而 **ReLU** 激活则为模型学习复杂信号特性提供了必要的非线性。

经过六个卷积阶段后，所得特性图通过一个**自适应平均池化层**，将空间分辨率压缩至固定大小 **1×1**，而与之之前的时序长度无关。这一设计尤为关键：通过将模型与固定时序维度解耦，该模型本质上也支持不同长度的时序输入，为未来的部署场景提供了灵活性，无需修改网络架构。

固定尺寸的池化特性图随后被展平为二维张量，并送入最后的全连接 (FC) 层，该层将压缩后的高维特性表示直接映射到目标类别的输出得分。

模型输出和后处理

导出的模型文件从最后一个全连接层输出 INT8 (8 位量化) 张量。由于原始输出为量化形式的每类置信度得分，因此需要后处理步骤才能得到最终的分类结果。通常，这是通过在输出向量上应用 **argmax** 操作来实现的，该操作会识别出置信度分数最高的类别索引作为预测的分类结果。

5.3 模型内存注意事项

在资源受限的嵌入式平台上开发边缘 AI 解决方案，需要平衡一个严格的三维优化组合：算法性能（准确性、延迟）、非易失性存储（FLASH/ROM）和运行时内存（SRAM）。

- **算法性能**：包括推理准确性和延迟，这两者主要由模型架构决定（包括网络深度、层宽度和算子复杂度）。
- **静态储存内存 (Flash/ROM)**：主要由模型的总参数量（权重和偏置）决定。在完全量化的 INT8 流水线中，每个参数直接对应一个字节的闪存存储。因此，通道宽度扩展的更深网络会线性地增加静态二进制文件大小。有时，输入特性维度的变化也会影响静态内存使用（例如，在全连接层中）。
- **动态运行时内存 (SRAM/RAM)**：由输入特性大小以及中间层产生的激活张量（特性图）决定。当时间序列切片在网络中传播时，处理核心必须分配临时工作区来存储层的输入和输出。更长的输入时间窗口或更高的特性维度会指数级地增加峰值运行时 RAM 需求。

受这些硬件现实驱动，边缘部署迫使人们摆脱传统的“准确性优先”思维。相反，成功的部署需要进行细致的权衡，将原始模型性能直接与静态闪存和峰值动态 SRAM 的硬性物理边界相平衡。

为了说明这种关系，表 5-5 量化了波形分类器在不同模型大小和输入特性配置下的内存占用和资源利用情况。

表 5-5. 模型内存分析

模型变体 (参数)	闪存 (ROM) 大小	RAM 大小@输入= 64	RAM 大小@输入= 128	RAM 大小@输入= 256
CLS_1K	约 5.6KB	约 2.7KB	约 5.3KB	约 10.4KB
CLS_4K	约 9.9KB	约 1.2KB	约 1.7KB	约 2.7KB
CLS_13K	约 21.4KB	约 2.3KB	约 3.3KB	约 5.3KB

在所有输入维度下，最小的模型 (CLS_1K) 所消耗的运行时 RAM 始终是较大模型 (CLS_4K 和 CLS_13K) 的两到四倍。这种看似不合常理的行为源于参数数量和运行时内存是由截然不同的架构决策所驱动的。CLS_1K 采用浅层拓扑结构，缺少早期下采样，导致大幅值、高分辨率的特性图在整个中间层持续占用内存。相比之下，CLS_4K 和 CLS_13K 采用更深层的架构，并在网络前端引入带步长的卷积，能够立即缩减时间维度，从而显著减小中间运行时张量的大小。

5.4 NPU/CPU 性能比较

边缘 AI 模型可通过 TI 神经网络编译器部署到硬件上，目标可以是专用的硬件 NPU，也可以是主机 CPU。为便于性能评估，MSPM0 SDK 同时提供了基于 NPU 和基于 CPU 的实现示例，供用户进行基准测试，参见：

- [基于 NPU 的波形分类器边缘 AI 示例](#)
- [基于 CPU 的波形分类器边缘 AI 示例](#)

表 5-6 总结了波形分类器应用中基于 NPU 与基于 CPU 的边缘 AI 设计之间的性能对比。

表 5-6. NPU/CPU 性能比较

性能指标	基于 NPU 的设计	基于 CPU 的设计
精度	100%	100%
闪存使用情况	21.5KB	21.7KB
RAM 使用情况	3.3KB	3.1KB
特性提取延迟	10.75ms	10.93ms
特性提取功耗 (AVG)	752.92uJ	755.8suJ
推理延迟	5.77ms	54.42ms
推理功耗 (AVG)	412.90uJ	4559.68uJ

在硬件计算效率的对比评估中，NPU 加速架构在性能和功耗指标上均显著优于基准 CPU 实现方案。

在纯推理阶段，基于 NPU 的设计展现出卓越的加速效果，执行延迟降低约 8 倍，同时单次推理能耗降低约 91%。

若将前端特性提取纳入端到端系统分析，由于信号预处理在 CPU 上产生的确定性工作负载，整体效率提升会呈现局部衰减；即便如此，基于 NPU 的流水线仍能保持约 3 倍的延迟压缩，并将单次执行能耗降低约 78%。

6 总结

本应用手册全面介绍了基于 MSPM0G5187 微控制器的边缘 AI 解决方案，涵盖硬件特性、软件工具链以及代表性应用示例。

通过两个应用验证了该平台的功能：数字识别方案实现了约 99% 的分类准确率，且硬件 NPU 相比基于 CPU 的实现，推理延迟降低 15 倍，能效提升 93%；波形分类器实现了 100% 的分类准确率，NPU 推理延迟降低 8 倍，能效提升 91%。

这些结果证实，搭载集成式 TinyEngine NPU 的 MSPM0G5187 是面向资源受限嵌入式器件高效部署 AI 的极具吸引力的平台。

7 参考资料

1. 德州仪器 (TI), [具有 128kB 闪存、32kB SRAM、USB 2.0 FS、I2S、ADC 和 TinyEngine™ NPU 的 80MHz Arm® Cortex®-M0+ MCU](#), 产品页面。
2. 德州仪器 (TI), [边缘 AI 加速的 Arm® Cortex®-M0+ MCU 如何为电子产品注入更强智能](#), 技术文章。
3. 德州仪器 (TI), [TI 的 TinyEngine™ NPU 为更多嵌入式系统解锁边缘 AI 加速能力](#), 产品页面。
4. 德州仪器 (TI), [在 MSPM0 低成本微控制器上利用边缘 AI 进行交流电弧故障检测](#), 应用手册。
5. 德州仪器 (TI), [MSPM0 G 系列 80MHz 微控制器](#), 技术参考手册。
6. 德州仪器 (TI), [使用 TI 微控制器 \(MCU\) 进行数字识别](#), 产品页面。
7. 德州仪器 (TI), [Tiny ML Tensorlab GitHub 存储库](#), 网页。
8. 德州仪器 (TI), [Tiny ML Tensorlab 用户指南](#), 网页。
9. 德州仪器 (TI), [MSPM0 边缘 AI 用户指南](#), 资源管理器。

重要通知和免责声明

TI“按原样”提供技术和可靠性数据（包括数据表）、设计资源（包括参考设计）、应用或其他设计建议、网络工具、安全信息和其他资源，不保证没有瑕疵且不做任何明示或暗示的担保，包括但不限于对适销性、与某特定用途的适用性或不侵犯任何第三方知识产权的暗示担保。

这些资源可供使用 TI 产品进行设计的熟练开发人员使用。您将自行承担以下全部责任：(1) 针对您的应用选择合适的 TI 产品，(2) 设计、验证并测试您的应用，(3) 确保您的应用满足相应标准以及任何其他安全、安保法规或其他要求。

这些资源如有变更，恕不另行通知。TI 授权您仅可将这些资源用于研发本资源所述的 TI 产品的相关应用。严禁以其他方式对这些资源进行复制或展示。您无权使用任何其他 TI 知识产权或任何第三方知识产权。对于因您对这些资源的使用而对 TI 及其代表造成的任何索赔、损害、成本、损失和债务，您将全额赔偿，TI 对此概不负责。

TI 提供的产品受 [TI 销售条款](#)、[TI 通用质量指南](#) 或 [ti.com](#) 上其他适用条款或 TI 产品随附的其他适用条款的约束。TI 提供这些资源并不会扩展或以其他方式更改 TI 针对 TI 产品发布的适用的担保或担保免责声明。除非德州仪器 (TI) 明确将某产品指定为定制产品或客户特定产品，否则其产品均为按确定价格收入目录的标准通用器件。

TI 反对并拒绝您可能提出的任何其他或不同的条款。

版权所有 © 2026，德州仪器 (TI) 公司

最后更新日期：2025 年 10 月