

Mustang-V100-MX4



Feature

- PCIe Gen 2 x 2 form factor
- 4 x Intel® Movidius™ Myriad™ X VPU MA2485
- Power efficiency, Approximate 15W.
- Operating Temperature -20°C~60°C
- Powered by Intel's OpenVINO™ toolkit
- Multiple cards supported



Introduction

The Mustang-V100-MX4 is a PCIe Gen 2 x 2 card included 4 Intel® Movidius™ Myriad™ X VPU, providing an flexible AI inference solution for compact size and embedded systems.

VPU is short for vision processing unit. It can run AI faster, and is well suited for low power consumption applications such as surveillance, retail, transportation. With the advantage of power efficiency and high performance to dedicate DNN topologies, it is perfect to be implemented in AI edge computing device to reduce total power usage, providing longer duty time for the rechargeable edge computing equipment.

Specifications

Model Name	Mustang-V100-MX4
Main Chip	4 x Intel® Movidius™ Myriad™ X MA2485 VPU
Operating Systems	Ubuntu 18.04.x 16.04.x LTS 64bit, CentOS 7.4 64bit, Windows® 10 64bit
Dataplane Interface	PCIe Gen 2 x 2
Power Consumption	Approximate 15W
Operating Temperature	-20°C~60°C
Cooling	Active fan
Dimensions	113 x 56 x 23 mm
Operating Humidity	5% ~ 90%
Dip Switch/LED indicator	Identify card number
Support Topology	AlexNet, GoogleNetV1/V2, MobileNet SSD, MobileNetV1/V2, MTCNN, Squeezenet1.0/1.1, Tiny Yolo V1 & V2, Yolo V2, ResNet-18/50/101 * For more topologies support information please refer to Intel® OpenVINO™ Toolkit official website. [Supported Models] https://docs.openvinotoolkit.org/latest/_docs_IE_DG_Introduction.html#SupportedFW [Supported Framework Layers] https://docs.openvinotoolkit.org/latest/_docs_MO_DG_prepare_model_Supported_Frameworks_Layers.html

Ordering Information

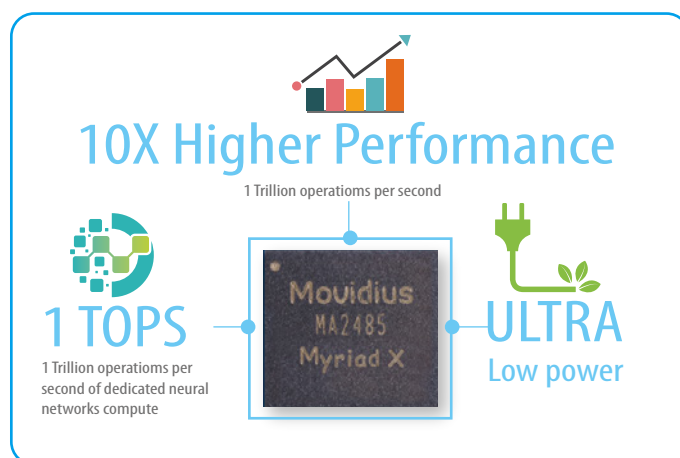
Part No.	Description
Mustang-V100-MX4-R10	Computing Accelerator Card with 4x Intel® Movidius™ Myriad™ X MA2485 VPU, PCIe Gen 2 x 2 interface, RoHS

Packing List

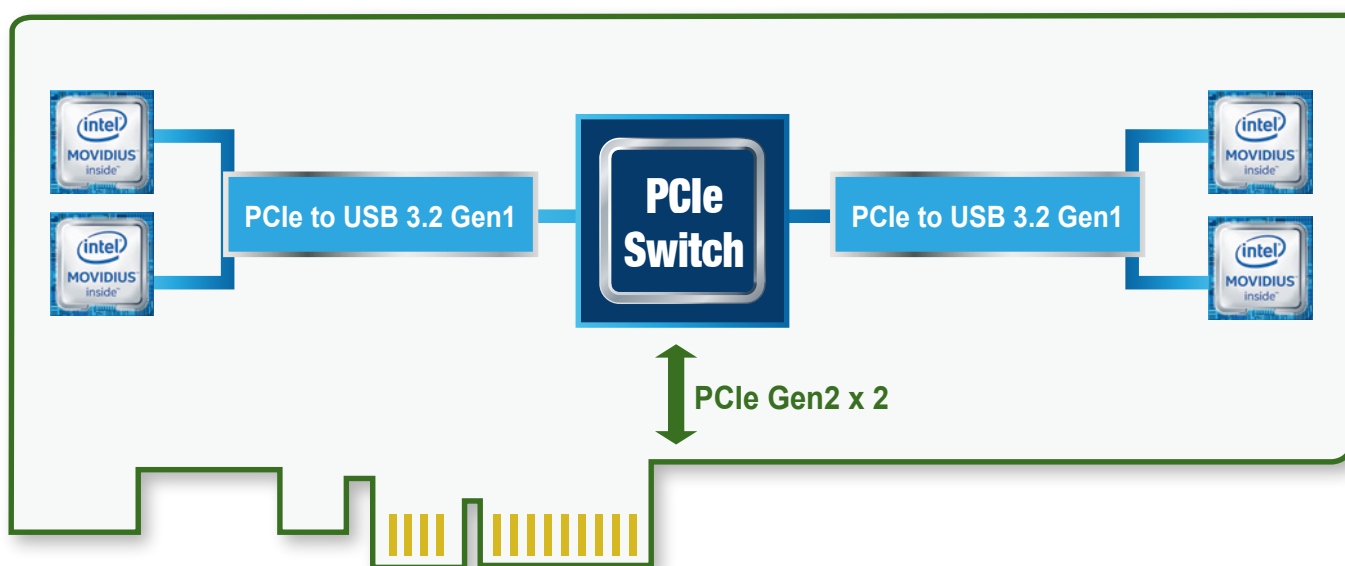
1 x Full height bracket

Key Features of Intel® Movidius™ Myriad™ X VPU:

- Native FP16 support
- Rapidly port and deploy neural networks in Caffe and Tensorflow formats
- End-to-End acceleration for many common deep neural networks
- Industry-leading Inferences/S/Watt performance



Mustang-V100-MX4 Block Diagram



Mustang-V100-MX4 Dimensions (Unit: mm)

