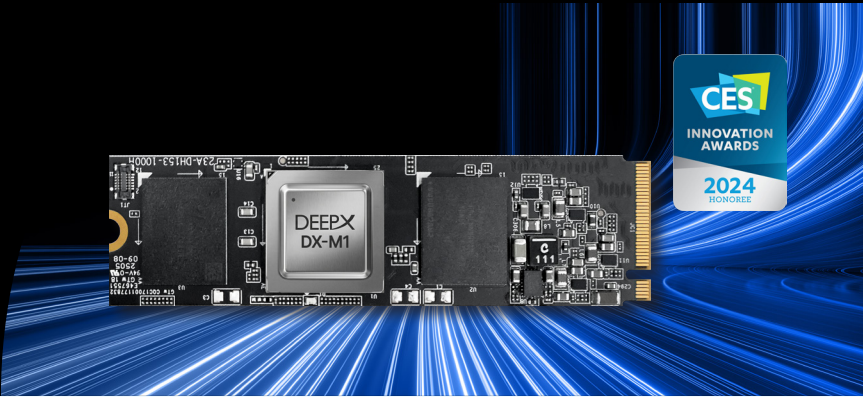


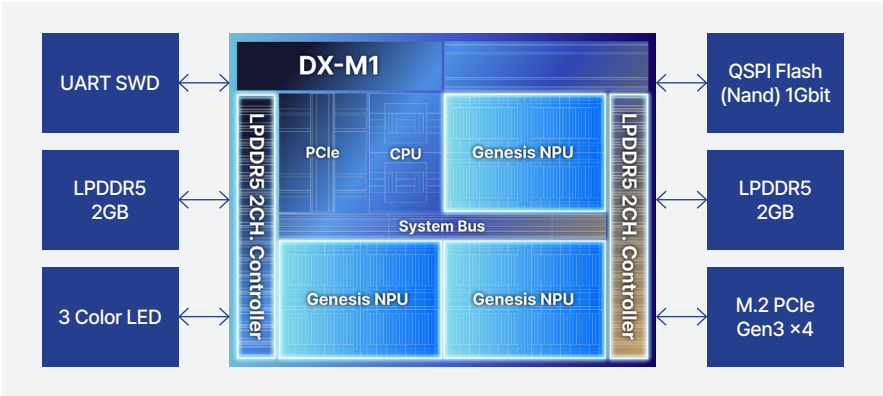
DX-M1 M.2 Module

Intelligence liberated — powerful, efficient, and universal at the edge



Type: AI Accelerator Performance: 25 TOPS / 5W (Typical)

Functional Block Diagram



Specifications

Features	Details
AI Performance	25 TOPS (INT8)
Host Interface	PCIe Gen3 x4 (Supports Gen 1/2/3 & x1/x2/x4)
Memory	4GB LPDDR5
Power Consumption	5W (Typical)
Operating Temperature	0 ~ 70°C (Commercial)
	-25 ~ 85°C (Extended)
	-40 ~ 85°C (Industrial)
Thermal Solution	Heatsink (Option)
Form Factor	M.2 2280(M Key), 22 × 80 × 4.72mm
OS Support	Windows 11/10
	Debian-based Linux (Ubuntu 24.04/22.04/20.04 LTS), Docker
	Yocto Linux
	Android
AI Frameworks	Ultralytics, TensorFlow, PyTorch, ONNX, Keras
System Support	x86, ARM Based Architecture

Overview

The DX-M1 M.2 is more than silicon, it is the declaration that intelligence belongs everywhere, not just in the cloud. At only **5W (Typical)**, it **delivers the force of 200 GPU-class TOPS**, transforming the edge from a passive receiver into an active decision-maker. at just 1-5W through breakthrough efficiency design.

Through its innovative memory design, DX-M1 allows **multiple AI models to work side by side**, not as a luxury but as a new standard. What once demanded costly, power-hungry systems is now within reach of **every device, every factory, every city**.

Target Applications

- Edge Cameras Systems
- Smart Mobility
- Smart Factory
- Smart Cities
- Robotics
- Drones
- Edge Computing
- Smart Homes
- Smart Retail

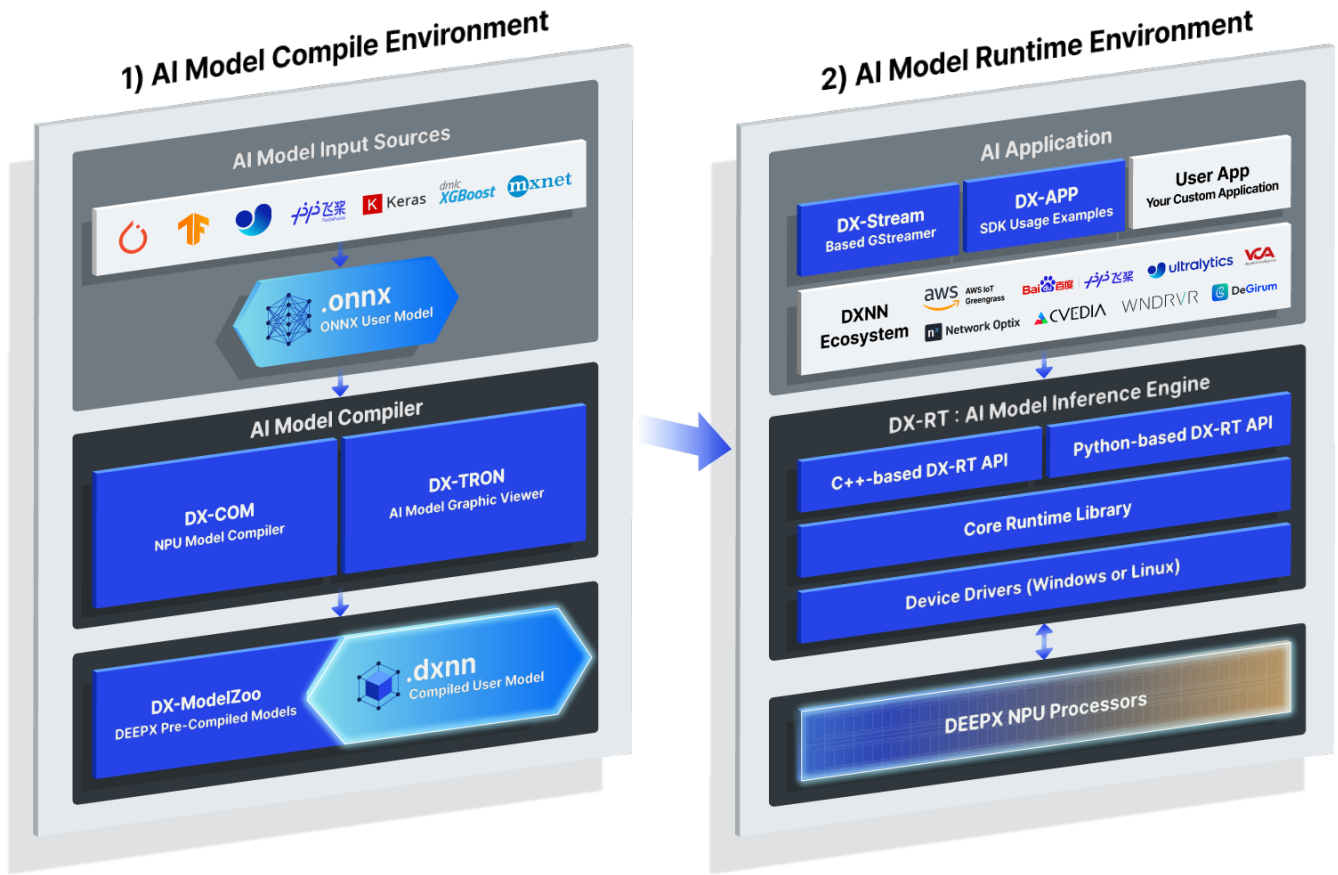


DXNN® - DEEPX SW Development Kit

DXNN® (DEEPX Neural Network) SDK streamlines the AI deployment pipeline on DEEPX NPUs. By integrating essential tools for compilation, optimization, simulation, and inference, it ensures high development efficiency. Experience a ready-to-use development environment with DX-AS (All Suite), a fully integrated and version-aligned package designed for rapid scaling.

DXNN SDK Full Stack Architecture

□ Third Party Space ■ User Space ■ DEEPX SDK Space



1) AI Model Compile Environment

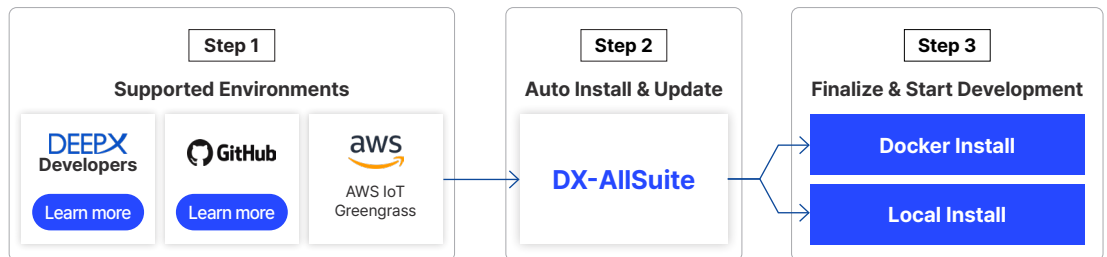
Optimizes models from various sources into DEEPX-specific .dxnn formats via the ONNX framework.

2) AI Model Runtime Environment

Deploys optimized models onto hardware via DX-Stream and DX-RT engines using C++ or Python APIs.

DX-AllSuite

DX-AllSuite: The Single Package for Your Complete DXNN Environment. Simplify setup, installation, and updates across local machines and Docker.



DEEPX HQ

3F, 20 Pangyoyeok-ro 241beon-gil,
Seongnam-si, Gyeonggi-do, South Korea

USA

1735 Technology Drive Suite
740. San Jose, CA U.S.A

China

No 0910, Yangguang Yuehai Building, Nanshan
District, Shenzhen, Guangdong, China

Taiwan

Nanjing E. Rd., Zhongshan Dist.,
Taipei City 104, Taiwan (R.O.C.)



Shop Now!
Contact Sales
sales@deepx.ai