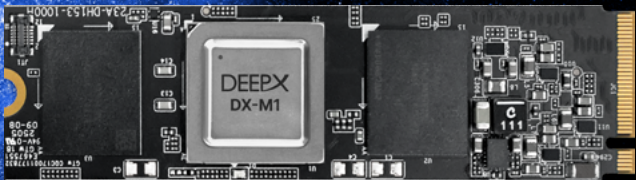


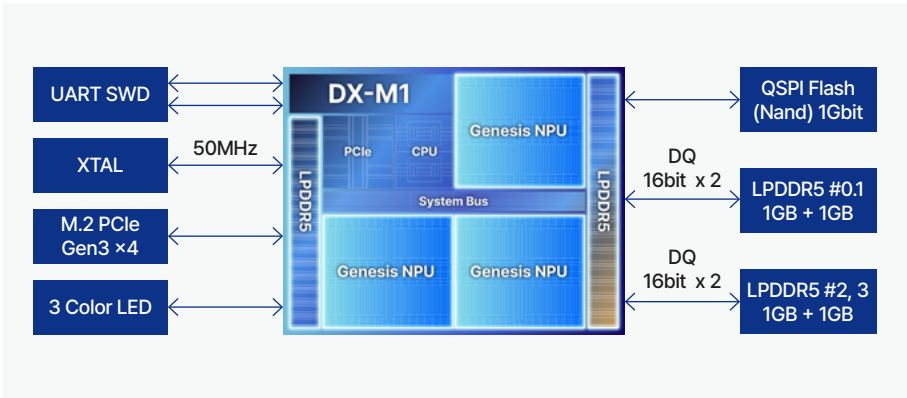
DX-M1 M.2 LPDDR5×2 for Every AIoT

The DEEPX DX-M1 M.2 module brings server-grade AI inference directly to edge devices. Delivering 25 TOPS of performance at just 2-5W, the module achieves 20x better performance efficiency (FPS/W) than GPGPUs while maintaining GPU-level AI accuracy.



“Integrating powerful AI vision processing and essential core features into a single chip, DEEPX drives innovative edge AI solutions for diverse intelligent systems.”

DX-M1 M.2 LPDDR5x2: Functional Block Diagram



Specifications

Feature	AI Accelerator	Details
Processor	INT8 Performance	25 TOPS (=200 eTOPS / INT8)
Signal Interface	PCI Express	PCIe Gen.3 x4 / Bandwidth: 4GB/s *Compatible to PCIe x1
Power	Power Consumption	2W min., 5W max. for DEEPX supported models
Operating	Temperature	-25 ~ 85°C (Throttling) -25 ~ 65°C (Non_Throttling)
Environment	Humidity	40 °C @ 85% relative humidity (non-condensing)
Thermal Solution	Cooling	Heatsink (Option)
Physical	Form Factor	M.2 2280 (Key M)
	Dimensions	22mm x 80mm x 4.1mm
	Power Range	3.3V ± 5%
Software Support	Windows	Windows 11, 10 64 bit
	Linux	Ubuntu 22.04, 20.04 LTS Support Yocto Project and Docker
	Framework	Support TensorFlow, TensorFlow Lite, ONNX, Keras, PyTorch by Dataflow compiler converted
System Support	CPU Platform	x86, ARM Based Architecture



Key Features

- > Type: AI Accelerator
- > Form Factor: M.2 M Key (22 × 80 mm)
- > Interface: PCIe Gen.3 x4
- > Memory: 4GB LPDDR5, QSPI 1Gbit NAND Flash
- > Host HW: x86, ARM Based Architecture

Support DXNN® SDK

DXNN® SDK is a comprehensive SW development environment for deploying AI on DEEPX NPUs. It integrates tools for compiling, optimizing, simulating, and inferring the latest AI models, such as YOLO, ViT, and VLMs. And it provides an optimized, ready-to-use environment as the DX-AII Suite package to support fast and efficient AI development.

Target Applications

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

DEEPX HQ
5F, 20 Pangyoyeok-ro 241beon-gil,
Seongnam-si, Gyeonggi-do, South Korea

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

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

Europe
Coming Soon



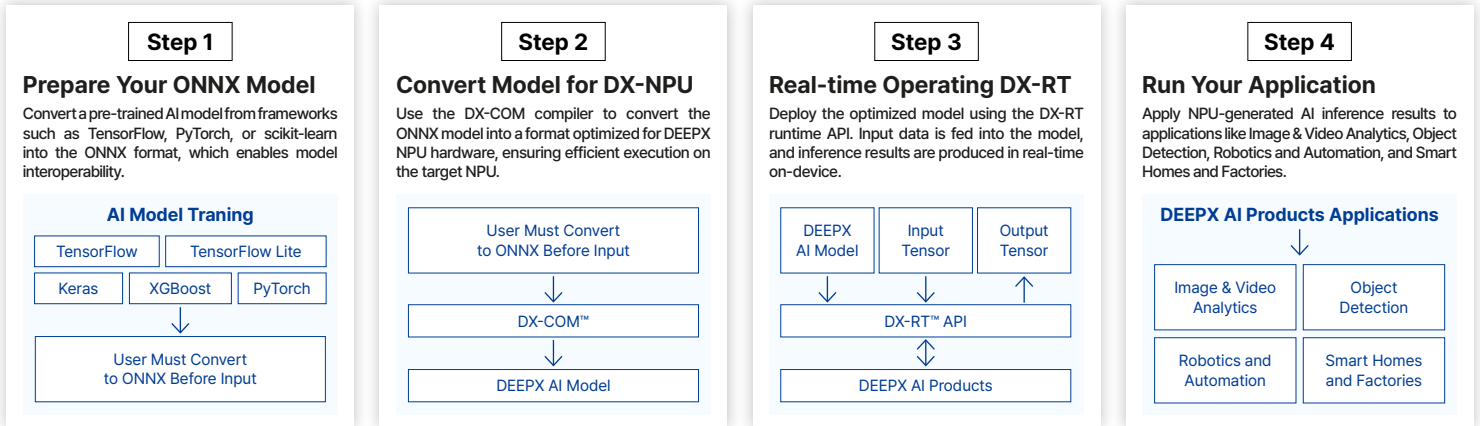
Buy Now! - [DX TechBridge Program](#)
Contact Sales - sales@deepx.ai

DXNN® SDK

DXNN® (DEEPX Neural Network) SDK streamlines AI deployment on DEEPX NPUs by integrating version-aligned tools for compilation, optimization, simulation, and inference. For efficient development, it's offered as the DX-AS (All Suite), a fully integrated and optimized package.

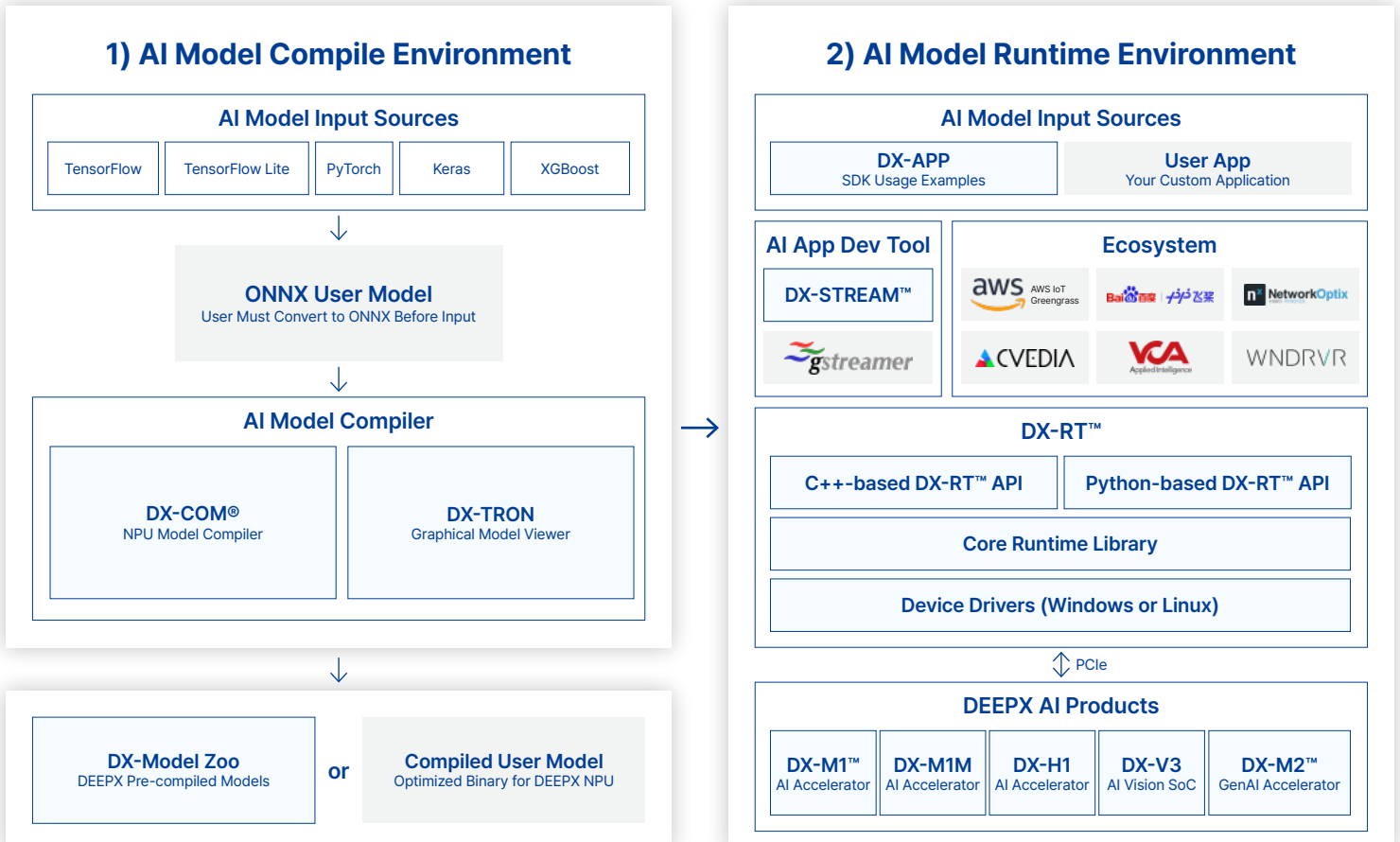


How It Works: 4-Step AI Deployment with DXNN® SDK



DXNN® Full Stack Architecture

DXNN® Full Stack Architecture streamlines AI model deployment onto DEEPX products using its two-stage AI Model Compile and Runtime Environments.



DEEPX HQ
5F, 20 Pangyoyeok-ro 241beon-gil,
Seongnam-si, Gyeonggi-do, South Korea

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

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

Europe
Coming Soon



Buy Now! - [DX TechBridge Program](#)
Contact Sales - sales@deepx.ai