

DX-AIPlayer N97

Comprehensive AI computing solutions designed for peak performance, reliability, and efficiency



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

Specifications

Category	Feature	Details
System	DX-M1 M.2 Module (AI Accelerator)	AI Perf: 25 TOPS / 5W (Typical)
		Form Factor: M.2 M Key (22 × 80 mm)
		Interface: PCIe Gen.3 × 4
		Memory: 4GB LPDDR5
	CPU	Intel® Processor N97
	Memory	RAM: 8GB LPDDR5
	Graphics	Intel® UHD Graphics for 12th Gen Intel® Processors H/W Decoder : H.264/265 8CH (1080P @ 30fps)
	Storages	eMMC Onboard 64GB
	Ethernet	Gb Ethernet (full speed) RJ-45 × 2
	WIFI/BT	Optional with M.2 2230 E-key x 1
	Expansion Slot	M.2 2230 E-Key x 1 (PCIe x1, USB2.0 ×1)
		M.2 2280 M-Key x 1 (PCIe x2)
	Security	Onboard TPM 2.0
	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
I/O Placements	USB	USB 2.0 × 2 (via 10 Pin Header x 1)
		USB 3.2 Gen 2 × 3 (Type-A)
	Display Port	HDMI 2.0b x 1
		DP 1.2 × 1
	Ethernet	Gb Ethernet (full speed) RJ-45 × 2
	COM	RS-232/422/485 × 2
Power Supply	Audio	Line out x 1, Mic in x 1
	Power Requirement	12V DC-in, 5A
	Power Supply Type	AT/ATX
Mechanical	Power Consumption	15W ~ 37W (Typical)
	Mounting	VESA Mounting
	Dimensions	95 mm x 95 mm x 55 mm
	Net Weight	450g
Environment	Gross Weight	600g
	Operating Temp.	0°C ~ 60°C / 0.5 airflow
	Storage Temp.	-20°C ~ 70°C
	MTBF	1,224,238 Hours
	Certification	CE, FCC, ROHS, UKCA, VCCI, KC

Overview

An integrated edge AI platform combining the Intel® N97 processor and DX-M1 NPU to deliver high-performance AI inference alongside low-power, low-heat reliability for throttling-free operation in demanding industrial environments.

Target Applications

- Edge Cameras Systems
- Smart Mobility
- Smart Factory
- Smart Cities
- Robotics
- Edge Cameras Systems
- Industrial Robot Arms
- Remote Patient Monitoring (RPM)
- Smart Homes
- Smart Retail



Learn more

AWS IoT Greengrass Qualification

AWS IoT Greengrass, enabling secure cloud connectivity and fleet-scale deployment. Supports local execution of components and containerized Docker applications.



Learn more

DEEPX x Ultralytics YOLO26

As a key partner, Ultralytics provides optimized support for YOLO26, enabling fast and seamless deployment of real-time AI solutions in robotics and factory automation.

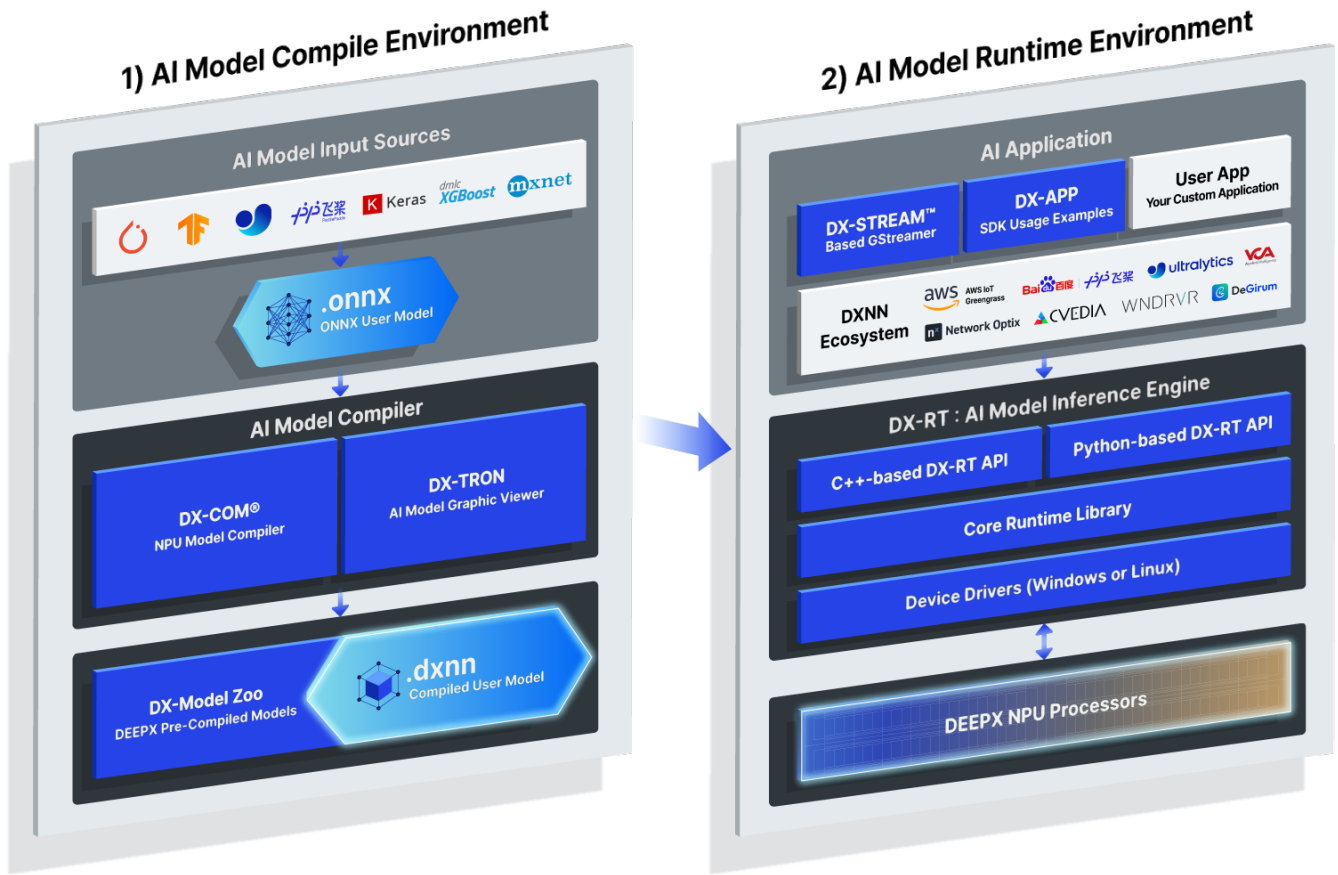


DXNN® - DEEPX SW Development Kit

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

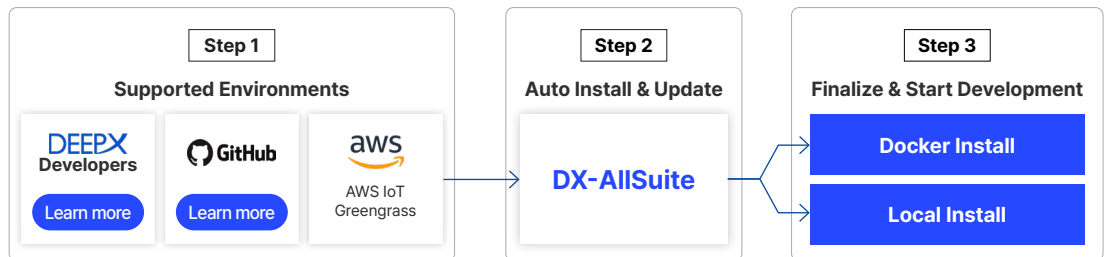
Optimizes models from various sources into DEEPX-specific .dxnn formats via the ONNX framework.

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