

AIR-075

AI Inference System Powered by NVIDIA® Jetson Thor™



Features

- Powered by NVIDIA Jetson T5000™ and Jetson T4000™, delivering up to 2070 TFLOPS FP4 inference performance.
- 19~36V wide power and -10~40 °C wide temp. supported.
- Multiple I/O Expansion: 2x USB-C 3.2 Gen 2, 2x USB-A 3.2 Gen 1, 1x M.2 B Key (5G/LTE), and 2x M.2 E Key (NVMe / Wi-Fi).
- Support 10GbE with RJ45 connectors.
- Support WEDA / SUSI / Edge AI SDK / Inference Kit.



Specifications

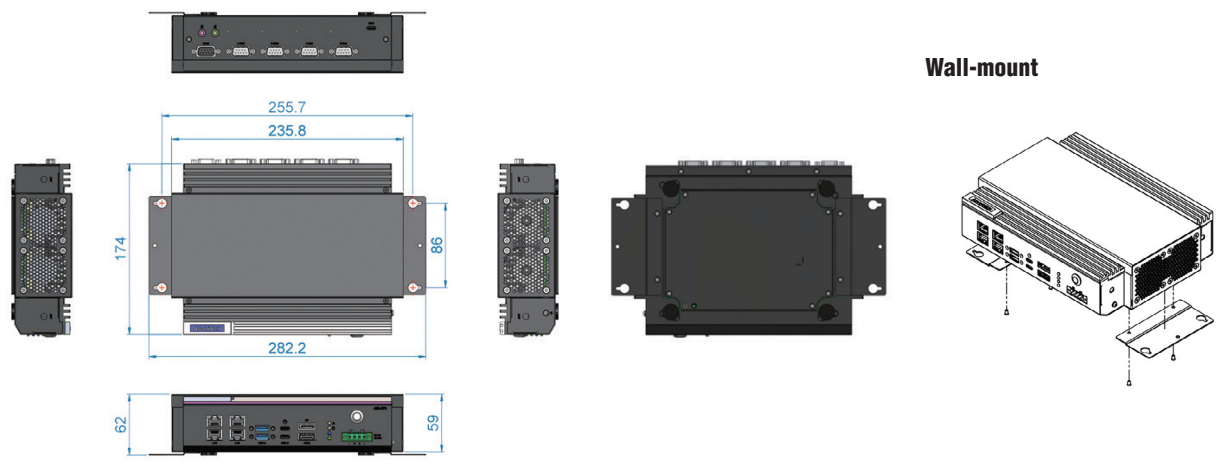
Processor Module	Module	NVIDIA Jetson T5000	NVIDIA Jetson T4000
	CPU	14-core Arm® Neoverse®-V3AE 64-bit CPU	12-core Arm® Neoverse®-V3AE 64-bit CPU
	GPU	2560-core NVIDIA Blackwell architecture GPU with 96 fifth-gen Tensor Cores	1536-core NVIDIA Blackwell architecture GPU with 64 fifth-gen Tensor Cores
	AI Performance	2070 TFLOPS (FP4—Sparse)	1200 TFLOPS (FP4—Sparse)
	Memory	128GB 256-bit LPDDR5X , 273 GB/s	64GB 256-bit LPDDR5X , 273 GB/s
Ethernet	Interface	RJ-45 optional PoE support (IEEE 802.3 af)	
	PHY	Marvell AQR113	
	Speed	4x 10GbE Ethernet (10M/100M/1G/2.5G/10Gbps)	2x 10GbE Ethernet (10M/100M/1G/2.5G/10Gbps)
Display	HDMI	1x HDMI 2.1, Max 3840 x 2160 (4K) at 60 Hz	
	DP	1x DP 1.4a, Max 3840 x 2160 (4K) at 60 Hz	
IO Ports	USB	2x USB 3.2 Gen1 Type C (USB 5 Gbps) 2x USB 3.2 Gen2 Type A (USB 10 Gbps) 1x USB 3.2 Gen1 Type C (RCM only)	
	CANBus	2x DB9	N/A
	DI/DO	1x DB9 (8 bit)	
	COM	2x DB9 (RS-232/RS-422/RS-485)	
	Audio	Line-out, Mic-in	
Expansion	M.2	1 x M.2 2230 E Key (PCIe x1; USB2.0) for WIFI/BT 1 x M.2 3052 B Key (USB3.2 x1; USB2.0) for 5G	
Others	TPM	TPM2.0	
Storage	M.2	1 x M.2 2280 M Key (PCIe x4) for NVME 1 x M.2 2242 M Key (PCIe x2) for NVME	
Power	Power Input	19~36V DC input	
	Power Type	ATX/AT mode, ATX default	
Power Consumption	Typical (OS idle mode)	13.2W	12W
	Max. (Full loading)	156W	104W
Environment	Op.Temp	-15 ~ 40°C (non-POE) / -15 ~ 35°C (POE) with 0.7 m/s airflow (TDP 130W)	
	Op. Humidity	95% @ 40 °C (non-condensing)	
	Vibration	3 Grms @ 5 ~ 500 Hz, random, 1 hr/axis	
Mechanical	Dimensions (W x D x H)	235.8 x 174 x 59 mm	
	Weight	2.89Kg	
	Mounting Support	Wall mounting	
Operating System		Ubuntu 24.04 LTS with NVIDIA Jetpack™ 7.1	
Software Support	Software API	WEDA / SUSI / Edge AI SDK / Inference Kit compatible	
Certifications	EMC/Safety	CE/FCC Class B, CB, UL, CCC and BSMI (No RED Certificate)	

Advantech SUSI is a device management and system monitoring suite for hardware configuration, control, and status monitoring.

SUSI information: <https://github.com/ADVANTECH-Corp/SUSI>

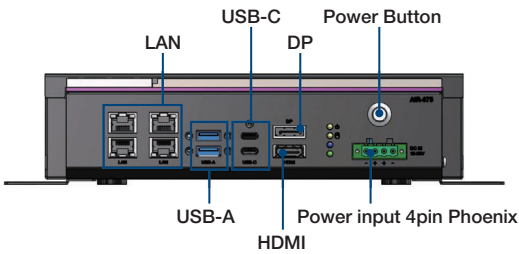
Dimensions

Unit: mm



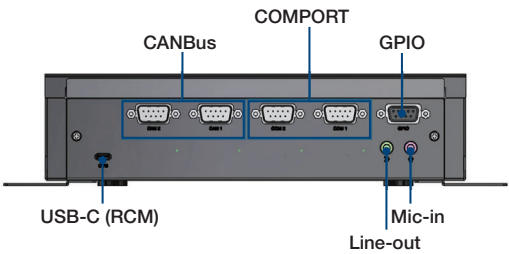
Front Panel I/O Placement

T5000 SKU

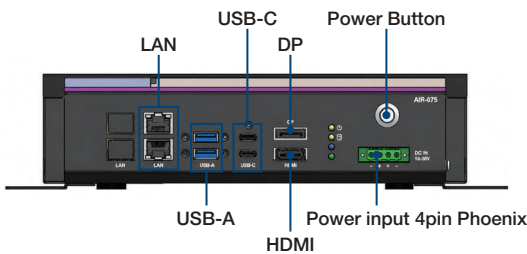


Back Panel I/O Placement

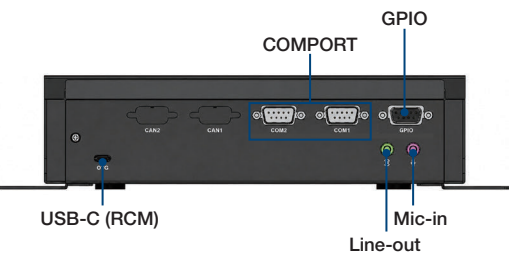
T5000 SKU



T4000 SKU



T4000 SKU



Ordering Information

Part No.	Module	CPU	GPU	AI Performance	Memory	Storage	10GbE	HDMI	DP	USB	CANBus	DIO	RS-232/422/485	Power Input	Operating Temperature	Made in	Remark
AIR-075A-S6A1U	NVIDIA Jetson T5000	14-core Arm® Neoverse®-V3AE 64-bit CPU	2560 NVIDIA Blackwell cores and 96 Cores	2070 TFLOPS (FP4—Sparse)	128GB	512GB	4	1	1	4	2	8 bit	2	19-36V _{DC}	-10~40°C	Taiwan	Made in TW
AIR-075A-S6B1U	NVIDIA Jetson T4000	12-core Arm® Neoverse®-V3AE 64-bit CPU	1536 NVIDIA Blackwell cores and 64 Tensor Cores	1200 TFLOPS (FP4—Sparse)	64GB	512GB	2	1	1	4	NA	8 bit	2	19-36V _{DC}	-10~50°C	Taiwan	Made in TW
AIR-075A-S6A2U	NVIDIA Jetson T5000	14-core Arm® Neoverse®-V3AE 64-bit CPU	2560 NVIDIA Blackwell cores and 96 Cores	2070 TFLOPS (FP4—Sparse)	128GB	512GB	4	1	1	4	2	8 bit	2	19-36V _{DC}	-10~40°C	Taiwan	PCBA CH, Assemble TW
AIR-075A-S6A3U	NVIDIA Jetson T5000	14-core Arm® Neoverse®-V3AE 64-bit CPU	2560 NVIDIA Blackwell cores and 96 Cores	2070 TFLOPS (FP4—Sparse)	128GB	NA	4	1	1	4	2	8 bit	2	19-36V _{DC}	-10~40°C	Taiwan	PCBA CH, Assemble TW, No SSD
AIR-075A-S6B3U	NVIDIA Jetson T4000	12-core Arm® Neoverse®-V3AE 64-bit CPU	1536 NVIDIA Blackwell cores and 64 Tensor Cores	1200 TFLOPS (FP4—Sparse)	64GB	NA	2	1	1	4	NA	8 bit	2	19-36V _{DC}	-10~50°C	Taiwan	PCBA CH, Assemble TW, No SSD
AIR-075A-S6B4U	NVIDIA Jetson T4000	12-core Arm® Neoverse®-V3AE 64-bit CPU	1536 NVIDIA Blackwell cores and 64 Tensor Cores	1200 TFLOPS (FP4—Sparse)	64GB	NA	2	1	1	4	NA	8 bit	2	19-36V _{DC}	-10~50°C	Taiwan	Made in TW, No SSD

Packing List

Part Number	Description	Quantity
AIR-075	NVIDIA AI Inference System	1
1652003234	Phoenix connector counterpart	1
-	Wall Mount Kit for AIR-075	1
-	Simplified Chinese User Manual	1
20706UJTNS0040	Image Ubuntu v300 AIR-075	1

*In case of any changes, the actual packaging shall prevail.

Optional

Part Number	Description
96PSA-A330W24P4-3	A/D 100-240V 330W 24V C14 TERMINAL BLOCK 4P
1702002600	Power Cord UL 3P 10A 125V 183cm (US)
1702002605	Power Cord EU 3P 10A 250V 183cm (EU)
1702031801	Power Cord BSI 3P 10A 250V 183cm (UK)
1700000237	Power Cord PSE 3P 12A 125V 183cm (Japan)
1700013977	Power Cord CCC 3P 10A 250V 200cm 90°(China)
AIW-169BR-GX1	Realtek Wi-Fi 6E M.2 2230 E-Key
1751000620-01	1x Cable Ant. L300mm for WIFI
1751000651-01	1x Antenna for WIFI
AIW-356DQ-JK1	Qualcomm 5G M.2 3052 B-Key (JK SKU)
AIW-356DQ-E01	Qualcomm 5G M.2 3052 B-Key (EU SKU)
AIW-356DQ-N01	Qualcomm 5G M.2 3052 B-Key (US SKU)
1751000623-01	1x Cable Ant. L300mm for 5G
1750009372-01	1x Antenna for 5G
MIOE-PSE-DPA1	MIOe-PSE POE module
1700034631-01	1pcs M cable USB-A 4P(M)/USB-C 24P(M) 90cm (Flashing Cable)
1700031582-01	1pcs F Cable D-SUB 9P(M)/1x5P-1.25 20cm (debug cable)

If a power supply other than Advantech adapter P/N: 96PSA-A330W24P4-3 is used, it must meet the following requirements: Continuous: 200W; Peak: 440W@6μs, 330W@1ms.

Inference Kit | Production-Ready AI Inference on Edge Devices

Provides a unified and hardware-aligned runtime for deploying and validating AI inference on edge devices

It simplifies integration across CPUs, GPUs, and AI accelerators while enabling performance benchmarking and compatibility verification on target hardware. Designed for production use, Inference Kit helps hardware partners ensure stable, scalable, and repeatable AI deployment across product lines.



EdgeAI SDK Inference Kit

Streamlined Edge Inference

- Ready-to-Run Inference Runtime
- Accelerator-Aware Optimization
- Stable Edge Production Stack
- Unified Inference Interfaces





Benefits and Features

 Unified Inference Runtime - Consistent inference across CPUs, GPUs, and accelerators - Vendor-optimized runtime integration - Built-in UniInfra acceleration framework - Optimized inference pipelines and runtime efficiency	 Hardware Validation - Benchmarking on target devices - OS and accelerator compatibility validation - Performance and stability verification
 Production-Ready Deployment - Stable, long-running inference operation - System monitoring and observability support - Designed for scalable edge deployment	 Global Customer Support - System reliability certification - Inference computing enablement - Edge-to-cloud scalability collaboration

