

MegaEdge x86-Based Edge AI Solution

Optimized x86 Architecture for Scalable,
High-Performance AI Inference

MegaEdge x86-Based Edge AI Solution



PCIe / MXM / M.2 series

Optimized x86 Architecture for Scalable, High-Performance AI Inference

Aetina MegaEdge is a series of expandable x86-based AI inference systems delivering high-performance computing for generative AI and scalable vision AI workloads. Supporting GPUs and AI accelerators via PCIe, MXM, and M.2, MegaEdge provides a versatile AI computing platform with rich I/O for seamless integration of sensors, cameras, displays, and robotic equipment. Designed for rapid deployment, it targets diverse AI applications across smart city, factory, retail, healthcare, security, banking, and smart agriculture.

Why Aetina MegaEdge

◆ Worry-free from Reliability

Providing high reliability and optimized performance with pattern thermal design and vigorous shock and vibration testing, avoiding downtime and hassles.

◆ Optimal Heterogeneous Computing Solutions

MegaEdge delivers diverse heterogeneous computing solutions, fusing CPU, GPU, and ASIC synergy for targeted computing excellence, and offering various extension along with GPU/AI Accelerator modules for varying AI project demands.

◆ Simplified Fleet Management

MegaEdge supports EdgeEye, Aetina's cloud management platform, which serves as a unified platform to help clients efficiently and systematically manage edge devices.

◆ AI-enabled Turnkey Solutions

By collaborating with Aetina's ecosystem partners, MegaEdge provides vertical, application-specific solutions. These comprehensive end-to-end offerings help developers and application integrators accelerate the time-to-market for AI deployment.

Product Highlights

◆ Maximum Storage Capacity

Up to 4x 2.5" SATAIII SSD with seamless swappable and installation.

◆ Rich Accessible Front- I/Os

Available wide range of I/O interfaces for high-speed networking and connectivity.



◆ Flexible Expansion Capability

Providing an abundance of PCIe/MXM/M.2 expansion slots allowing for versatile functionality options.

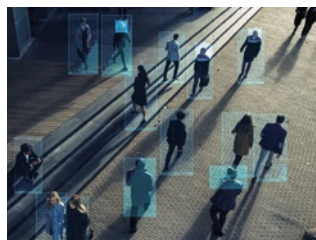
◆ Out-of-Band Management

Built-in with Out of Band embedded module makes remote power on/off/reset possible even if the system crashes.

Targeted Applications



Roadside Vehicle Monitoring



AI Surveillance



AI Defect Inspection



Agentic AI

PCIe Series

AIP-FR68



Scale Out LLM Inference & Enterprise AI

MegaEdge AIP-FR68 provides unparalleled AI performance and acts as an AI workstation designed to meet state-of-the-art edge AI technology. Certified for NVIDIA NCS (NVIDIA Certified System), the AIP-FR68 has integrated the power of Intel 12/13/14th Gen Core processors and NVIDIA Data Center PCIe AI cards. This x86 edge AI workstation is crafted for the most challenging AI tasks, operating seamlessly and reliably even under intensive workloads.

◆ Features

- Powered by Intel 12/13/14th Gen Core™ i5/i7/i9 processors, TDP up to 65W
- Support 2x NVIDIA Passive AI card - AI Inference and Mainstream Computing
- 2 x M.2 M-Key 2280 supports NVME, AI accelerator card
- 1 x 10GbE LAN and 3 x 2.5GbE LAN Ports
- Up to 4 x 2.5" SSD and support RAID 0/1/5/10
- Built-in Out-of-Band Remote management module

◆ Obtain AI Surge with NVIDIA RTX Power

Achieve AI excellence with AIP-FR68, featuring NVIDIA RTX™ PRO 6000 Max-Q GPUs for high-performance LLM inference and AI rendering, or Qualcomm® Cloud AI 100 Ultra accelerators for highly efficient, cost-optimized on-premises LLM deployment.



Unprecedented
AI Performance



96GB GDDR7
Memory



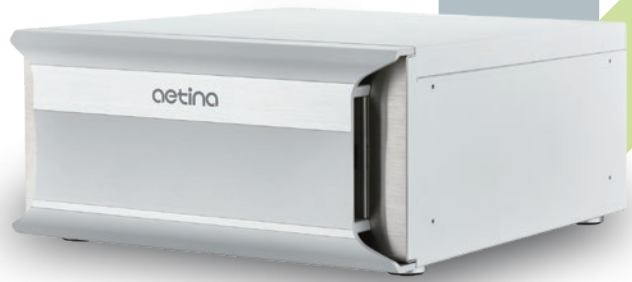
New-Gen CUDA
Computing



Ultimate MLPerf
Inference

PCIe Series

AIP-KQ67



Intensify Computer Vision & AI Inference

Unleashing peak AI performance at the edge, the MegaEdge AIP-KQ67 is a GPU-expansion platform which equips you with the latest Intel 12/13/14th Core i9/i7/i5 processors and powerful NVIDIA RTX PRO series GPUs for cutting-edge computer vision and unmatched AI inference processing. Its scalable design, diverse expansion options, and user-friendly screw-less chassis empower customers to revolutionize their AI deployments across smart city, security, automation, and manufacturing.

◆ Features

- Powered by Intel 12/13/14th Gen Core™ i5/i7/i9 processors, TDP up to 125W
- Support Four DDR5 U-DIMM up to 192GB
- 1 x PCIe16 Gen5, 2 x PCIe4 Gen4, 1 x PCIe2 Gen3, 1x M.2 M-Key, 1x M.2 E-Key
- PCIe x16 Graphics card support up to 300W paired with 850W PSU
- Support Aetina EdgeEye, enhancing device remote monitoring

◆ Elevate AI Edge Inference with NVIDIA RTX Superiority

AIP-KQ67 incorporates unparalleled expansion flexibility with GPUs, offering accelerated AI inference computing for edge AI solutions. Certified by NVIDIA NCS certification with NVIDIA A2 Tensor Core GPU, it also supports a spectrum of NVIDIA RTX PRO GPUs, including the RTX PRO 6000 Max-Q / PRO 5000 PCIe GPUs, contributing superior AI-driven advantages and performance to cater to your diverse AI applications at the edge.



Efficient Price-
Performance



On-premises AI
Inference



Computer Vision
Optimization



Diverse GPU
Options



PCIe Series

AIP-FR68S

Unlock Flagship On-Prem LLM AI Engine

MegaEdge AIP-FR68S AI Workstation is the flagship addition to the MegaEdge series, purpose-built for real-time LLM inference, multimodal analysis, and advanced rendering. It supports both data center and workstation-class GPUs—including NVIDIA RTX PRO 6000 Blackwell—as well as up to three Qualcomm Cloud AI 100 Ultra accelerators (870 TOPS INT8 each). With a PCIe Gen 5 high-speed switch, AIP-FR68S unlocks exceptional AI performance for demanding enterprise workloads.

◆ Features

- Intel 12/13/14th Gen Core™ i9/i7/i5 processor, TDP up to 65W
- PCIe Gen5 x16 Support: Enables seamless integration with the latest GPU
- Integrated PCIe Gen5 Switch: Supports unified computing across multiple Qualcomm AI100 Ultra cards
- Swappable 2000W Power Supply: Delivers ample, replaceable power for high-performance GPUs and AI accelerators
- Built-In GPU Holder: Ensures GPU stability, even in rugged operating environments

◆ Enterprise-Grade GenAI Engine for Agentic Intelligence

Unlock on-prem agentic AI, advanced vision AI with AIP-FR68S—supporting NVIDIA RTX™ Pro 6000 Blackwell and Qualcomm Cloud AI 100 Ultra. Deliver scalable, high-performance edge computing built to accelerate enterprise AI across industries.



On-Prem Gen
AI Performance



192GB
DDR5 Memory



Next-Gen AI Core



Streamline AI
Development

PCIe Series

AIP-RQ87



Optimized Rack Design System for AI Video and Robotics

MegaEdge AIP-RQ87 rack system, powered by Intel® Core™ Ultra i9K/i9/i7 processors (Arrow Lake-S), combines enterprise-grade rack management with a front-access 2U design. Optimized for advanced video analysis and robotic control, it supports NVIDIA RTX™ PRO 6000 Blackwell Max-Q, PRO 5000, and Qualcomm Cloud AI 100 Ultra accelerators (870 TOPS INT8) via high-speed PCIe Gen 5.

Features

- Powered by Q870 chipset with PCIe Gen5 x16 for Blackwell Pro 6000 and Qualcomm AI100 Ultra.
- High-Capacity DDR5 Memory: Supports 4x DDR5 U-DIMM, offering massive scalability up to 256GB.
- Flexible Storage Options: drive bays supporting up to 5x 2.5" SATA SSDs or 3x 3.5" HDDs for diverse data needs.
- 2U Front-Access Design: User-friendly front I/O with status LEDs, 2x USB 2.0 and reset button for effortless maintenance.
- Rugged GPU Stabilization: Built-in heavy-duty holder ensures GPU stability and signal integrity in vibration-prone environments.
- Smart Thermal Management: Autonomous fan control optimizes cooling for high-performance AI computational loads.

Power Video and Robotic Intelligence with Optimized Rack-Ready System

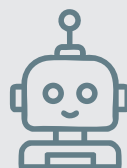
Enable low-latency video processing and robotics workloads with AIP-RQ87, featuring NVIDIA RTX™ PRO 6000 Blackwell and Qualcomm Cloud AI 100 Ultra accelerators.



Optimized 2U
Rack Design



256GB
DDR5 Memory



Video &
Robotic Intelligence



Enterprise-grade
rack management

MXM Series

AIP-SQ67



Enriched Edge Inference & Data Connectivity

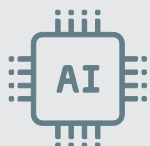
MegaEdge AIP-SQ67 is an expandable AI inference platform that gives AI developers and system integrators superior graphics and edge AI capabilities. Benefiting from 12/13/14th Gen Intel Core™ processors and one extension AI accelerator MXM module slot built for boosting superlative AI acceleration, driving massive performance-driven tasks with up to 65W TDP. Low power profiles also reduce power consumption, which can help developers meet AI-multitasking and product sustainability goals.

◆ Features

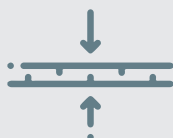
- Powered by Intel 12/13/14th Gen Core™ i5/i7/i9 processors, TDP up to 65W
- 5 x DP++ (4 from MXM GPU Card, 1 from CPU)
- 2 x M.2 M-Key with NVMe/SATA
- 5 x 2.5GbE LAN Ports
- 6 x USB 3.2 Gen2 Type-A & 1 x USB 3.2 Gen2 Type-C
- Built-in out-of-band remote management module
- Swappable drive bay for 2 x 2.5" SATA/III SSD

◆ Unleash AI Brilliance with MXM Module

Equipped with an extended MXM GPU for top-tier AI processing that allows advanced sophisticated AI applications right at the edge, experiencing rapid, insightful, and smart AI decision making where it matters most.



Extended AI Power



Thin & Robust



Ready-to-Use



Conformal Coating

M.2 Series**AIP-MURE**

Highly Efficient & Real-time AI Visual Inference at the Edge

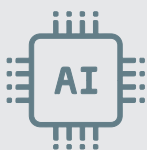
The AIP-MURE is a fanless x86 edge AI inference system powered by energy-efficient Intel 13th Gen i7/i5/i3 processors (15W). It enhances AI inference performance with three accelerators, delivering up to 300 TOPS of power while maintaining low energy consumption. Ideal for real-time AI inference and video stream analysis, the AIP-MURE excels in smart surveillance, security, and smart city applications. Its compact, reliable, and low-power design makes it perfect for demanding edge environments, offering a cost-effective and flexible solution for diverse Edge AI applications.

◆ Features

- Feature with 3x M.2 M-key for AI Accelerator cards, boosting extreme AI Inference performance.
- Fanless & rugged chassis for harsh environment.
- Dual PSE support IP Camera for real time video streaming.
- Support out-of-band (OOB) management tool for remote maintenance of edge AI system.

◆ Empower Edge Intelligence with Efficient & Cost-effective M.2 Acceleration

AIP-MURE elevates edge AI by integrating three powerful M.2 AI accelerators with a reliable system design — a compact, fanless, and rugged chassis — along with real-time visual inference capabilities. This powerful yet cost-effective solution is designed to unlock the full potential of AI visual computing at the edge, propelling your AI initiatives to new heights with unmatched speed and efficiency.



Extended AI Power



Fanless System



Cost Efficiency



Low-power Consumption

Selection Guide



Model Number	AIP-FR68	AIP-KQ67	AIP-FR68S
CPU	Intel 12/13/14th Core i5/i7/i9 Processor up to 65W	Intel 12/13/14th Core i5/i7/i9 Processor up to 125W	Supports 12/13/14th Gen. Intel Core i9/i7/i5 Processor up to 65W
Chipset	Intel R680E Chipset	Intel Q670E Chipset	Intel R680E Chipset
GPU(Optional)	NVIDIA: RTX Blackwell PRO 6000 Max-Q 96G/ PRO 5000 48G/PRO 4500 32G/ PRO 4000 24G/PRO 2000 16G Qualcomm: AI 100 Ultra/AI 80 Ultra	RTX Blackwell PRO 6000 Max-Q 96G/ PRO 5000 48G/PRO 4500 32G/ PRO 4000 24G/PRO 2000 16G	NVIDIA: RTX Blackwell PRO 6000 Max-Q 96G/ PRO 5000 48G/PRO 4500 32G/PRO 4000 24G/ PRO 2000 16G Qualcomm: AI 100 Ultra/AI 80 Ultra
Memory	4 x DDR5 U-DIMM up to 192G Support 4400MT/s	4 x DDR5 U-DIMM up to 192G (4000MT/s with 4 DIMM)	4 x DDR5 288-pin U-DIMM up to 192G
Storage	4 x 2.5" SATAIII SSD (Support RAID 0/1/5/10) 1 x M.2 M-Key Slot, Support PCIe Gen 4 x4 for NVME/SATA, Size 2280 1 x M.2 M-Key Slot, Support PCIe Gen 4 x4 for NVME, Size 2280	2 x 2.5" SATAIII SSD (500W PSU SKU) 1 x M.2 M-Key Slot, Support PCIe Gen4 x4 for NVME/SATA, Size 2280 (500W & 850W PSU SKU)	2 x 2.5" SATAIII SSD (RAID 0/1) 1 x M.2 M-Key Slot, Support PCIe Gen4 x4 for NVME/SATA, Size 2280 1 x M.2 M-Key Slot, Support PCIe Gen4 x4 for NVME, Size 2280
Front I/O	8 x USB 3.2 Gen2 Type-A(10G) 1 x USB 3.2 Gen2 Type-C(20G) 3 x 2.5GbE LAN (RJ-45) 1 x 10GbE LAN (RJ-45) 1 x RS-485(RJ-45) 2 x RS-232/422/485 (DB9) 1 x Line out 1 x Mic in 1 x 8bits GPIO (DB15) 2 x DP++1.4 1 x HDMI 2.0 1 x Power Button	1 x Power button 2 x USB 2.0	8 x USB 3.2 Gen2 x1 (Type-A) 1 x USB 3.2 Gen2 x2 (Type-C) 2 x 2.5G RJ-45 Ethernet 2 x RS-232/422/485 (COM1/2, DB9, Default RS-232) 1 x HDMI 2.0 2 x DP++1.4
Rear I/O	6 x SMA Antenna hole	6 x USB 3.2 Gen2 Type-A(10G) 1 x USB 3.2 Gen2 Type-C(20G) 3 x 2.5GbE LAN (RJ45) 1 x 10GbE LAN (RJ45) 2 x DP++1.4 1 x HDMI 2.1 3 x Jack support Line out/Line in/Mic in 2 x SMA Antenna hole	6 x SMA Antenna hole
Internal I/O	N/A	N/A	N/A
Expansion	1x PCIe Gen4 x16 or Gen4 x8 (PCIe x16 slot) 1x PCIe Gen4 x8 (PCIe x16 slot) 1x PCIe Gen3 x4 (PCIe x4 slot) 1x PCIe Gen3 x4 (PCIe x16 slot) 1x PCIe Gen3 x1 (PCIe x1 slot) 1 x M.2 2230 E-Key 1 x M.2 3052 B-Key 1 x Nano SIM slot	1 x PCIe Gen5 x16 (PCIe x16 slot) 2 x PCIe Gen4 x4 (PCIe x4 slot) 1 x PCIe Gen3 x2 (PCIe x4 slot) 1 x M.2 2230 E-Key	1x M.2 E-Key 2230 (USB2.0/PCIe) 1x M.2 B-Key 3052(USB2.0/USB3.2 SKU1 : 3x PCIe x16 (PCIe x16 slots) (w/PCIe Gen5 Switch) SKU2 : 1x PCIe Gen5 x16 or Gen5 x8 (PCIe x16 slot) 1x PCIe Gen5 x8 (PCIe x16 slot) 1x PCIe Gen4 x4 (PCIe x4 slot)
MISC. Function	OOB (out of band), Built-in Innodisk – InnoAgent	N/A	OOB (out of band), Built-in Innodisk – InnoAgent
Power Consumption	Full loading: up to 600 Watts	Full loading: up to 500 Watts/850 Watts	Full loading: up to 2000 W
Power Input / Connector	DC-in 24 to 48VDC, 4-pin Terminal Block	500W: AC-in 110-240VAC,60~50 Hz, 8-4A(FLEX ATX PSU) 850W: AC-in 110-240VAC,60~50 Hz, 10-6A(FLEX ATX PSU)	AC-in 110~220VAC
Dimension (W x D x H)	340 x 213 x 279 mm	315 x 413 x 159mm	240 x 400 x 206 mm
Mounting	Wallmount / Deskmount	Desktop	Wallmount / Deskmount
Net Weight	11 Kg	10 kg	14.45 kg
Vibration	1Grms ,IEC60068-2-64,Random, 5 ~ 500 Hz ,1Hr / Axis (operation, w/o PCIe Card)	1Grms ,IEC60068-2-64,Random, 5 ~ 500 Hz ,1Hr / Axis (operation, w/o PCIe Card)	1Grms ,IEC60068-2-64,Random, 5 ~ 500 Hz ,1Hr / Axis (operation, w/o PCIe Card)
Shock	10G, IEC 60068-2-27, Half Sine, 11 ms Duration (operation, w/o PCIe Card)	10G, IEC 60068-2-27, Half Sine, 11 ms Duration (operation, w/o PCIe Card)	10G, IEC 60068-2-27, Half Sine, 11 ms Duration (operation, w/o PCIe Card)
Temperature	Operating Temp. : 0°C ~ +50°C (32°F ~ 122°F) (w/o GPU) Operating Temp. : 0°C ~ +50°C (32°F ~ 122°F) (w/ 165W Passive GPU) Operating Temp. : 0°C ~ +40°C (32°F ~ 104°F) (w/ 300W Active GPU) Storage Temp. : -40°C ~ +85°C (-40°F ~ 185°F)	Operating Temp.: -10°C ~ +50°C, with 0.5m/s air flow (w/ 120W Passive GPU) Operating Temp.: -10°C ~ +40°C, with 0.5m/s air flow (w/ 300W Passive GPU) Storage Temp.: -40°C ~ +85°C (-40°F ~ 185°F)	Operating Temp.: 0°C ~ +50°C (32°F ~ 122°F)(w/o GPU) Operating Temp.: 0°C ~ +40°C (32°F ~ 122°F)(w/ GPU) Storage Temp.: -40°C ~ +85°C (-40°F ~ 185°F)
Humidity	95% @ 40°C Related Humidity (non-condensing)	95% @ 40°C Related Humidity (non-condensing)	95% @ 40°C Related Humidity (non-condensing)
OS Support	Windows 10 IOT / Windows 11 IOT/ Ubuntu 22.04	Windows 11 IOT/ Ubuntu 22.04	Windows 10 IOT / Windows 11 IOT/ Ubuntu 22.04
Certification	CE / FCC Class A / UKCA / LVD / RoHS / BIS	CE / FCC Class A / UKCA / LVD / RoHS	CE / FCC Class A / UKCA / LVD / RoHS



Model Number	AIP-RQ87	AIP-SQ67	AIP-MURE
CPU	Intel Core™ Ultra i9K/i9/i7 Processor(Arrow Lake-S), up to 125W	Intel 12/13/14 th Core i5/i7/i9 Processor up to 65W	Intel 13 th Core i3/i5/i7 Processor up to 15W
Chipset	Intel Q870 Chipset	Intel Q670E Chipset	SoC
GPU(Optional)	NVIDIA: RTX Blackwell PRO 6000 Max-Q 96G/ PRO 5000 48G/PRO 4500 32G/PRO 4000 24G/ PRO 2000 16G Qualcomm: AI 100 Ultra/AI 80 Ultra	Blackwell Embedded: RTX PRO 5000 Ada Lovelace: RTX 2000/3500/5000 Ampere: A4500/A2000/A1000/A500	AXELERA AI ACCELERATOR M.2 2280 MODULE
Memory	4 x 288-pin CU-DIMM DDR5 6400 MHz, U-DIMM DDR5 5600 MHz, up to 256GB (64GB per DIMM)	2 x DDR5 SO-DIMM up to 64G Support 4800MT/s	2 x DDR5 SO-DIMM up to 96G Support 5200MT/s
Storage	5 x 2.5" SATAIII SSD (Support RAID 0/1/5/10) or 3 x 3.5" HDD 1 x M.2 (Key M, 2242/2280) with PCIe Gen5 x4 for SSD 1 x M.2 (Key M, 2242/2280) with PCIe Gen4 x4 for SSD	2 x 2.5" SATAIII SSD 2 x M.2 M Key Slot, Support PCIe Gen 4 x4 for NVME/SATA, Size 2280	1 x 2.5" SATAIII SSD 2 x M.2 M-Key Slot, Support PCIe Gen4 x4 for NVME, Size 2280
Front I/O	1 x Power button 2 x USB 2.0 1 x Reset Button	6 x USB3.2 Gen2 Type-A(10G) 1 x USB3.2 Gen2 Type-C(20G) 5 x 2.5GbE LAN (RJ-45) 4 x RS-232/422/485 (DB9) 1 x 7bits GPIO(DB9) 1 x Power Button	2 x USB 3.2 Gen2 Type-A(10G) 2 x USB 2.0 Type-A 2 x 1GbE LAN (RJ-45, 2 x support Dual POE 802.3 at) 1 x DP++1.4 1 x HDMI 2.0 1 x Line out 1 x Mic in 2 x SMA Antenna hole
Rear I/O	1 x HDMI 2.1 1 x DP 1.4a++ 1 x 1 Gigabit LAN 1 x 2.5 Gigabit LAN 7 x USB 3.2 Gen2 1 x USB4/Thunderbolt™4 (5V/3A, supports DP 2.1 display output) 1 x COM1 (RS-232/422/485)	5 x DP++ (HBR2) 1 x Line Out 1 x Mic in 2 x CAN Bus-Isolation 2.0B (optional) 2 x USB2.0 Type-A	1 x Power Button 1 x Rest Button 1 x 8bits GPIO (DB15) 2 x RS-232/422/485 (DB9) 1 x VGA 4 x SMA Antenna hole
Internal I/O	N/A	N/A	1 x USB 2.0 Type-A
Expansion	2 x PCIe Gen5 Slots (PCIe1/PCIe4:single at x16 (PCIe1); dual at x8 (PCIe1)/x8 (PCIe4)) (PCIe1: Support Riser card x8/x8) 2 x PCIe x4 (Gen4) 1 x M.2 (Key E, 2230) with PCIe Gen4 x1	1 x MXM Slot Gen4 x16 (Type B+)	1 x M.2 2230 E-Key 1 x M.2 2242/3042/3052 B-Key 1 x Nano SIM slot
MISC. Function	OOB (out of band), Built-in Innodisk – InnoAgent(Optional)	OOB (out of band), Built-in Innodisk – InnoAgent	OOB (out of band)(optional), Built-in Innodisk – InnoAgent
Power Consumption	Full loading: up to 2000W	Full loading: up to 600 Watts	Full loading: up to 130 Watts
Power Input / Connector	AC-in 110~220VAC	DC-in 24VDC / 4-pin Terminal Block	DC-in 12~24VDC / 4-pin Terminal Block
Dimension (W x D x H)	438 x 500 x 88 mm	270 x 280 x 150mm	270 x 195 x 80 mm
Mounting	Desktop	Deskmount	Din-rail / Wallmount / Deskmount
Net Weight	TBD	5.5 kg	3.82Kg
Vibration	1Grms ,IEC60068-2-64,Random, 5 ~ 500 Hz ,1Hr / Axis (operation, w/o PCIe Card)	2Grms ,IEC60068-2-64,Random, 5 ~ 500 Hz ,1Hr / Axis (operation)	1Grms ,IEC60068-2-64,Random, 5 ~ 500 Hz ,1Hr / Axis (op mode) 1.15Grms , ISTA 2A Random with package, 1 ~ 200 Hz , 30 min / Axis (non-op mode)
Shock	10G, IEC 60068-2-27, Half Sine, 11 ms Duration (operation, w/o PCIe Card)	5G, IEC 60068-2-27, Half Sine, 11 ms Duration (operation)	10G, IEC 60068-2-27, Half Sine, 11 ms Duration (op mode) 30G, IEC 60068-2-27, Half Sine, 1 ms Duration (non-op mode)
Temperature	Operating Temp.: -20°C ~ +40°C (32°F ~ 104°F)(w/ 600W Active GPU) Storage Temp.: -40°C ~ +85°C (- 40°F ~ 185°F)	Operating Temp. : 0°C ~ +50°C (With MXM) Storage Temp. : -40°C ~ +85°C (-40°F ~ 185°F)	Operating Temp.: -30°C ~ +60°C (-22°F ~ 140°F) Storage Temp.: -40°C ~ +85°C (-40°F ~ 185°F)
Humidity	95% @ 40°C Related Humidity (non-condensing)	95% @ 40°C Related Humidity (non-condensing)	95% @ 40°C Related Humidity (non-condensing)
OS Support	Windows 10 IOT / Windows 11 IOT/ Ubuntu 22.04	Windows 11 IOT/ Ubuntu 22.04	Windows 11 IOT/ Linux Ubuntu 22.04/24.04
Certification	CE / FCC Class A / UKCA / LVD / RoHS	CE / FCC Class A / UKCA / LVD / RoHS	CE / FCC Class A / UKCA / LVD / RoHS

AI On-Prem Solution



Unleash Next-Level On Prem Gen AI with Aetina Workstation and Qualcomm® Cloud AI 100 Ultra Card

Transform your enterprise AI capabilities with Aetina's pioneering on-premises solution. The Aetina MegaEdge AIP-FR68, featuring the Qualcomm® Cloud AI 100 Ultra inference card, delivers exceptional performance for generative AI, computer vision, and LLMs up to 70B parameters. This compact desktop, wall-powered AI workstation supports real-time AI agents, workflow automation, and intelligent search while ensuring secure, low-latency processing at industry-leading power efficiency and superior total cost of ownership.

► Flagship Base LLM Inference



MegaEdge AIP-FR68

- Intel 12/13/14th Gen Core™ i9/i7/i5/i3 CPU (65W)
- Support dual passive Qualcomm Cloud AI 100 Ultra cards
- 4x 2.5" SATAIII SSD, supports RAID configurations
- 2x M.2 2280 M-Key slots, supporting Gen4 x4 NVMe
- 1x 10G Ethernet and 3x 2.5G RJ-45 Ethernet
- OOB management tool for remote maintenance

Qualcomm Cloud AI 100 Ultra Card

- PCIe FH3/4L, Single slot
- Support 870 TOPS ML capacity (INT8)
- PCIe Gen4 x16
- 128GB on board DRAM (with ECC)
- 548GB/s memory bandwidth
- 150W power consumption

► Flagship Advanced LLM Inference at Scale



MegaEdge AIP-FR68S

- Intel 12/13/14th Gen Core™ i9/i7/i5 CPU (65W)
- Integrated PCIe Gen5 Switch: Supports unified computing across multiple Qualcomm AI100 Ultra cards
- Swappable 2000W Power Supply: Delivers ample, replaceable power for high-performance GPUs and AI accelerators
- Built-In GPU Holder: Ensures GPU stability, even in rugged operating environments

Qualcomm Cloud AI 100 Ultra Card

- PCIe FH3/4L, Single slot
- Support 870 TOPS ML capacity (INT8)
- PCIe Gen4 x16
- 128GB on board DRAM (with ECC)
- 548GB/s memory bandwidth
- 150W power consumption

Accelerate your AI Deployment with Local AI Services

Aetina's expertise in industrial computing joins forces with Qualcomm Technologies' cutting-edge inference accelerators and advanced software stack, including the Qualcomm® AI Inference Suite—featuring ready-to-use AI applications, tools, and libraries for scaling generative AI. This powerful AI On-Prem Solution supports a wide range of local AI services, from computer vision to generative AI, offering capabilities such as OpenAI-compatible APIs with user management, voice agents in-a-box, and offloading of large language model (LLM) and retrieval-augmented generation (RAG) functions for intelligent indexed search, summarization, image and code generation, and camera processing.

- AI agents
- RAG
- Transcription & Translation
- Code development
- Chatbots
- Text-to-images



RAG-enabled AI Chatbots



AI Agents



AI Workflow Automation

Turbocharge Your Edge AI Operations with AI On-Prem Solution



On-Premise AI
Processing



Superior AI Performance
On a Single Platform



Ready-for-Use
AI Inference Tools



Real-Time AI
Analytics

AI Vision Solutions



Ignite Next-level AI Inference with NPU-Powered Tech

AI Vision Solutions, fueled by the cutting-edge expertise of Aetina's edge AI and Axelera's NPU tech, delivers a revolutionary Edge AI vision solution that merges Aetina's high-performance edge AI systems with Axelera's top-tier AI accelerators. Engineered to meet a broad spectrum of AI demands, it provides unmatched AI inference performance in image classification, object detection, and pose estimation, unlocking new possibilities in smart security, retail, and manufacturing.



High Performance with
Energy Efficiency



Vision
Acceleration



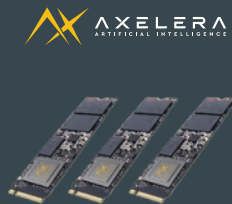
Validated systems
for Quick Adoption



High Flexibility
for Expansion

Edge AI Inference System

M.2



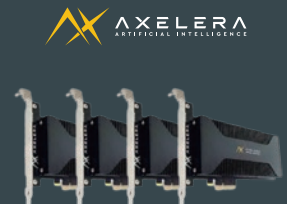
AIP-MURE-A1AXE

**3xM.2
AI accelerator**

- Intel 13th Gen Core™ i7/i5/i3 CPU (15W)
- Up to **642TOPS** AI performance(INT8)
- Pre-installed VOYAGER SDK with Large Model Zoo
- Industrial grade fanless system
- Dual PSE for real-time IP camera streaming
- OOB management tool for remote maintenance

Edge AI Workstation

PCIe



AIP-FR68-A1AXE

**4xPCIe
AI accelerator cards**

- Intel 12/13/14th Gen Core™ i9/i7/i5/i3 CPU (65W)
- Up to **856TOPS** AI performance(INT8)
- Pre-installed VOYAGER SDK with Large Model Zoo
- OOB management tool for remote maintenance

Aetina AI On-Prem Solution Accelerates Enterprise AI Agents for Smarter, Faster On-Prem Document Processing

Enterprises process large volumes of documents daily, such as purchase orders, reimbursement forms, and work logs, often relying on manual or semi-automated workflows that are time-consuming and error-prone. These inefficiencies create operational bottlenecks and limit business growth. To scale AI-driven document intelligence while controlling costs, enterprises increasingly seek compact, cost-effective edge AI solutions, capable of high-performance LLM (large language models) inference on premises, without the complexity of data center-class systems.

Aetina addresses these challenges with its AI On-Prem Solution — built on the AIP-FR68 edge AI x86 workstation, accelerated by Qualcomm Cloud AI100 Ultra accelerator. Delivering strong LLM inference with superior cost and power efficiency, the solution integrates AI assistants, OCR, and ERP management software to automatically extract and structure text and numerical data from documents of various formats. The processed data is seamlessly synchronized with ERP systems, significantly reducing manual input and improving productivity for sales and sourcing teams.

Benefits

- AI On Prem Solution delivers efficient AI inference in scale with improved cost and power efficiency compared to traditional GPU-based solutions
- The space-saving on-prem LLM solution reduces total cost of ownership (TCO) compared to data center-class deployments

Results

- Accurately boosting document processing efficiency by **96%**
- Improved scalability for enterprise document AI deployments
- Faster processing throughput with lower infrastructure and operational costs
- Enhanced overall operational efficiency through optimized edge inference



AI On-Prem
Solution

AIP-FR68
AI Workstation



Qualcomm
Cloud AI 100 Ultra

AI Voice-Controlled Robotic Arms Streamline Multi-variety & Small-batch Production Costs

With the era of intelligent manufacturing and consumer-oriented personalized demands, flexible production mode on multi-variety and small batch has gradually taken the place of mass production on specific single product specification, which better meets the personalized demand from market. However, multi-variety and small batch production mode brings challenges for manufacturers on production transfer costs, such as labor resource and manufacturing setting.

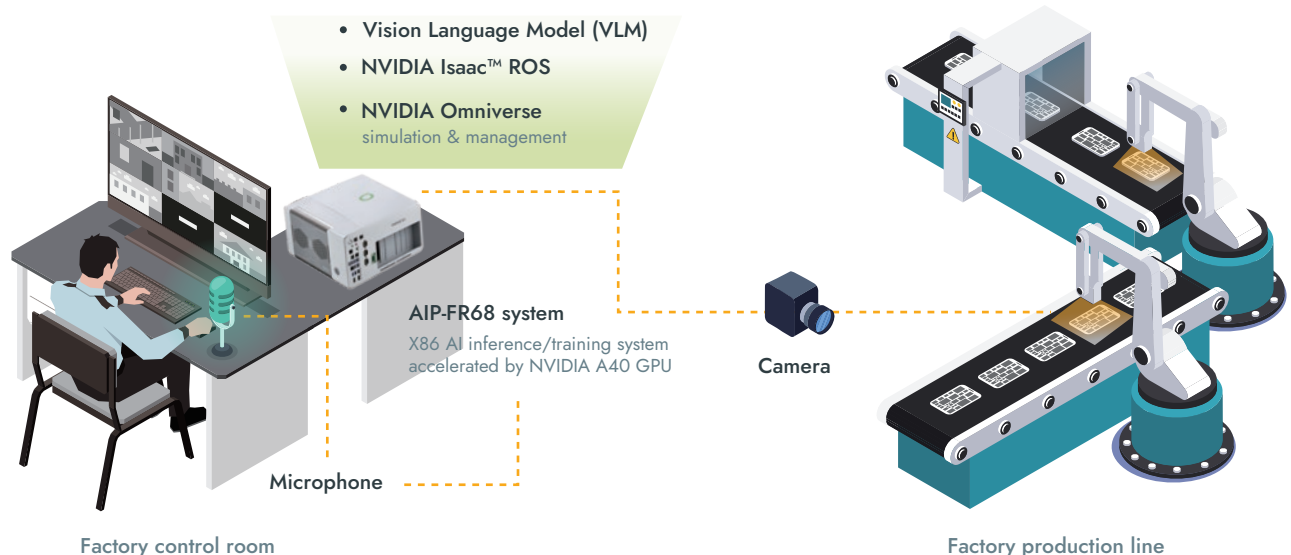
AI-enabled robotic arms are increasingly integrated into production workflows for enhanced productivity. Aetina's MegaEdge AIP-FR68 AI computing system with NVIDIA A40 GPU, combined with large language models (LLMs), empowers robotic arms to understand voice commands, enabling production managers to adjust production line parameters without the need of complicated coding. This reduces changeover costs and simplifies transitions between production lines. Creating digital twins in NVIDIA Omniverse allows robotic arms to interact with the factory environment, enabling managers to simulate operations and ensure safety through real-time monitoring. Once activated, the robotic arms collaborate autonomously to optimize production efficiency. The AIP-FR68 efficiently tackles demanding real-time AI Inference tasks, including neural networks computing and intensive 3D images processing. The solution brings more efficient and productive AI robotic arms, empowering manufacturers to streamline production costs of multi-variety and small-batch production mode and speed up the workflow.

Benefits

- Aetina's NCS AI platforms assure dependable performance, ensuring stability and consistency for critical applications
- AIP-FR68 supports NVIDIA data center GPUs for demanding AI workloads
- Support NVIDIA Omniverse

Results

- No need for complicated code when transitioning the production line, which speeds up the changeover workflow.
- Collaboration between robotic arms enhances production efficiency



◆▶ **Aetina Corporation | Headquarters**

17F, No.237. Sec.1, Datong Rd.,
Xizhi Dist., New Taipei City 221, Taiwan
Tel: +886-2-7709-2568
Fax: +886-2-7746-1102
Email: sales@aetina.com

◆▶ **Aetina Europe B.V.**

Pisanostraat 57, 5623 CB, Eindhoven, The Netherlands
Tel: +31-0-40-3045-400
Email: eusales@aetina.com

◆▶ **Aetina Japan Corporation**

2-35-4 2nd floor, Nihonbashi-Hamacho, Chuo-ku,
Tokyo. Nihonbashi-Hamacho Park building
Mobile: +81-70-3148-9655
Email: sales@aetina.com

◆▶ **Aetina (Shenzhen) Artificial Intelligence Co., LTD**
— Shanghai Office

Room 210, 2F, Building 1, Hansheng Technology
Industries Park, No.3888, Humin Highway, Minhang
District, Shanghai, China
Mobile: +86-153-1699-0980
Email: cnsales@aetina.com

◆▶ **Aetina USA Corporation**

42996 Osgood Road, Suite A, Fremont CA 94539, USA
Tel: +1-510-996-5659
Email: usasales@aetina.com



Website



Linkedin



Facebook



Youtube