

fogwise AIRbox Q900

Compact Industrial Edge AI Box



Powered by Qualcomm Dragonwing IQ-9075

- 8x Big-Core Architecture for Flagship Edge Performance
- Adreno 663 GPU
- 200 TOPS@INT8(Sparse) **200 TOPS**

840MHz

Adreno Neural Network

Kryo Gen6 Up to 2.36GHz	Kryo Gen6 Up to 2.36GHz	Kryo Gen6 Up to 2.36GHz	Kryo Gen6 Up to 2.36GHz
Kryo Gen6 Up to 2.36GHz	Kryo Gen6 Up to 2.36GHz	Kryo Gen6 Up to 2.36GHz	Kryo Gen6 Up to 2.36GHz



OpenGL ES3.2

Vulkan 1.3

OpenCL 2.0 FP

Flagship Video Engine

Multi-Stream UHD, Real-Time Processing

32 cameras
120 FPS encode
240 FPS decode

- AV1, HEVC (H.265), H.264, VP9, MPEG-2 video decoder up to 1x 8Kp60, 4x 4Kp60 or 32x 1080p30
- H.264, H.265, HEIF/ HEIC video encoder up to 2x 4Kp60 or 16x 1080p30
- Concurrent 2x 4Kp60 encode and 2x 4Kp60 decode

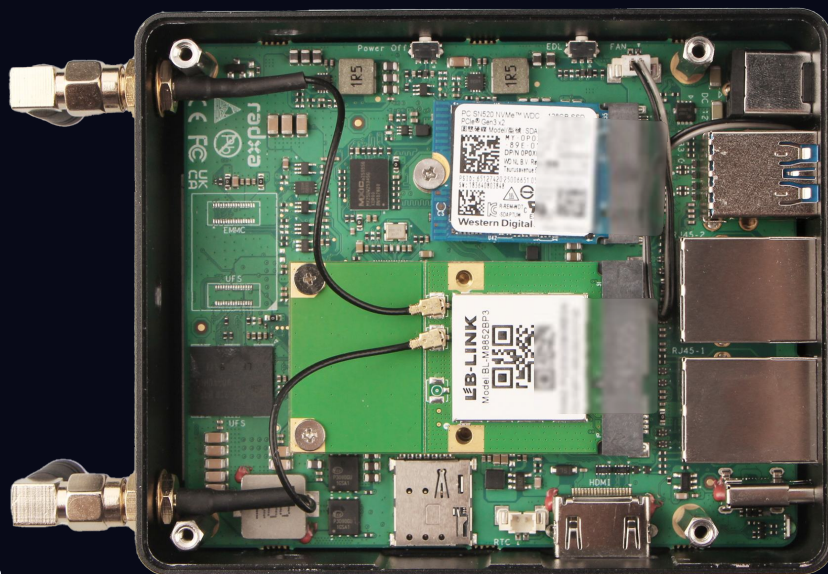
H.265 HEVC

H.264 MPEG-4 AVC

AV1

VP9

High-Speed Memory & Storage Expansion



36GB
LPDDR5@6400

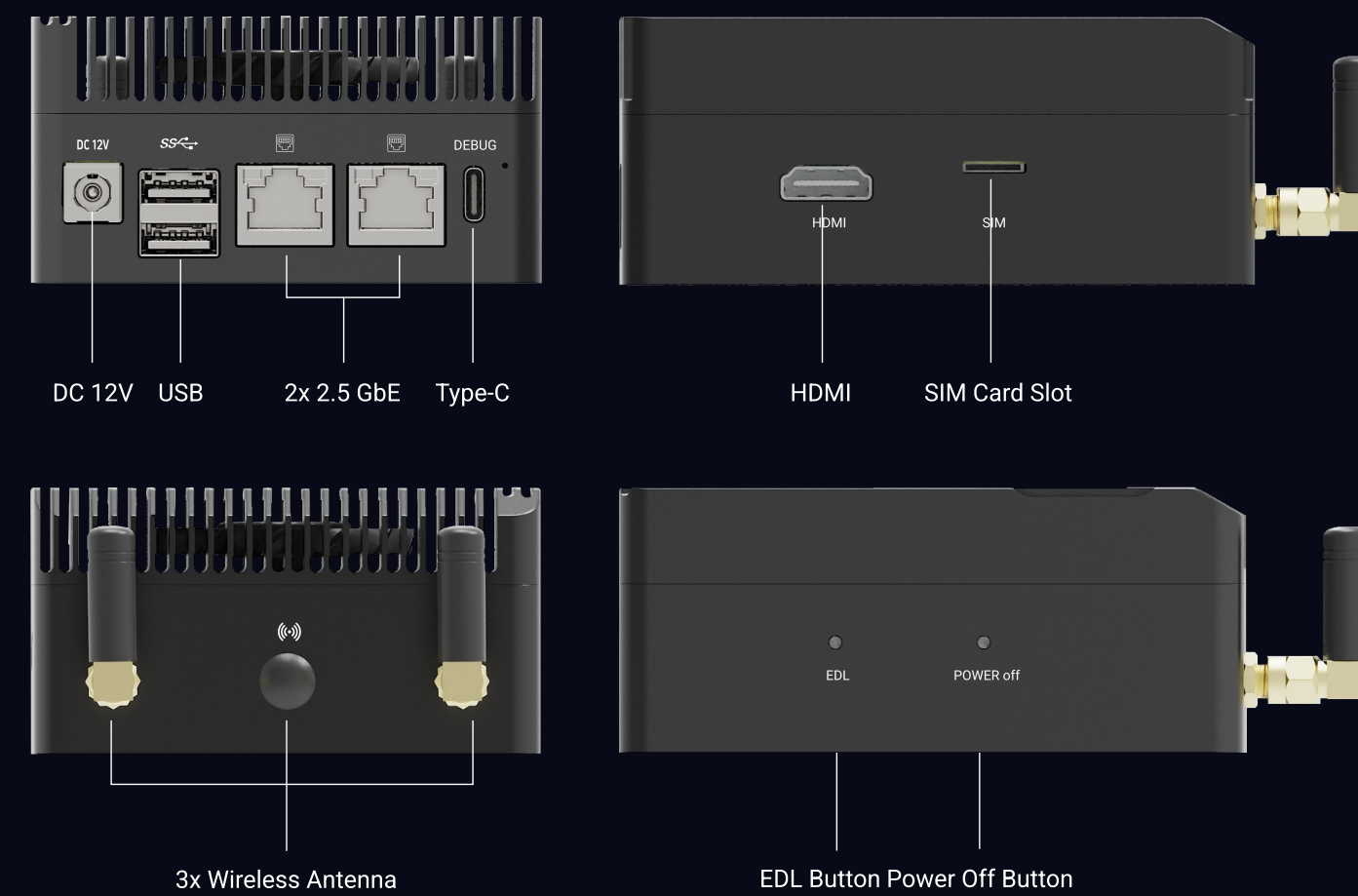
128GB
UFS 3.1 Gear4

M.2 NVMe
PCIe Gen4 SSD

Airbox is

- AI Ready**
Plug-and-play AI deployment
- AI Recognition**
Vision, voice, and data recognized in milliseconds
- AI Reasoning**
On-device inference and intelligent decisions
- AI Realtime**
Millisecond responses, no cloud delay
- AI Reliability**
Industrial-grade reliability, built to endure

Rich Connectivity & Expansion



Cooling & Power

All-metal chassis with PWM-controlled fan interface and 12 V/5.4 A DC power input.

HDMI 2.0

4K@60fps display output

USB

USB 3.1 Host + USB 3.1 OTG Type-A, plus Type-C serial console port

Storage

128 GB UFS 3.1 onboard; M.2 M-Key 2230 for NVMe SSD expansion

Networking

Dual RJ-45 2.5 GbE ports with TSN support

Wireless/Cellular

Mini PCIe slot (half/full size) supports Wi-Fi 6 or 4G/5G modules—plus a built-in SIM slot

Maintenance

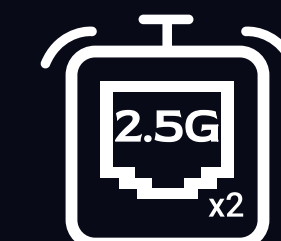
Includes EDL (Emergency Download) button for flashing/reset



WiFi 6



BT 5



Ethernet



5G Network



Full-Metal

Ready to Compute Out of the Box

Support Yocto or Ubuntu and ready for AI deployment.

Ubuntu

casaos

yocto

Extensive Compatibility

Compatible with major frameworks: TensorFlow, PyTorch, ONNX, Paddle, Caffe, DarkNet, etc.

TensorFlow

PyTorch

ONNX

Caffe

飞桨

Low-latency Performance

e.g., LLaMA-7B inference in ~0.6 s first token and ~12 tokens/s thereafter.

Stable-Diffusion

FFNet

Whisper

DeepLabV3

Llama

Depth-Anything

Qwen

MobileNet

EfficientNet

GoogLeNet

Ideal For Local AI APPs

Run popular open-source apps with ease

FRIGATE

DeepSpeech

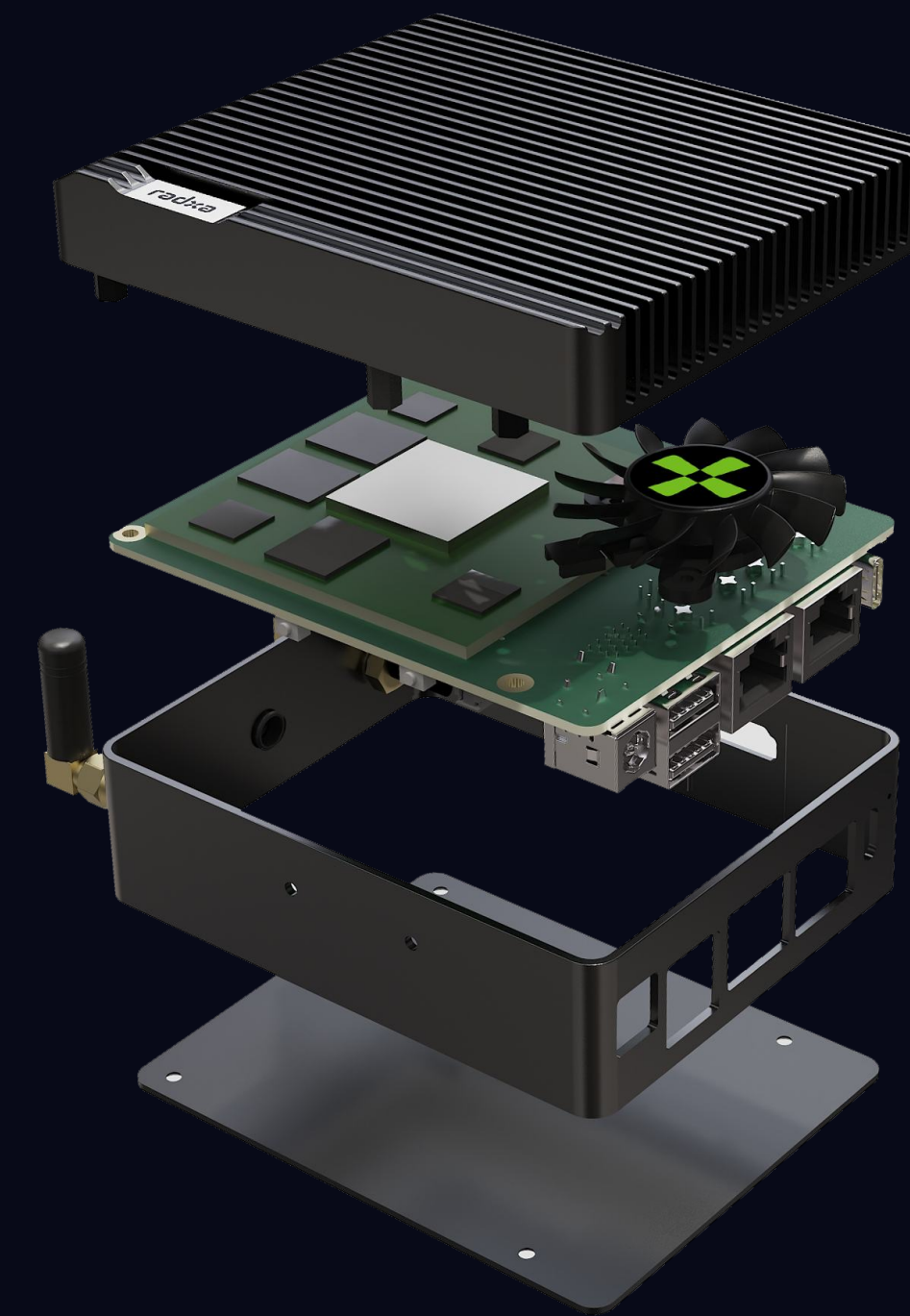
Real-ESRGAN
PRACTICAL RESTORATION

Stable Diffusion

Open WebUI

Ollama

Durability Metal housing, active cooling, industrial reliability



Full Metal

Active Cooling

Compact Size



Edge AI, Ready from Day One