

ADVANTECH ARK-3534

with Axelera Edge 130p



NEVIS



Security



Industry 4.0



Retail



Mobility



Logistics



Robotics



Medical



Hospitality



Utilities

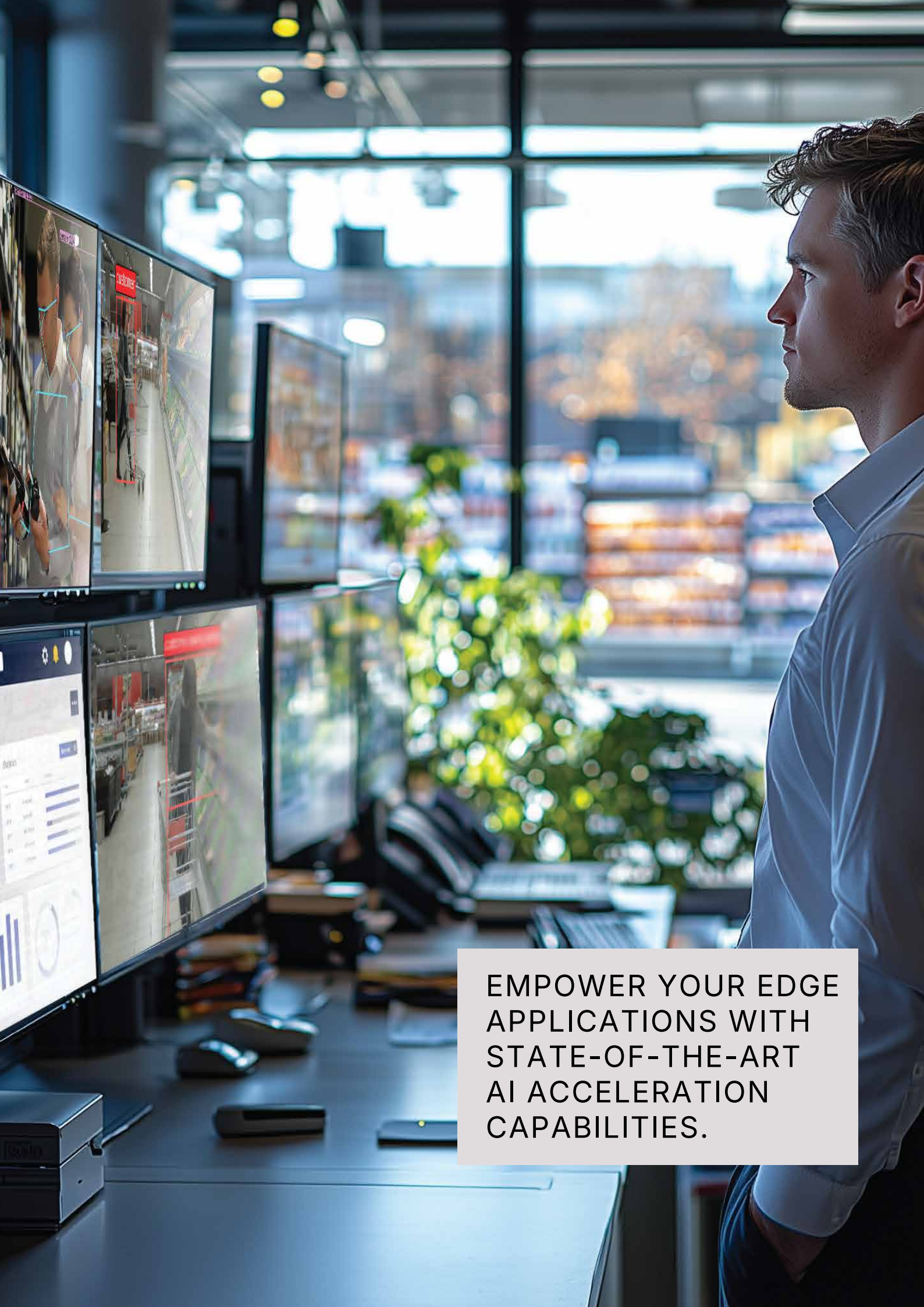


Agritech



AXELERA[®]
ARTIFICIAL INTELLIGENCE

ADVANTECH



EMPOWER YOUR EDGE
APPLICATIONS WITH
STATE-OF-THE-ART
AI ACCELERATION
CAPABILITIES.

KEY FEATURES

- Supported by Advantech ARK-3534, designed for advanced computing and automation including for industrial environments.
- Inference at the Edge powered by Metis® AIPU delivering 214 TOPS/s with industry-leading energy-efficiency.
- Everything you need to get started in a single box pre-integrated Axelera Edge 130p and pre-installed Voyager® SDK.
- A simplified Edge AI inference - run a preloaded neural network in under ten minutes and leverage industry-standard API combined with built-in evaluation tools.
- A growing Model Zoo with a broad coverage of computer vision networks for object detection, image classification, pose estimation and segmentation.

KEY TECHNICAL SPECIFICATIONS

Metis	1x Axelera Edge 130p with 4 GB memory
Processor	Intel Core i5 12500E
Memory	16 GB DDR5 5600 MT/s, SO
Storage	M.2 512 GB
Operating System	Linux Ubuntu 22.04
PCIe	Half-height Gen 3 PCIe x4 slot (Metis AIPU PCIe card)

For more information, please visit or scan the QR code:
<https://www.advantech.com/en/products/>



PROVEN IN KEY MARKETS

Companies in multiple market segments have already adopted Axelera Edge-based AI acceleration for different applications such as:



Security: reducing the time to detect and resolve incidents (abandoned baggage, intrusion, fall) thanks to high resolution, high throughput processing of tens of camera feeds.



Industry 4.0: improve accuracy and speed in defect detection and quality control. Increase worker safety with automated PPE control.



Retail: improve operational efficiency and customer experience with in-store customer behavior, stock monitoring and automated checkout systems.



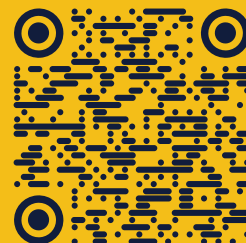
Mobility: real time vehicle detection, identification and tracking from multiple traffic cameras to enhances tolling, enforcement and parking management.



Logistics: monitor the movement of goods and personnel to improve operational efficiency and safety of logistic centers by improving resource allocation and detecting safety hazards.

PERFORMANCE BENCHMARKS

See how Axelera Edge 130p Card compares against other AI accelerators across real-world neural networks. Scan to view the latest results.



<https://axelera.ai/axelera-products/edge#130p>

EASY TO INTEGRATE

Axelera® technology integrates seamlessly with host CPUs based on both x86 and ARM architectures. Our team actively tests different systems from vendors making it easy for embedded developers to prototype AI applications.

VOYAGER SDK

Thanks to Voyager® Software Development Kit (SDK), users have a simple software integration path for AI inference at the edge:

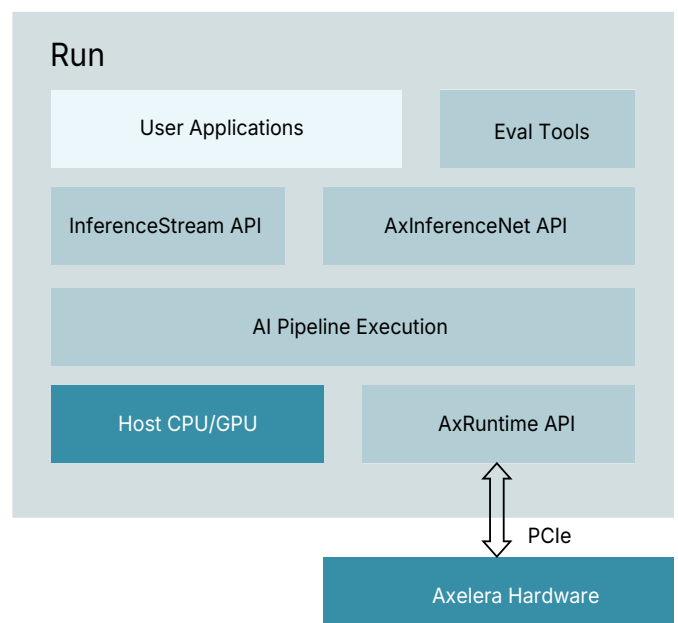
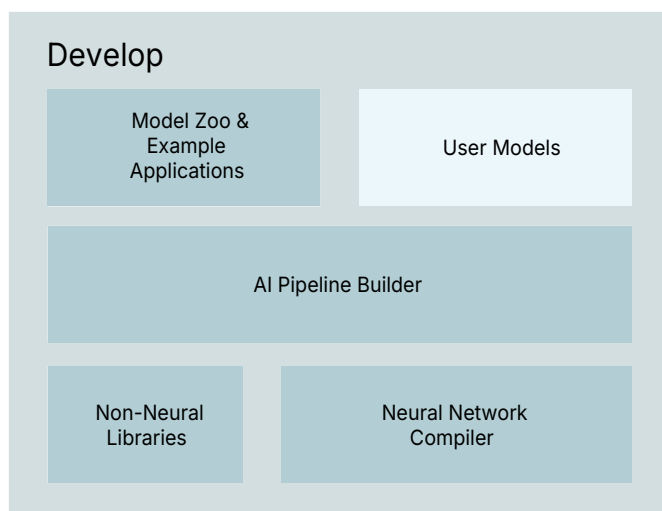
- **Great out-of-the-box experience:** The SDK's built-in tools and models allows evaluating Axelera performance, accuracy and power consumption in a few minutes.
- **Fast end-to-end integration path:** The SDK provides a high-level pipeline description framework that allows building optimized end-to-end AI applications with custom inputs, datasets, models and business logic with very few lines of code.
- **Low-level knobs and APIs:** For users that have their own pipelines and software infrastructure, the SDK includes low-level APIs to directly control the inference hardware.

Voyager is a simple yet feature rich SDK:

- Large [Model Zoo](#) supporting, among others:
 - Small Language Models (Phi3-mini, Llama-3.1 8B etc.)
 - Image Classification (EfficientNet, ResNet etc.)
 - Object Detection (YOLO models, RetinaFace etc.)
 - Semantic Segmentation (U-Net FCN)
 - Instance Segmentation (YOLO models)
 - Keypoint Detection (YOLO models)
- Compiler support for models from Pytorch and ONNX. The compiler automatically manages quantization and graph optimization without user intervention and achieves optimal performance and accuracy.



- Libraries including all pre- and post-processing required to run end-to-end pipelines: scaling; cropping; normalization; format conversion; non-maximal suppression (NMS) and more.
- A YAML description file is used to automatically generate the AI pipelines. The pipeline can then be run as a plugin to GStreamer or within an inference server.
- Built-in tools to test accuracy and performance of models running on Metis AIPU.



SCAN ME

Ordering information

To order the Advantech ARK-3534 with Axelera Edge 130p, please visit: store.axelera.ai/products

Part numbers are listed in the product datasheet.

