

NoctuaEdge™ – Full Technical Specifications

Complete edge AI deployment specifications for enterprise and industrial applications.

1. Hardware & Deployment Targets

- **Architecture Support:** ARM Cortex-A/R/M, x86_64, RISC-V
- **Compatible Devices:** NVIDIA Jetson (Nano, Xavier, Orin), Intel Movidius VPU, AMD Ryzen Embedded, custom ASIC/FPGA
- **Memory Footprint:** Minimum 512MB RAM (Recommended: 2GB+)
- **Storage Requirements:** Minimum 256MB installation; SD, eMMC, SSD supported
- **Edge Cluster Support:** Native support for Kubernetes, Docker Swarm, and standalone orchestrators
- **Deployment Modes:** Embedded module, microcontroller, industrial gateways, ruggedized edge clusters

2. Supported Frameworks & Runtime

- **Inference Engines:** ONNX Runtime, TensorFlow Lite, PyTorch Mobile, TVM, OpenVINO
- **Model Optimization:** Quantization-aware training (QAT), Post-training quantization, pruning, knowledge distillation
- **Operating System Compatibility:**
 - Linux (Ubuntu Core, Yocto, Debian, Alpine)
 - Android Things
 - Real-Time OS (FreeRTOS, Zephyr)
- **Edge Orchestration Integration:** K3s, Azure IoT Edge, AWS Greengrass, BalenaOS
- **Languages Supported:** Python 3.7+, C++, Rust, Go (via SDK)

3. Communication Protocols

- **Messaging Standards:** MQTT 5.0, AMQP, Kafka, ZeroMQ
- **Industrial Interfacing:** OPC UA, Modbus TCP/IP, CAN Bus, EtherCAT
- **Web & API Interfaces:** REST, GraphQL, WebSockets, gRPC
- **Remote Management:** WebRTC P2P, secure OpenSSH tunneling, edge fleet CLI toolkit

4. Security & Privacy

- **Model Protection:** AES-256 encrypted model files, runtime encryption & decryption
- **Root of Trust:** TPM 2.0, Secure Boot, ARM TrustZone, Intel SGX
- **Federated Learning:** Flower, PySyft, TensorFlow Federated – local training with privacy-preserving protocols
- **Data Privacy Features:**
 - Differential privacy libraries
 - Synthetic data injection support
 - On-device anonymization APIs
- **Audit Logging:** Secure append-only logs using local hash trees or decentralized ledger (blockchain optional)

5. Performance Benchmarks

- **Inference Latency:**
 - <10ms on NVIDIA Orin for real-time CV models
 - <30ms on Raspberry Pi 5 with quantized NLP models
- **Throughput Capacity:**
 - 60+ FPS image analysis on Xavier
 - 3K+ TPS for text classification
- **Power Consumption:**
 - 1.8W (microcontroller-class inference)
 - Up to 15W (high-performance edge AI cluster)
- **Failover & Redundancy:**
 - <100ms switch-over time with mesh-aware orchestration
 - Local cache fail-safe & inference continuity

6. Integration & DevOps

- **DevOps Toolchains:** Jenkins, GitHub Actions, GitLab CI/CD, Azure DevOps
- **Container Support:** Docker, Podman, OCI runtime-compliant images
- **Monitoring Integrations:** Prometheus, Grafana, Datadog, edge syslog pipeline
- **Provisioning & Updates:**
 - Secure OTA updates via Mender, BalenaCloud
 - Version rollback & dependency snapshotting

7. Compliance & Standards

- **Regulatory Compliance:** GDPR, HIPAA, ISO/IEC 27001, ISO/IEC 24029 (AI Trustworthiness)
- **Explainability & Bias Tools:** SHAP, LIME, integrated bias diagnostics module
- **AI Lifecycle Management:**
 - Built-in model registry
 - Data lineage tracking
 - Edge-to-cloud pipeline version sync

Contact & Support

For tailored specifications, industrial integrations, or regional deployment packages:



edge@noctua.solutions



www.noctua.solutions/noctuaedge