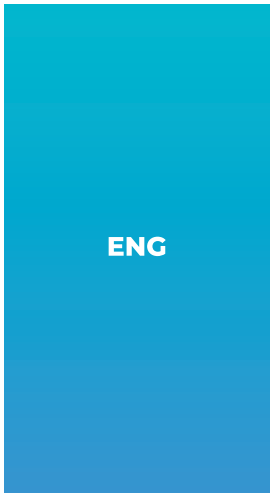


NEUROCLE

DEEP LEARNING VISION SOFTWARE

MAKING
DEEP LEARNING
VISION TECHNOLOGY
MORE ACCESSIBLE



ENG

The next level of deep learning vision inspection

Our Company

With the vision of ‘Making deep learning vision technology more accessible’, Neurocle is developing an innovative vision inspection software ecosystem with its advanced technology and strong R&D team.

Domestic Awards



Selected for the Emerging AI+X Top 100 for four consecutive years (2023–2026)



Won the AI Korea Grand Prize (Minister of SMEs and Startups Award)



Won the K-Startup Challenge Excellence Award (Minister of Culture, Sports and Tourism Award)



Won the Korea ImpaCT-ech Grand Prize

Global Awards & Certifications



Won the Innovators Awards for five consecutive years (2021–2025)



Selected for the Cowen Startup Challenge Top 10 Finalist



Achieved ISO 27001 Certification



Selected as Gartner Cool Vendor in AI for Computer Vision

Our Technologies

Deep Learning Vision Technology

Our deep learning vision software integrates deep learning and computer vision technologies within the framework of artificial intelligence. The models generated using this technology resemble the functionality of the human brain.

Vision Inspection Solution

Deep learning technology enables vision models to conduct high-quality inspections. Among systems utilizing this technology, Neurocle excels by both creating and implementing these advanced models.



Neurocle machine vision ecosystem

Our Partners

We are proud to collaborate with leading machine builders, and our software is trusted by numerous prestigious global companies.

Semiconductor · Electronics · Battery



Automotive OEMs · Suppliers



Industrial Infrastructure · Defence · Steel



Consumer Products · Healthcare



Neurocle's Coverage

Neurocle provides a total solution for vision inspection in collaboration with top industry partners globally.

39+

Greece, Netherlands, Norway, Germany, Romania, Luxembourg, Lichtenstein, North Macedonia, Bulgaria, Belgium, Cyprus, Serbia, Switzerland, Spain, Albania, Ireland, Italy, Turkey, Portugal, France, Finland, Hungary

United States, Mexico, Brazil, Chile, Colombia, Costa Rica, Peru

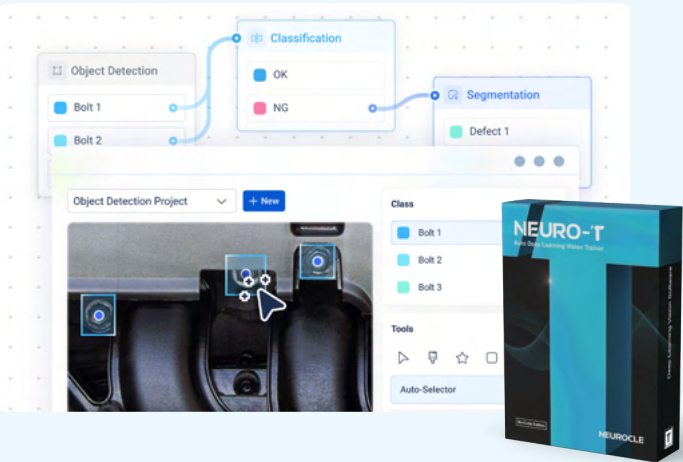
South Korea, Taiwan, Malaysia, Vietnam, Singapore, India, Indonesia, Japan, China, Thailand



6 years of refinement, industry-proven Auto Deep Learning Software

NEURO-T No-Code, GUI-Based Auto Deep Learning Model Trainer

- Creates high-performance models without AI expertise, using an auto optimization algorithm, the Auto Deep Learning
- Supports the entire training process, including image preprocessing, labeling, model training, and evaluation



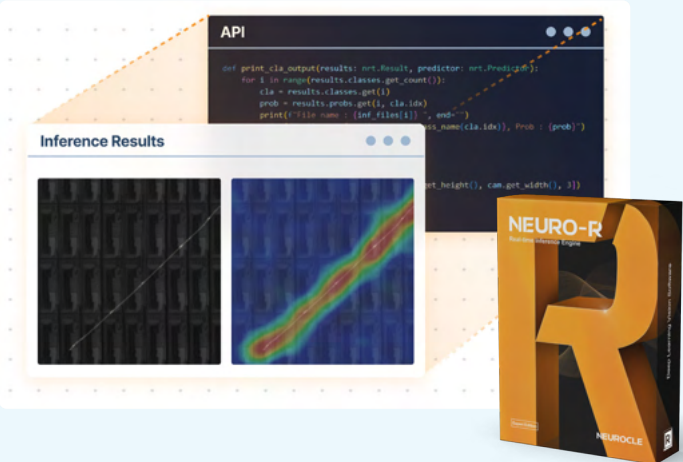
NEURO-T ENGINE Auto Deep Learning Model Training Engine for MLOps Implementation

- Supports API and CLI execution modes, enabling integration of high-performance training modules across diverse systems
- Easily integrates with existing equipment to build optimized, automated training pipelines for on-site environments



NEURO-R Runtime Library for Real-Time Model Deployment

- Provides APIs to deploy models from Neuro-T and Neuro-T Engine to on-site inspection systems
- Inspects images and video from cameras in real time and delivers detection results



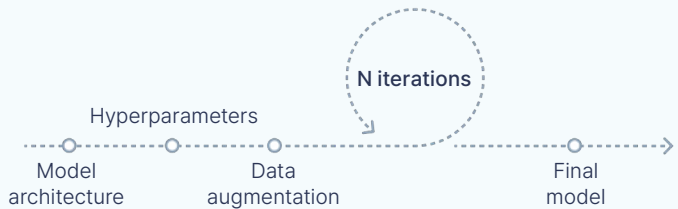
Train high-performance models with Auto Deep Learning Algorithm

Core technology

Our Auto Deep Learning Algorithm automatically optimizes model architectures and hyperparameters to create high-performance vision inspection models.

Limitations

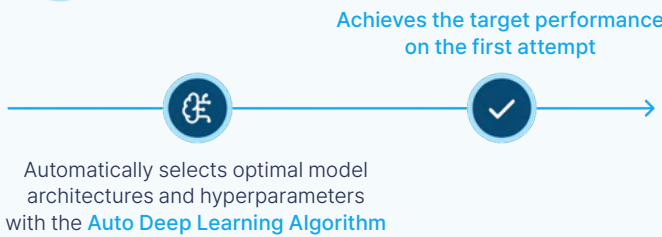
Conventional deep learning requires continuous retraining by skilled engineers to meet the performance standards.



VS

Solutions

Our technology simplifies this process, enabling anyone, regardless of their deep learning expertise, to build a high-performance model with a single click.



Benefits

Achieves high inspection accuracy of initial model

Reduces model development time

Minimizes required development resources

99.9%

4 months → 2 weeks

10 people → 1 person

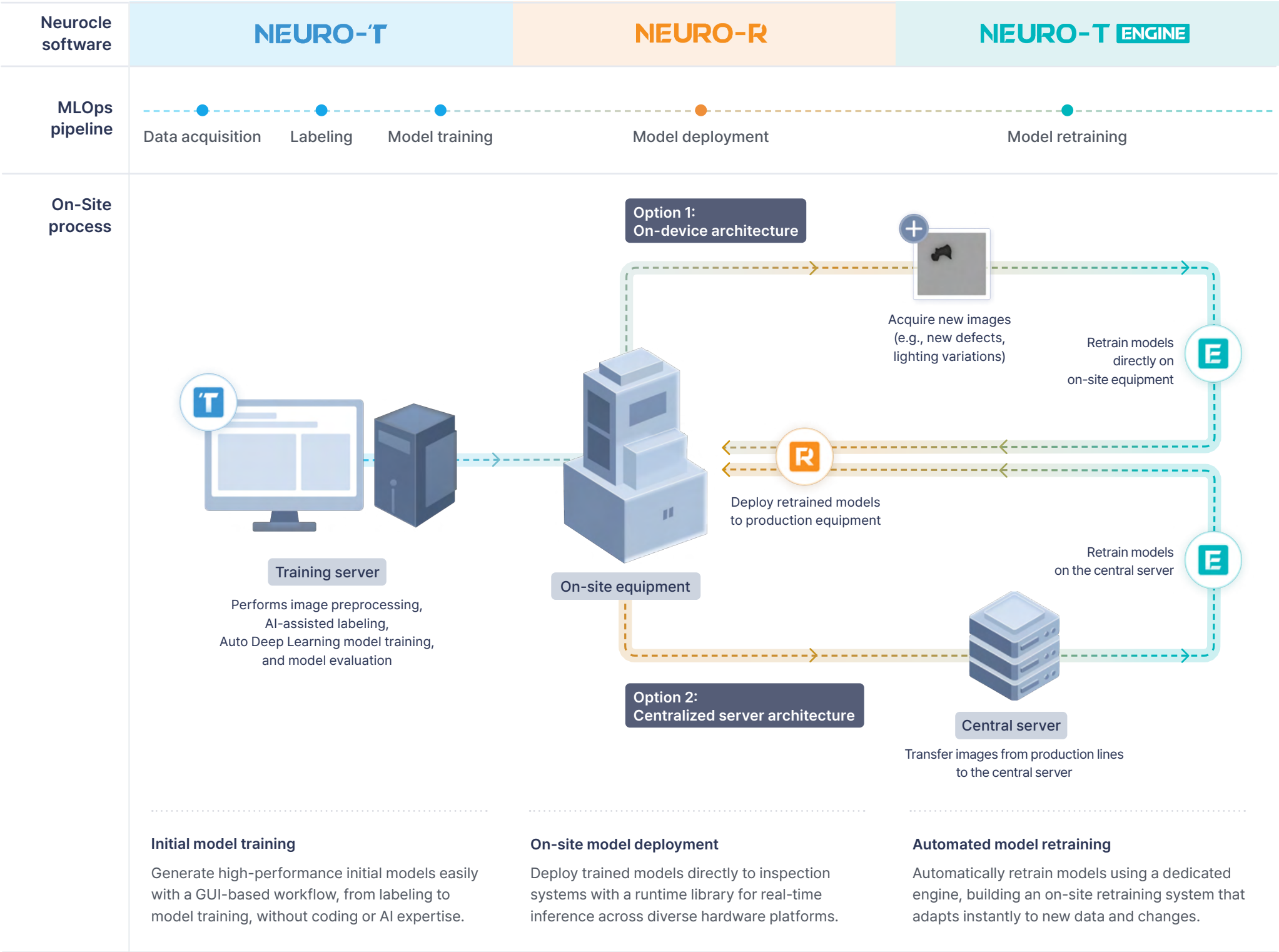
“By leveraging high-performance models trained with Auto Deep Learning algorithms, we were able to **reliably detect microscopic defects, reduce over-inspection**, and improve the overall reliability of the inspection process.”

Vision Engineer | Semiconductor Manufacturer A

“When developing an irregular tire defect inspection system, we generated high-performance models without deep learning experts, **achieving significant cost savings with minimal labor required.**”

Quality Specialist | Tire Manufacturer B

Auto Deep Learning MLOps pipeline:
From initial model training to model maintenance



Advantages of Neurocle MLOps

Fast Implementation
& Flexible Architecture

- 01

Fast implementation with modularized software

Build MLOps quickly using only the required components from model training software, runtime library, and retraining modules.
- 02

Flexible on-device and centralized architectures

Allows flexible architecture design based on deployment environments without fixed MLOps structures.

- On-device architecture: Simultaneously perform real-time inspection and retraining on on-site equipment
 - Centralized server architecture: Transfer images from on-site equipment to a central server for centralized training and model management
- 03

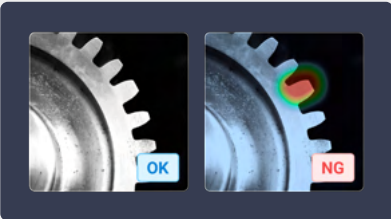
Flexible single and multi-GPU training configuration

Configure GPU resources flexibly based on deployment environments and task requirements.

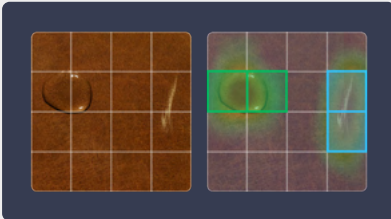
- Train a single task on multiple GPUs
 - Train multiple models in parallel on multiple GPUs
 - Train multiple models sequentially or in parallel on a single GPU

11 types of deep learning vision models

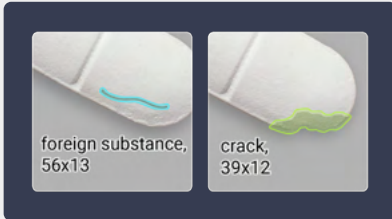
Supervised models



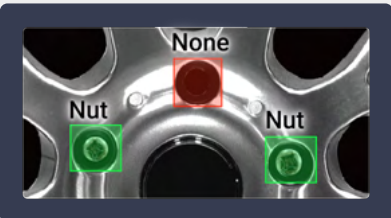
01 Classification
Classifies images to multiple defect categories (an image can contain only one defect)



02 Patch Classification
Classifies high-resolution images by dividing them into small patches (an image can contain multiple defects)



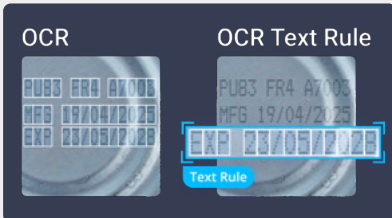
03 Segmentation
Detects the precise shape and location of defects at the pixel level (an image can contain multiple defects)



04 Object Detection
Recognizes the number of objects and identifies their locations within the image



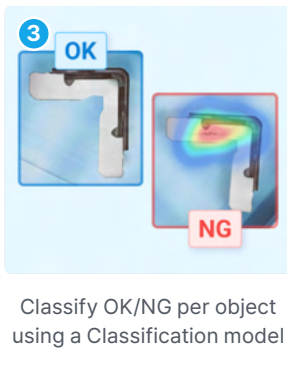
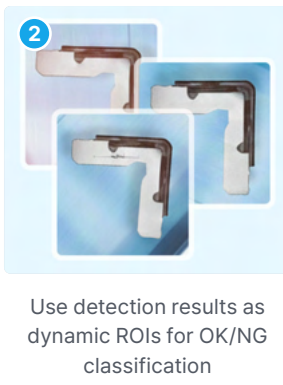
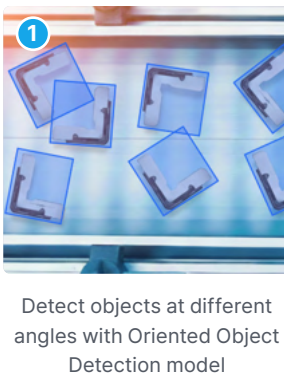
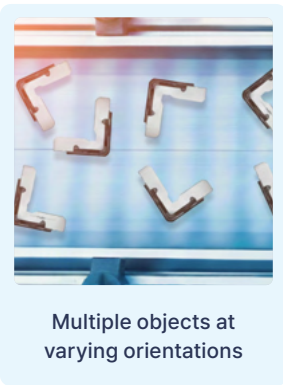
05 Oriented Object Detection
Recognizes the number of objects and identifies their locations, even with varying orientations



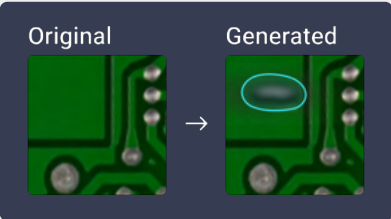
06 OCR / OCR Text Rule
Recognizes text within an image / Applies customized rules to OCR model

Inspection scenarios using multiple deep learning models

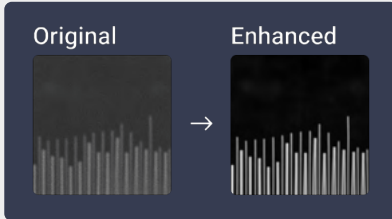
Case 1 Oriented Object Detection + Classification



Generative models



07 GAN (Defect Generator)
Generates artificial defect images that resemble real defects



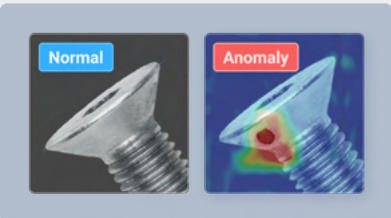
08 Image Enhancement
Converts low-quality images into high-quality images to improve inspection accuracy

Alignment

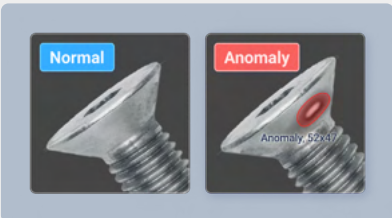


09 Rotation
Automatically rotates original images to the correct orientation

Unsupervised models

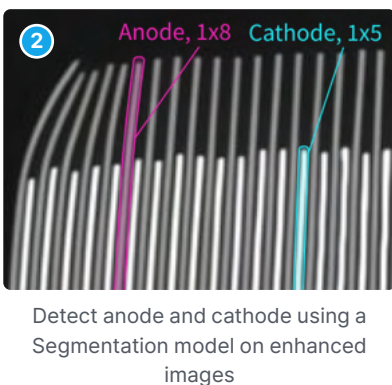
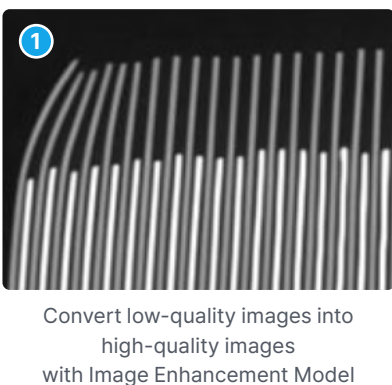
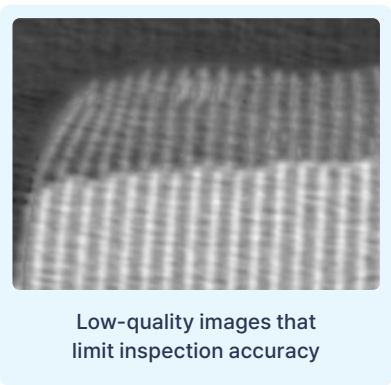


10 Anomaly Classification
Provides binary* classification in the form of a heatmap, trained solely on normal images (*Normal/Anomaly)



11 Anomaly Segmentation
Detects defective areas at the pixel level, trained only on normal images

Case 2 Image Enhancement + Segmentation



Key features of Auto Deep Learning software

NEURO-T

- 01 Exclusive models
- Defect Generator GAN model for synthetic defect generation
 - Patch-based training for high-resolution images
- 02 Single and multi-model workflows
- Single Model Workflow:
 - Quick Learning, Auto Deep Learning for Initial model training
 - Fast Retraining, Transfer Learning for model enhancement
 - Multi-Model Workflow: Flowchart, Inference Center
- 03 AI-assisted labeling tools
- Auto-Labeling, Auto-Selector, Keyword Labeler, Shape Converter

NEURO-T ENGINE

- 01 Dual execution modes(API/CLI) for any integration environment
- 02 Flexible single and multi-GPU configuration for training and inference

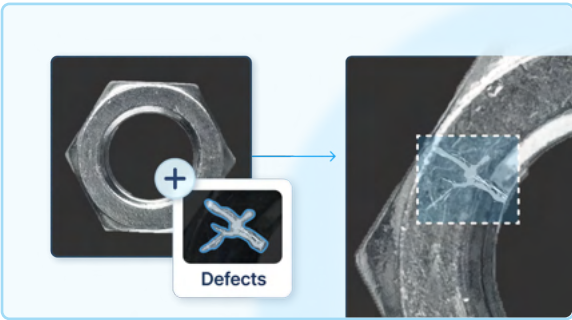
NEURO-R

- 01 Wide range of hardware support
- 02 Fast real-time inspection for inline production
- 03 Flexible implementation and multi-language support
- 04 Reduced development effort with modular APIs

Exclusive models: Generate synthetic data and detect ultra-fine defects

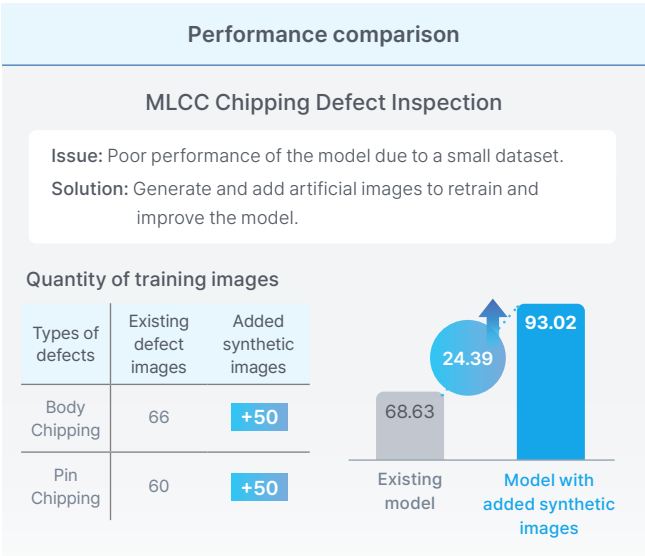
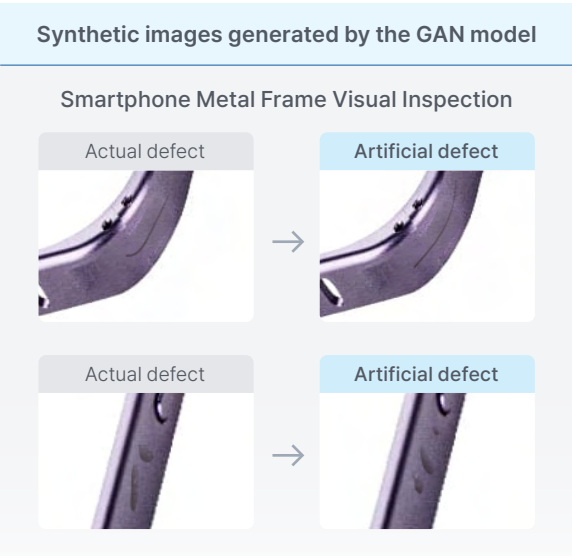
Generates synthetic defect images

GAN Model (Defect Generator)



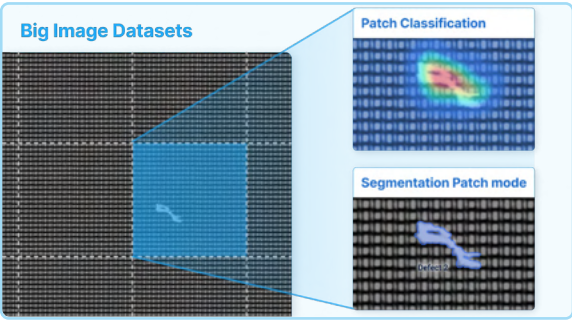
The GAN model addresses the issue of defect data shortage by generating artificial defect images (synthetic data) that closely mimic real defects.

The high similarity between the generated and actual defect images significantly improves the model's accuracy and effectiveness in identifying defects.



Trains on tiny defects without resizing

High-resolution patch-based training



Patch Classification Model

High-resolution images are divided into multiple patches for training. This technique helps to identify and distinguish between normal and defective areas within images.

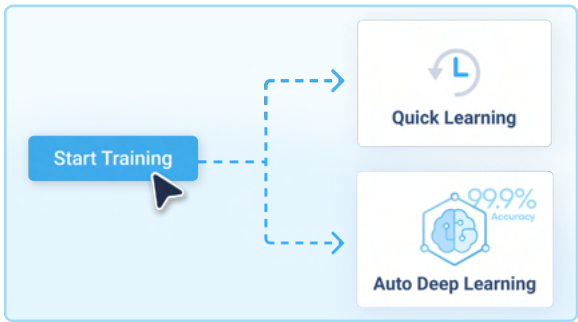
Segmentation Model Patch Mode

During segmentation model training, high-resolution images are split into multiple patches. These techniques minimize missed defects and ensure precise inspection results.

High-performance training from single to multi-model workflows

Maximize model performance on the first attempt

Initial model training



Quick Learning

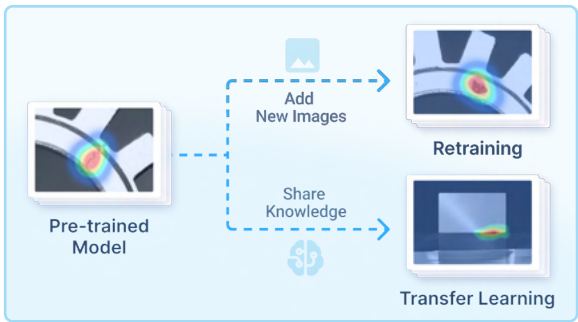
A fast training approach for rapid model training. Quickly assess model performance and reduce trial and error for initial training.

Auto Deep Learning

An automated training approach that finds the optimal model architecture and hyperparameters. Ensures high inspection accuracy with minimal manual intervention.

Adapt to new images or reuse knowledge

Two retraining methods



Fast Retraining

A retraining approach that quickly incorporates new data into existing models. Maintains existing performance while rapidly adapting to new defects.

Transfer Learning

A training approach that reuses pre-trained models for similar inspection tasks. Continuously improves performance and enables reuse across various applications.

Validate in the trainer before runtime

Multi-model design and validation



Flowchart

Connect multiple deep learning models and design complex inspection projects with intuitive flowcharts.

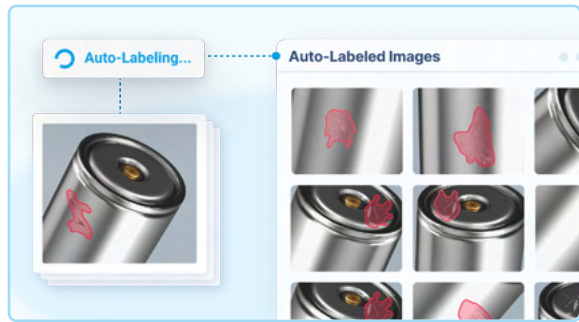
Inference Center

Predict and validate the performance of flowchart models before deployment. This step ensures effective deployment during the proof of concept (POC) phase.

17x faster AI-assisted labeling with clicks, drags, and simple keyword inputs

Automatic labeling feature for large datasets

Auto-Labeling

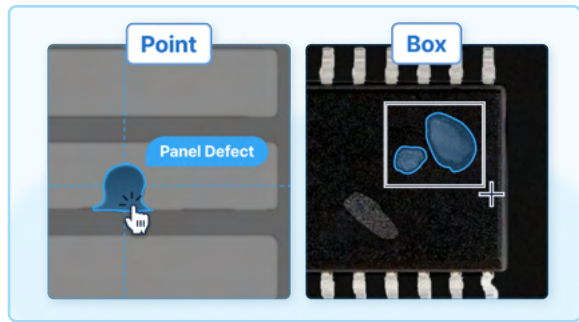


Labels are automatically recommended after the user labels a small number of images. The initial labeled images set the standard, reducing labeling effort while ensuring accuracy.

Available Models	All models in Neuro-T
Labeling Criteria	Pre-defined range by user

Polygon labeling with just one click and drag

Auto-Selector

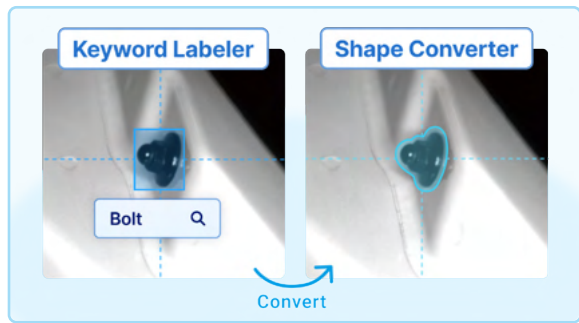


By clicking or dragging to select areas for labeling, our AI detects similar characteristics and identifies related objects. You can easily add or remove labeling areas with a single click, ensuring accurate labeling in the shortest time possible and significantly reducing manual effort.

Available Models	Segmentation, Object Detection, Oriented Object Detection, GAN, Anomaly Segmentation
------------------	--

Labeling by entering simple keywords

Keyword Labeler & Shape Converter



Keyword Labeler

Simply enter keywords, and the AI will label boxes around the areas in the image that match the keywords.

Available Models	Segmentation, Object Detection, Oriented Object Detection, GAN, Anomaly Segmentation
------------------	--

Shape Converter

Convert box labels into shapes that fit the objects with a single click.

Build on-site MLOps with flexible execution modes and GPU configurations

Choose based on system integration environment
API / CLI execution modes



API Mode

Train models on the Neuro-T Engine server via REST API calls. Seamlessly integrate with existing inspection system UIs without modifications. Built-in Task Manager and Queue modules minimize deployment workload.

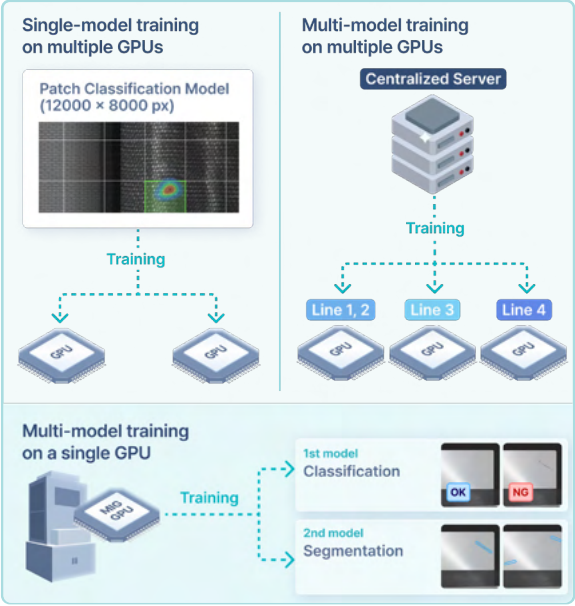
Recommended: Windows-based equipment systems with data transfer and management capabilities.

CLI Mode

Train and test models quickly with a single command. Follow command guidance and monitor training progress and performance metrics in real-time directly in the terminal.

Recommended: Linux-based systems operated primarily through scripts.

Adapt to site conditions and task requirements
Multi-single GPU configuration



Multi-GPU

Single-model training: Handles large datasets or high-resolution images from multiple production lines with improved batch processing compared to a single GPU.

Multi-model training: Efficiently train multiple models simultaneously, either across several production lines on a central server or on a single line with different inspection targets.

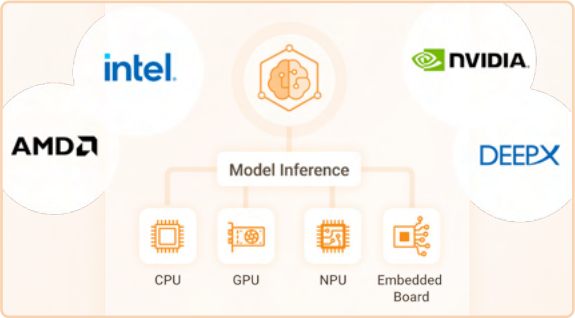
Single-GPU

Multi-model training: Train multiple models on a single on-site device, ideal when several models must be trained simultaneously.

*Utilized in GPU environments with memory partitioning capabilities.

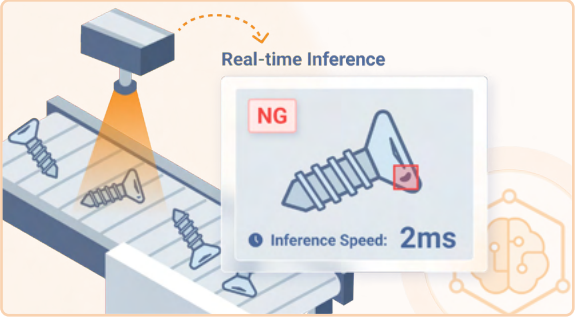
High-speed inspection across diverse hardware environments

Expand processor and hardware options
Broad hardware compatibility



Supports models optimized for various processors and hardware platforms. Models can be deployed immediately without manual conversion or optimization, reducing development effort while expanding platform flexibility.

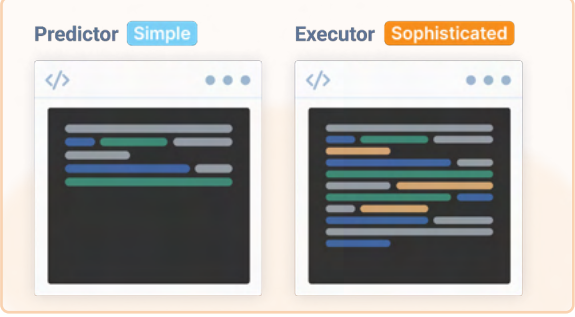
Deploy models inline, in real-time
High-speed inference



Delivers reliable high-speed inspection on any device, with inference speed optimization and model lightweighting features.

Inference speed across different processors				
Model	Image Input Size	GPU	CPU	NPU
Classification	256×256	1.1ms	2.4ms	2.4ms
Segmentation	256×256	2.2ms	4.3ms	4.4ms

Streamline development while retaining integration control
Reduce effort, increase flexibility



Two execution modes:

- Predictor structure for simple model calls.
- Executor structure for finer control and customization, enabling more sophisticated system implementations.

Single API for multi-model execution: Export the final flowchart model designed in Inference Center as a single file and run with just one API.

Multiple programming languages: Supports C++, C#, and Python.

Inspection solution for all applications across different industries

Use cases by industry

Electronics



Accurately detect defects in small electronic components.

Example

- LED panel inspection
- Soldering bubble detection in PCB board
- Surface inspection for FPCB / MLCC
- Speaker mesh inspection

Battery

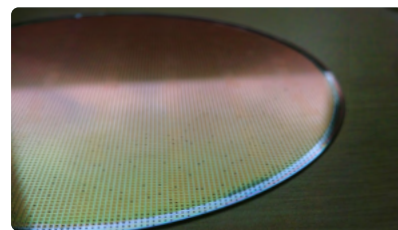


Improve battery yield with precise detection.

Example

- Pouch type battery inspection
- Cylindrical battery cap inspection
- Battery X-ray inspection
- Battery CT inspection

Semiconductor



Accurately recognize irregular defects, replacing visual AOI inspections.

Example

- Wafer notch detection
- Wafer surface inspection
- Wire bond defect inspection
- Integrated circuit leads inspection

Display



Improve display quality through precise panel inspection.

Example

- Bear glass inspection
- OLED mask inspection
- Pattern defect inspection
- Panel damage inspection

Automotive



Enhance production speed which has been automated early on.

Example

- Assembly detection
- Automotive parts and tire inspection
- VIN Number recognition
- Sunroof primer recognition

Metal



Detect machining and welding that are hard to inspect with the naked eye.

Example

- Steel surface inspection
- Press products surface detection
- Steel plates edge detection in rolling process
- Color-coated steel plates inspection

Pharmaceutical



Minimize undetected rates in pharmaceutical industries with strict regulations.

Example

- Pill surface defect inspection
- Medical kit inspection
- Syringe rubber parts defect inspection
- Catheter laser punching detection

Bio & Healthcare



Detect anomalies that are difficult to identify with the naked eye.

Example

- Number of microorganism detection
- Cell detection inside of organs
- Virus infection classification
- Contact lens inspection

Food & Beverage



Provide fresh products to consumers by detecting contaminants.

Example

- Ramen quality classification
- Dumpling number counting
- Tofu contaminant detection
- Beverage foreign object detection

Packaging



Ensure best packaging quality through checking the package condition.

Example

- Glass bottle exterior defect inspection
- Package sealing detection
- Label alignment inspection
- Container printing defect inspection

Logistics



Maximize distribution speed by automating goods sorting and label tracking.

Example

- Container text identification
- Invoice identification and product classification
- Pallet damage classification
- Conveyor belt inspection

Agriculture



Improve yield and quality by establishing personalized management.

Example

- Crop disease infection classification
- Ripeness classification
- Animals abnormalities detection
- Forest damage and fire detection

Choose a lineup that works with you

Trainer Lineups

Essential Compact essential features for vision inspection	Pro All features included for optimal model performance
- Support for Auto Deep Learning Training 6 Different deep learning models* - AI-assisted Labeling	11 Different deep learning models * - Inference Center and Generation Center - Multi-model Export

*Essential: Classification, Segmentation, Object Detection, OCR, Rotation, Anomaly Classification
*Pro: Includes all essential models plus Patch Classification, Oriented Object Detection, GAN, Image Enhancement, Anomaly Segmentation

Product	Lineup	License	Account	GPU(Max)	Parallel task
Neuro-T Neuro-T Engine	Essential	Basic	1	1	1
	Pro	Standard	3	2	2
		Team*	5	4	4

*If you wish to use licenses higher than Team (more than 6 accounts, more than 5 GPUs), please contact us.

Runtime Lineups

Embedded Reliable AI performance for compact embedded systems	NPU High-efficiency inference for low-power environments	PC High-performance configuration leveraging CPUs, GPUs, and iGPUs
Recommended for space-constrained or compact equipment applications	Ideal for power efficiency-sensitive environments	Recommended for environments requiring flexible resource utilization or scalability

Product	Lineup	License	GPUs / NPUs (Max)
Neuro-R	Embedded	Embedded	1
	NPU	Single	1
		Double	2
		Multi	4
	PC	Single	1
		Double	2
		Multi	4

Hardware Requirements

Product	Category			Minimum	Recommended
Neuro-T • Neuro-T Engine	Server ¹	OS		Windows 10 64-bit (Build 17763 or higher), Windows 11	
		CPU		Core : 4 Thread : 6	Single GPU : i7, Ultra 7 Multi GPU : i9, Ultra 9
		RAM		32GB (2x larger than GPU memory)	64GB
		GPU		NVIDIA GeForce RTX 4070 (CUDA Compute Capability: 7.0)	NVIDIA GeForce RTX 5080, GeForce RTX 5090, RTX PRO 5000 Blackwell
	Client	Browser		Chrome (Recommended), Microsoft Edge	
Neuro-R ²	PC	OS		Windows 10 64-bit (Build 17763 or higher), Windows 11 Linux Ubuntu 20.04, 22.04	
		Development ³ Environment		Visual Studio 2019	Visual Studio 2022 or higher
		Processor	CPU	Intel Core i5, Core Ultra 5 AMD Ryzen 5	Intel Core i7, Core Ultra 7 AMD Ryzen 7
			GPU ⁴	NVIDIA GeForce RTX 4060, RTX 2000 Ada AMD Radeon RX 7600 Intel Arc B570	NVIDIA GeForce RTX 5070, RTX PRO 4000 Blackwell AMD Radeon RX 9070 Intel Arc B580, Arc Pro B60
	NPU ⁵		Intel NPU, DEEPX DX-M1 (To be supported)		
	Embedded	Supported Platform	Jetson ⁶	Jetson Orin Series (JetPack 6.2) Jetson Xavier Series (JetPack 5.1.5)	
			LattePanda	LattePanda 3 Delta (Windows 11) LattePanda Mu (Windows 11)	

- 1 Model training is recommended on a desktop for stable performance. Specifications are validated in a desktop environment.
- 2 Supports C++, C#, and Python (versions 3.9 to 3.12) as programming languages.
- 3 Visual Studio requires the v142 toolset and Windows SDK version 10.0.19041.0 or later, and supports builds only for x64 CPU architecture.
- 4 NVIDIA GPUs are supported on products with Compute Capability 7.5 or higher (GeForce RTX 20 Series or above).
- 5 Intel NPUs use OpenVINO as the inference backend, and OpenVINO is supported on Windows only.
- 6 For Jetson Orin Series, models with at least 16GB of unified memory are recommended. Supported JetPack version varies depending on the Neuro-R version.

Operating environment and system architecture

On-premise environment	Client-server architecture
Neuro-T, Neuro-T Engine, and Neuro-R operate entirely in on-premises environments. They run on local servers without the use of external cloud services, ensuring high data security.	Neuro-T supports a client-server architecture. The central server manages the training environment, allowing multiple users to connect simultaneously and perform project creation, model training, and management tasks.

*Licenses are provided using USB hardware dongles.



Robovision
Machine Vision Experts

Main:

1st Kifisias str
56532 Thessaloniki, GR
T: +30 2310672436

contact@robovision.gr

Branch:

81st Platonos str
10441 Athens, GR
T: +30 2105157861

www.robovision.gr