



Advantech Responsible AI Program and Commitments

Guided by our core corporate values and our Environmental, Social, and Governance (ESG) sustainability goals, we apply a rigorous governance framework and take concrete action to ensure that, throughout the development, deployment, and use of AI technologies, Advantech consistently upholds ethical standards, maintains transparency, and proactively protects both users and the environment. The following outlines the core commitments and practical guidelines of our Responsible AI Program:

I. Upholding AI Ethics and Ensuring a High Degree of Transparency

We are committed to building an AI ecosystem that earns the full trust of our customers, partners, and society at large. Transparency is the cornerstone of trust, and we rigorously govern the boundaries within which AI technologies may be applied.

- Restricting Access to and Use of Sensitive AI Functions:

We recognize that certain AI technologies may pose potential threats to personal privacy and human rights. Accordingly, Advantech enforces strict access controls and review mechanisms for highly sensitive AI functions — those whose decisions or outputs could seriously harm or affect an individual's rights, finances, health, or privacy. We are committed to never applying AI technology in ways that infringe upon fundamental human rights, enable unlawful surveillance, or violate societal ethical standards.

- Clearly Labeling AI-Generated Content and Decisions:

To safeguard users' right to be informed and to prevent confusion or misinformation, we provide clear and explicit labeling for AI-generated text, images, audio, and video content delivered to external users, as well as for key recommendations and decisions driven by AI algorithms. This enables users to readily identify the source of such content and understand the role AI plays in it.

II. Pursuing Excellence in Performance and Absolute Fairness

The stability and fairness of our AI systems reflect our commitment to quality. We have established a comprehensive lifecycle management mechanism to ensure that AI models continue to operate efficiently and fairly after deployment.

- Assessing Model Fairness and Eliminating Bias:

We have formed a cross-departmental AI Governance Team to confirm the fairness and bias testing conducted on externally provided AI models. Our current AI vendor's report documents in detail its "Responsible Data and Development" process, which includes rigorous filtering of personally identifiable information (PII) and safety review during pretraining, as well as alignment through Reinforcement Learning from Human Feedback (RLHF) to significantly reduce bias and stereotyping related to gender, race, and political

views. ([Source: Gemma: Open Models Based on Gemini Research and Technology \(Google DeepMind\)](#)) the vendor report is also used as one of the reference bases for our evaluation.

- **Detecting and Correcting AI Model Performance Degradation (Model Drift):**

To ensure the quality and reliability of AI services, Advantech is establishing a dedicated AI Operations team to build a dual-track mechanism combining continuous monitoring and periodic validation. On one hand, we continuously track key indicators such as model response accuracy and user satisfaction to identify performance changes in real time. On the other hand, we conduct regular regression testing through standardized benchmarks to objectively assess whether model performance has degraded, and initiate optimization procedures accordingly to ensure the long-term, stable operation of our AI systems.

- **Periodic Review and Evaluation of AI Models:** Taking into account the risk level and complexity of AI models, Advantech has established an AI Operations team within its project organization to conduct fairness and bias testing on AI models, and to monitor and correct AI model performance degradation. In addition, approved third-party AI tools will be re-reviewed whenever there are material changes to a vendor's terms of service or data processing practices.

III. Protecting User Rights and Strengthening Employee AI Literacy

Responsible AI is not merely a technical matter — it is a concrete expression of a human-centered approach. We empower users to question machine-driven outcomes, and we begin internally, comprehensively strengthening employees' AI literacy.

- **Establishing an Appeal and Remedy Mechanism for AI Decisions or Outputs:** We allow users or affected third parties to challenge AI decisions or outputs. The company provides an audit appeal and incident reporting channel; appeals or reports may be submitted by email to AI.Governance@advantech.com. All appeals are reviewed and finally adjudicated by the AI Governance Team.

- **Company-wide AI Ethics and Security Training:**

We have incorporated “AI Ethics and Security” into the mandatory training curriculum for all employees. Through regular training, we ensure that developers and all colleagues have a thorough understanding of the potential cybersecurity risks and ethical boundaries of AI. In the first half of 2026, Advantech completed company-wide AI security training for all employees in the China region, requiring them to apply the principles of responsible AI use in their daily work.

IV. Advancing Green AI and Environmental Sustainability

As demand for AI computing grows exponentially, we remain mindful of the burden this places on the global environment. Advantech is committed to closely integrating AI development with our climate change response strategy.

- **Reducing the Carbon Footprint of AI Infrastructure and Models:**

We work to minimize the carbon footprint of our AI infrastructure by optimizing model architectures to reduce computational energy consumption and by improving the power usage effectiveness (PUE) of our server facilities. We also hold our cloud service and

hardware suppliers to the same environmental standards. Our current AI vendor's report states its commitment to achieving net-zero emissions across all of its operations and value chain by 2030, and to deploying 24/7 carbon-free energy (CFE) at its data centers. Between 2024 and 2025, the vendor reduced the direct operational carbon emissions (Scope 1 & 2) of its infrastructure by approximately 12%. (Source: [Google Reduces Data Center Emissions](#))

- Quantifying the Impact of AI Initiatives on Sustainable Development (ESG):

Beyond reducing AI's own carbon emissions, we are also actively applying AI technology to address sustainability challenges. We plan to establish a quantitative evaluation framework to ensure that the benefits of AI can be objectively verified, and to regularly track and disclose the tangible contribution of each AI initiative to sustainability metrics such as energy efficiency, waste reduction, and supply chain optimization. We aim for AI to become Advantech's strongest driver in achieving our Net Zero goal.

- Advantech Case Study: The LEO-L50 outdoor asset management device received dual recognition at the 2025 iF Design Award and the Taiwan Excellence Award, helping a shipping company usher in a new era of sustainable asset tracking. Through real-time tracking and AI-driven big-data route planning and management, the solution enables precise control over mobile assets. The LEO-L50's rugged, reliable design (IP67; -30°C to 80°C; C1D2 certified) and smart connectivity (GNSS, Cat.1) improved theft prevention by 10% and increased asset utilization by 25%.
- Advantech Case Study: Smart Strawberry Growing — AI-Driven Automation Transformation in Japanese Greenhouse Agriculture. By monitoring and automatically controlling the greenhouse environment, the solution improves the operation of greenhouse equipment and reduces carbon dioxide emissions. Concrete results include reduced heavy-oil consumption for heaters and CO2 supply equipment, with an estimated annual reduction of approximately 20% in heavy-oil use. Through remote control of the greenhouse environment and equipment, along with detection of trespassers and pests, the solution also reduces manual labor intensity — lowering it by approximately 60% compared with the period before the system was implemented.