

OECD AI Transparency Report

Organization: Milestone (DNK)

Reporting Period: Q4 2025

Published: October 15, 2025

Section 1 - Risk identification and evaluation

a. How does your organization define and/or classify different types of risks related to AI, such as unreasonable risks?

AI risks are classified based on their potential impact on security, safety, societal impacts, and human rights and as defined by the EU AI Act, focusing first on High Risk AI that could affect humans and introduce an unreasonable risk. The risk assessment and mitigating actions are documented in our AI Governance and GRC Tool.

b. What practices does your organization use to identify and evaluate risks such as vulnerabilities, incidents, emerging risks and misuse, throughout the AI lifecycle?

Milestone has put in place processes to identify possible known and reasonably foreseeable risks to health, safety, and fundamental rights that could follow from the use of relevant AI systems throughout their lifecycle. We have developed policies to ensure high-quality training, validation, and testing datasets for relevant AI systems. Milestone has defined processes, defined in our Responsible Development Policy, to employ and document comprehensive risk and impact assessments from ideation through development, deployment, and usage. This includes assessing potential vulnerabilities, emerging risks, misuse, considering human rights and ethical impacts. Milestone is also a member of the AI Pact and already now start taking High Risk AI requirements from the EU AI Act into account. Besides the initial risk assessments throughout the development process incl. foreseeable risks, Milestone also plan to continuously perform diverse testing measures to monitor the performance of the AI throughout the AI lifecycle. The ability to identify, monitor and evaluate all the performance and risks throughout the AI lifecycle is limited by the extend that Milestone control the deployment and runtime environment. For environments outside Milestones control, we are depending on the HRDD process and the assessments and reporting performed on these environments by our partners and customers.

c. Describe how your organization conducts testing (e.g., red-teaming) to evaluate the model's/system's fitness for moving beyond the development stage?

As part of our planned MLOps process and CI/CD toolchain, our organization conducts testing measures to evaluate the system's fitness for deployment. These tests assess the accuracy, safety, security, and resilience of the systems against various metrics. As an AI Pact member, Milestone already plan to use accuracy metrics from the EU AI Act to ensure early conformity.

d. Does your organization use incident reports, including reports shared by other organizations, to help identify risks?

Yes

e. Are quantitative and/or qualitative risk evaluation metrics used and if yes, with what caveats? Does your organization make vulnerability and incident reporting mechanisms accessible to a diverse set of stakeholders? Does your organization have incentive programs for the responsible disclosure of risks, incidents and vulnerabilities?

Yes, qualitative metrics are used, with caveats such as the complexity of accurately predicting risks. Through our AI Governance and GRC Tool we have accessible reporting mechanisms for various stakeholders. We do not have an incentive program for responsible disclosure.

Incidents are reported through our Whistleblower and AI incident reporting webpage and the reports are used to identify new risks when they get reported. The Incident Reports are automatically generated in our Incident Management System and the reports are automatically assigned a dedicated AI Incident Manager. Milestone does not yet use incident reports shared by other organisations, but instead the latest news about AI risks are used.

f. Is external independent expertise leveraged for the identification, assessment, and evaluation of risks and if yes, how? Does your organization have mechanisms to receive reports of risks, incidents or vulnerabilities by third parties?

No, we do not leverage external independent expertise for risk identification and evaluation. We have an internal task force with deep understanding and experience in AI for the identification, assessment, and evaluation of risks. We also have mechanisms in place to receive reports of risks, incidents, and vulnerabilities from third parties, both through due diligence processes and assessments, but also through our whistleblower function.

g. Does your organization contribute to the development of and/or use international technical standards or best practices for the identification, assessment, and evaluation of risks?

Yes, Milestone actively develop and use international standards as active members of Working Group 3 (WG3) under CEN-CENELEC JTC 21, that focuses on engineering aspects of creating standards for the EU AI Act. This also gives us early insights into the EU AI Act requirements on risk management for AI systems as developed in Working Group 2, with a clear and actionable guidance on how risk can be addressed and mitigated throughout the entire lifecycle of the AI

system. These requirements build the coming EU AI Act standards from CEN-CENELEC that again build on other standards like **EN IEC 1010:2019** *Risk management - Risk assessment techniques*, **ISO/IEC TS 4213:2022** *Information technology — Artificial intelligence — Assessment of machine learning classification performance*, **ISO/IEC Guide 51:2014** *Safety aspects — Guidelines for their inclusion in standards*, **EN ISO/IEC 22989:2023** *Information technology - Artificial intelligence - Artificial intelligence concepts and terminology (ISO/IEC 22989:2022)*, **EN ISO 9000:2015** *Quality management systems - Fundamentals and vocabulary (ISO 9000:2015)*.

h. How does your organization collaborate with relevant stakeholders across sectors to assess and adopt risk mitigation measures to address risks, in particular systemic risks?

Milestone is directly involved in developing standards as described earlier, collaborating across the sector to assess and adopt risk mitigating measures. Milestone also help to develop the General-Purpose AI Code of Practice led by the EU AI Office. The General-Purpose AI Code of Practice will detail the AI Act rules for providers of general-purpose AI models and general-purpose AI models with systemic risks, including requirements for downstream providers. As part of this work Milestone collaborate with relevant stakeholders across sectors to assess and adopt risk mitigation measures to address risks, in particular systemic risks of General-Purpose AI involving nearly 1000 stakeholders, as well as EU Member States representatives, European and international observers.

Any further comments and for implementation documentation

No answer provided

Section 2 - Risk management and information security

a. What steps does your organization take to address risks and vulnerabilities across the AI lifecycle?

Milestone has implemented a Responsible Development Policy that specify roles and responsibilities for developing AI and addressing the risks and vulnerabilities. We have also implemented an AI Governance, Risk and Compliance tool to effectively govern data and AI. This platform help Milestone to establish clear data and AI policies, ensuring full visibility into our data products, AI models, metadata, and risk assessments. We gain a comprehensive understanding of our AI footprint and document how to mitigate AI risks, unlock AI-ready data for optimized performance and to develop and govern out AI/ML projects with clearly defined controls throughout the AI lifecycle.

b. How do testing measures inform actions to address identified risks?

Testing measures inform actions by identifying vulnerabilities and risks, which are then addressed through mitigation strategies and controls implemented during the development and deployment phases using our AI Governance platform.

c. When does testing take place in secure environments, if at all, and if it does, how?

Testing takes place in secure environments, following our Secure Software Development Lifecycle policy, during development and before deployment. Secure environments are used to ensure that models and systems are resilient to various security threats and vulnerabilities.

d. How does your organization promote data quality and mitigate risks of harmful bias, including in training and data collection processes?

Our organization has established comprehensive data governance processes, including documenting dataset properties, ensuring data quality, and mitigating unwanted biases. This ensures that the data used in AI systems is high-quality, relevant, and unbiased. We promote data quality and mitigate risks of harmful bias through transparency, privacy-preserving training techniques, rigorous testing, and fine-tuning data collection processes. We ensure that data collection methods are inclusive and representative of all relevant subgroups. This helps in improving the overall quality of data and reduces the risk of sampling bias. Through our Responsible Technology Program and Responsible Development Policy, we incorporate ethical considerations and legal compliance (GDPR, AI Act) into the design and deployment of AI systems. We also monitor and evaluate model performance to identify and address biases as they emerge. Our Responsible Development Policy also enhance the transparency and explainability of AI decision-making processes. This involves identifying and anonymizing personal or sensitive features in datasets and ensuring that the cause-and-effect relationships within models are clear and understandable. In a new project Milestone also acts as a trusted librarian of AI video data, carefully curating, tagging, and delivering ethically sourced, regulation-ready video data for AI model training. By leveraging Project Hafnia, AI learns from high-quality, verified video data, ensuring not only precision and compliance but also the protection of citizen privacy.

e. How does your organization protect intellectual property, including copyright-protected content?

Our organization protects intellectual property through AI literacy, robust AI policies and security controls and compliance with legal frameworks, ensuring that IP is not shared with third-party tools and that copyright-protected content is not used.

f. How does your organization protect privacy? How does your organization guard against systems divulging confidential or sensitive data?

We protect privacy by implementing privacy-preserving training techniques, rigorous testing to prevent systems from divulging confidential or sensitive data, and compliance with applicable privacy laws.

g. How does your organization implement AI-specific information security practices pertaining to operational and cyber/physical security?
ul
lii. How does your organization assess cybersecurity risks and implement policies to enhance the cybersecurity of advanced AI systems?
liii. How does your organization protect against security risks the most valuable IP and trade secrets, for example by limiting access to proprietary and unreleased model weights? What measures are in place to ensure the storage of and work with model weights, algorithms, servers, datasets, or other relevant elements are managed in an appropriately secure environment, with limited access controls in place?
liiii. What is your organization's vulnerability management process? Does your organization take actions to address identified risks and vulnerabilities, including in collaboration with other stakeholders?
liiv. How often are security measures reviewed?
liv. Does your organization have an insider threat detection program?
ul

Milestone assess cybersecurity risks through a structured approach, conducting risk analyses to identify vulnerabilities in AI systems, including adversarial threats, accidental errors, and environmental risks and implementing policies such as zero-trust architecture, encryption, and multi-factor authentication to mitigate risks.

Milestone already follow the requirements outlined in Executive Order 14028 and NIST standards for cybersecurity, NIST Secure Software Development Framework (SSDF) and we are preparing for EU NIS2.

Milestone on-prem products are designed to align with FIPS 140-2 standards, using Microsoft's Cryptography API, Next Generation (CNG) for secure cryptographic functions. This means they meet the high security standards required by government and regulated industries in the US and Canada, making them a reliable choice for organizations aiming for FIPS compliance.

i. Assessing Cybersecurity Risks and Enhancing AI System Security

Milestone implements a comprehensive Security Development Lifecycle (SDL) that includes practices such as threat modeling, secure coding, and vulnerability testing, which can be applied to AI systems to ensure operational and cyber/physical security.

https://doc.milestonesys.com/cybersecurity/pdf/en-US/Milestone_SecurityDevelopmentLifecycle_en-US.pdf

Milestone's toolchain ensures that:

- Everyone involved with all aspects of product development in Milestone adheres to our source control procedures and only uses approved tools.

- A registry of third-party libraries used in our products will be maintained, as well as a blacklist of libraries not allowed.
- Automated vulnerability and version update scanning of all libraries used are executed and reported for each release.
- A SBOM is generated, signed and stored alongside the software product for each version, documenting all components used for building the product.

ii. Protecting Valuable Intellectual Property (IP) and Trade Secrets

To secure proprietary assets like model weights, algorithms, datasets, and servers:

- Milestone enforce strict access controls using role-based authentication and encryption.
- Data and model weights are stored in secure environments with both physical- and cyber-security measures.
- Policies ensure unreleased model weights are accessible only to authorized personnel under stringent monitoring protocols.

Milestone Systems protects valuable IP and trade secrets through measures such as:

- Access Controls: Limiting access to proprietary information to authorized personnel only.
- Encryption: Using encryption for data at rest and in transit.
- Secure Storage: Ensuring that sensitive data is stored in secure environments with strict access controls.
- AI Usage policy: Sharing IP and trade secrets with unauthorized GenAI/GPAI tools are prohibited

iii. Vulnerability Management Processes

As a CVE Numbering Authority (CNA) under the MITRE domain, Milestone follows the industry's best practices in managing and responding to security vulnerabilities discovered in our products.

Milestone guarantee that we make a thorough effort to identify and mitigate potential vulnerabilities in our software, reducing the customer's risk of deploying or using Milestone's software products or services in their environment.

Milestone has a structured vulnerability management process that includes:

- Identification: Regular vulnerability testing and penetration testing.
- Documentation: Maintaining detailed records of identified vulnerabilities and actions taken.
- Patch Management: Providing software patches to address vulnerabilities and critical bugs.

Milestone scores reported vulnerabilities using the industry vulnerability scoring system CVSSv3.1 Common Vulnerability Scoring System and provide patches according to the scores described in our Vulnerability management Policy:

https://doc.milestonesys.com/cybersecurity/pdf/en-US/Milestone_VulnerabilityManagement_en-US.pdf

In case the person reporting the vulnerability has disclosed their contact information, Milestone will collaborate with them on details, such as the CVSSv3.1 score, content of security advisory, and date for the external disclosure.

Milestone Systems takes proactive actions to address risks and vulnerabilities, including collaborating with stakeholders such as development teams and external security experts and our security response team can be contacted directly through our website if any vulnerabilities have been detected:

<https://www.milestonesys.com/support/help-and-documentation/cyber-security/vulnerability-form/>

iv. Frequency of Security Measure Reviews

Security measures are reviewed periodically to adapt to evolving threats and the SDL is updated to align with industry standards, emerging threats, and organizational policies.

v. Insider Threat Detection Program

Milestone Systems' comprehensive security practices, including access controls and monitoring contribute to detecting and mitigating insider threats.

h. How does your organization address vulnerabilities, incidents, emerging risks?

We address vulnerabilities, incidents, emerging risks, and misuse across the AI lifecycle using the planned MLOps and CI/CD with rigorous testing, continuous monitoring, and implementing appropriate mitigation measures throughout development, deployment, and post-deployment phases, when possible, as described earlier.

Incidents can be reported in our Whistleblower and Incident Portal:

<https://www.milestonesys.com/company/about-milestone/csr/whistleblower>

Our security response team can be contacted directly through our website if any vulnerabilities have been detected:

<https://www.milestonesys.com/support/help-and-documentation/cyber-security/vulnerability-form/>

Any further comments and for implementation documentation

No answer provided

Section 3 - Transparency reporting on advanced AI systems

a. Does your organization publish clear and understandable reports and/or technical documentation related to the capabilities, limitations, and domains of appropriate and inappropriate use of advanced AI systems?

- i. How often are such reports usually updated?**
- ii. How are new significant releases reflected in such reports?**
- iii. Which of the following information is included in your organization's publicly available documentation: details and results of the evaluations conducted for potential safety, security, and societal risks including risks to the enjoyment of human rights; assessments of the model's or system's effects and risks to safety and society (such as those related to harmful bias, discrimination, threats to protection of privacy or personal data, fairness); results of red-teaming or other testing conducted to evaluate the model's/system's fitness for moving beyond the development stage; capacities of a model/system and significant limitations in performance with implications for appropriate use domains; other technical documentation and instructions for use if relevant.**

Technical documentation related to our advanced/High-Risk AI systems is part of our AI Governance and GRC tool. Information related to the capabilities, limitations, and domains of appropriate and inappropriate use of advanced AI systems is made publicly available as part of the EULA and product documentation.

i. How often are such reports usually updated? ii. How are new significant releases reflected in such reports?

The technical documents and product information is updated regularly and reflect the changes in the new releases.

iii. Which of the following information is included in your organization's publicly available documentation?

Our documentation is planned to include details and results of evaluations for potential risks, assessments of safety and societal impacts, results of testing, model capacities, limitations, and mitigating actions. The AI capabilities, limitations, and domains of appropriate and inappropriate use are documented as part of the product documentation and EULA deployed with the system.

b. How does your organization share information with a diverse set of stakeholders (other organizations, governments, civil society and academia, etc.) regarding the outcome of evaluations of risks and impacts related to an advanced AI system?

We only share information with partners and customers. The AI capabilities, limitations, and domains of appropriate and inappropriate use is part of the product documentation and EULA deployed with the system.

c. Does your organization disclose privacy policies addressing the use of personal data, user prompts, and/or the outputs of advanced AI systems?

Yes, we share publicly our Human Rights Policy, Data Ethics Policy and Global Data Protection & Privacy Policy. Our AI-policy is shared internally for now.

<https://www.milestonesys.com/company/about-milestone/about-us/responsible-technology/>

d. Does your organization provide information about the sources of data used for the training of advanced AI systems, as appropriate, including information related to the sourcing of data annotation and enrichment?

The complete data lineage and data annotation process is documented as part of our planned MLOps process and in the AI Governance tool.

e. Does your organization demonstrate transparency related to advanced AI systems through any other methods?

Yes, we have launched Project Hafnia. Project Hafnia acts as the trusted librarian of AI video data, carefully curating, tagging, and delivering ethically sourced, regulation-ready video data for AI model training. By leveraging Project Hafnia, AI learns from high-quality, verified video data, ensuring not only precision and compliance but also the protection of citizen privacy.

<https://www.milestonesys.com/resources/content/articles/project-hafnia-game-changer-ai-model-training/>

Any further comments and for implementation documentation

No answer provided

Section 4 - Organizational governance, incident management and transparency

a. How has AI risk management been embedded in your organization governance framework? When and under what circumstances are policies updated?

AI risk management is embedded in our AI governance- and GRC tool. Policies are updated based on new developments, and regular reviews scheduled in our GRC tool to ensure compliance and effectiveness.

b. Are relevant staff trained on your organization's governance policies and risk management practices? If so, how?

All employees are trained in general AI literacy and the EU AI Act, according to requirements in the EU AI Act, while employees working with the development of AI products and models are following an advanced individual training program, following High-Risk AI requirements from the EU AI Act.

c. Does your organization communicate its risk management policies and practices with users and/or the public? If so, how?

As part of our Responsible Development Policy, our risk management policy is shared with all employees involved in AI development, but not shared to the public.

d. Are steps taken to address reported incidents documented and maintained internally? If so, how?

Yes, steps taken to address reported incidents are documented and maintained internally through comprehensive records and reporting mechanisms as part of our AI Governance and GRC Tool.

e. How does your organization share relevant information about vulnerabilities, incidents, emerging risks, and misuse with others?

Milestone disclose identified vulnerabilities at least 90 days after the vulnerability has been communicated to Milestone, or, alternatively, not before a mutually agreed date.

f. Does your organization share information, as appropriate, with relevant other stakeholders regarding advanced AI system incidents? If so, how? Does your organization share and report incident-related information publicly?

Milestone disclose identified vulnerabilities at least 90 days after the vulnerability has been communicated to Milestone, or, alternatively, not before a mutually agreed date.

g. How does your organization share research and best practices on addressing or managing risk?

We share research through publications, collaborations, and engagement with universities.

h. Does your organization use international technical standards or best practices for AI risk management and governance policies?

Yes, as mentioned earlier, Milestone actively develop and use international standards as active members of Working Group 3 (WG3) under CEN-CENELEC JTC 21, that focuses on engineering

aspects of creating standards for the EU AI Act. This also gives us early insights into the EU AI Act requirements on risk management for AI systems as developed in Working Group 2

Any further comments and for implementation documentation

No answer provided

Section 5 - Content authentication & provenance mechanisms

a. What mechanisms, if any, does your organization put in place to allow users, where possible and appropriate, to know when they are interacting with an advanced AI system developed by your organization?

As of now, the users are not directly interacting with the AI system, this is running in the background to i.e. recognizing objects in the video etc. The output of the AI is clearly marked as identified objects within a bounding box.

b. Does your organization use content provenance detection, labeling or watermarking mechanisms that enable users to identify content generated by advanced AI systems? If yes, how? Does your organization use international technical standards or best practices when developing or implementing content provenance?

Milestone digitally sign (watermark) original content from cameras to ensure that the end-user know that this data has not been tampered with.

Milestone also use advanced AI to generate synthetic data, but this content is only used for training and not presented to the users. For AI that is generating new content that is presented to the end-user and that they could misconceive as human generated content, we plan to use watermarking or similar techniques that highlight this content as not original.

Any further comments and for implementation documentation

No answer provided

Section 6 - Research & investment to advance AI safety & mitigate societal risks

a. How does your organization advance research and investment related to the following: security, safety, bias and disinformation, fairness, explainability and interpretability, transparency, robustness, and/or trustworthiness of advanced AI systems?

Milestone prioritize research to mitigate societal, safety and security risks. This includes conducting, collaborating on and investing in research that supports the advancement of AI safety, security, and trust, and addressing key risks, such as prioritizing research on upholding democratic values, respecting human rights, protecting children and vulnerable groups, safeguarding intellectual property rights and privacy, and avoiding harmful bias, mis- and disinformation, and information manipulation.

Example: Milestone is an active research partner in the REPAI project with Aalborg University; the Responsible AI for Value Creation project is dedicated to positioning Denmark at the global fore front of responsible and impactful AI development. The project focuses on producing world-class interdisciplinary research with tangible, practical benefits. By examining real-world use cases, the project aim to offer new insights, foster a shared interdisciplinary language, and propose necessary updates to existing AI legislation. The goal is to assist private companies and public organizations in building trusted, ethical, and compliant AI solutions. Additionally, we are committed to advancing interdisciplinary collaboration, benefiting the next generation of researchers, and through this influencing students who will propagate findings and values into society.

<https://vbn.aau.dk/en/projects/responsible-ai-for-value-creation>

b. How does your organization collaborate on and invest in research to advance the state of content authentication and provenance?

Milestone use already existing techniques for content authentication, hence we are not doing research in this area right now.

c. Does your organization participate in projects, collaborations, and investments in research that support the advancement of AI safety, security, and trustworthiness, as well as risk evaluation and mitigation tools?

Yes, as mentioned earlier we actively participate in projects, collaborations, and investments to advance AI safety, security, trustworthiness, and risk mitigation.

d. What research or investment is your organization pursuing to minimize socio-economic and/or environmental risks from AI?

Milestone efforts in both research and development focus on minimizing the compute used for AI, if possible, and we always prioritize to use less compute intensive models if they solve a problem with the needed accuracy.

As a large company, Milestone adhere to the EU Corporate Sustainability Reporting Directive (CSRD). CSRD plays a crucial role in addressing socio-economic and environmental risks associated with AI by mandating comprehensive sustainability reporting. The CSRD requires Milestone to report on both how sustainability issues affect their performance and how their activities impact the environment and society. This principle ensures that companies consider

the broader implications of their AI investments, including potential socio-economic and environmental risks. By mandating detailed disclosures, the CSRD enhances transparency, allowing stakeholders to evaluate a company's sustainability performance. The reported information must align with the EU Taxonomy, which classifies environmentally sustainable economic activities. This helps Milestone to focus our AI research and investments on areas that contribute positively to sustainability goals. Milestone must identify and manage ESG risks, including those related to AI. This involves setting targets, policies, and KPIs that address the socio-economic and environmental impacts of AI technologies.

Overall, the CSRD provides a robust framework for Milestone to disclose our sustainability efforts, including those related to AI, thereby promoting transparency, accountability, and sustainable investment practices.

Any further comments and for implementation documentation

No answer provided

Section 7 - Advancing human and global interests

a. What research or investment is your organization pursuing to maximize socio-economic and environmental benefits from AI? Please provide examples.

As mentioned earlier, Milestone efforts in both research and development focus on minimizing the compute used for AI, if possible, and we always prioritize to use less compute intensive models if they solve a problem with the needed accuracy.

b. Does your organization support any digital literacy, education or training initiatives to improve user awareness and/or help people understand the nature, capabilities, limitations and impacts of advanced AI systems? Please provide examples.

Not yet, but we are in dialogues with affected communities to help the understand the potential impact of our AI systems.

c. Does your organization prioritize AI projects for responsible stewardship of trustworthy and human-centric AI in support of the UN Sustainable Development Goals? Please provide examples.

Yes, we prioritize AI projects that support responsible stewardship of trustworthy and human-centric AI in support of the UN Sustainable Development Goals. For example, Milestone participate in the REPAI project with Aalborg University as mentioned earlier with a goal to assist in building trusted, ethical, and compliant AI solutions.

<https://vbn.aau.dk/en/projects/responsible-ai-for-value-creation>

d. Does your organization collaborate with civil society and community groups to identify and develop AI solutions in support of the UN Sustainable Development Goals and to address the world's greatest challenges? Please provide examples.

Milestone promote the UN Sustainable Development Goals and engage with civil society and community groups through our participation in the REPAI project with Aalborg University where one of the goals is to assist public organizations in building trusted, ethical, and compliant AI solutions.

<https://vbn.aau.dk/en/projects/responsible-ai-for-value-creation>

Milestone also engage with civil society and community groups through our role in the Danish Data Ethics Council.

<https://dataetiskraad.dk/om-dataetisk-raad/medlemmer/peter-damm>

The Data Ethics Council monitors technological developments and is especially interested in ethical issues and dilemmas resulting from the use of data.

To obtain the many advantages offered by the use of data, the Data Ethics Council seeks to support development in an ethical manner that takes into consideration the citizens' fundamental rights, legal certainty and fundamental values of society.

Any further comments and for implementation documentation

No answer provided