

Responsible Artificial Intelligence

The Company recognizes both the opportunities and risks associated with Artificial Intelligence (AI) in supporting business operations, enhancing productivity, and driving innovation. At the same time, the Company places strong emphasis on risk management, transparency, cybersecurity, personal data protection, and the ethical use of AI to ensure that AI adoption creates sustainable value for the organization, society, and all stakeholders.

The Company has established a policy related to responsible Artificial Intelligence, together with digital technology development and usage, which are integrated into the Company's Information Technology and Cyber Security Policy. These policies provide a governance framework for the adoption and use of AI solutions developed by external providers and implemented within the organization. The framework is guided by key principles including Accuracy and Reliability, Fairness, Transparency, Accountability, Human Oversight, Privacy and Security, as well as consideration of environmental and social impacts.

1. Risk-Based AI Governance

The Company adopts a risk-based approach to AI governance by classifying AI tools according to their risk profile and suitability for organizational use.

- **Green List (Approved AI Services):** AI services procured and approved by the Company following assessments of cybersecurity, data protection, and risk management requirements. These solutions may be used within approved business contexts and defined operational boundaries.
- **Yellow List (Conditionally Approved AI Services):** General-purpose AI services that may be used under specific restrictions. The processing of confidential business information, sensitive corporate information, or personal data through these services is strictly prohibited.
- **Red List (Prohibited AI Services):** AI services deemed to present unacceptable security, privacy, compliance, or business risks. Access to these services is restricted through technical and administrative controls.

The Company regularly communicates and updates the list of approved AI tools through internal communication channels and the corporate intranet to ensure employees remain informed and compliant with applicable requirements. This approach enables the organization to manage AI usage according to risk levels, reduce the likelihood of sensitive information being processed through unapproved services, and promote responsible AI adoption across the enterprise.

2. Responsible Selection and Management of AI Services

The Company does not develop its own foundation AI models. Instead, it focuses on selecting technologies and service providers that demonstrate internationally recognized standards and strong commitments to sustainability. This includes leading cloud and AI platform providers whose governance frameworks support cybersecurity, privacy protection, reliability, responsible AI practices, and environmental stewardship.

The procurement and deployment of AI services are governed through enterprise architecture, risk management, and license management processes to ensure cost-effectiveness, value realization, and compliance with corporate requirements.

Employees who are granted access to Company-provided AI services are required to acknowledge and comply with the Company's Acceptable Usage Policy (AUP), which reinforces responsibilities related to data protection, ethical conduct, information security, and responsible AI usage.

The Company seeks to leverage AI to improve productivity, reduce repetitive tasks, enhance knowledge accessibility, and support more effective business decision-making. In addition, the Company developed methodologies and performance indicators to assess the impact of AI adoption on economic, social, and environmental outcomes in a structured and measurable manner.

3. AI Risk Management and Governance

All AI-related initiatives are subject to a formal risk assessment and control evaluation process prior to deployment. The assessment encompasses system architecture, data storage and data flow considerations, cybersecurity risks, personal data protection requirements, and potential impacts on business operations. This process helps ensure that AI solutions are implemented in a secure, responsible, and well-governed manner, aligned with the Company's risk management framework and business objectives.

The Company's Solution Architecture Committee reviews project suitability, associated risks, and control measures before approval to ensure that AI-enabled systems are designed and governed appropriately.

For significant systems, post-implementation reviews and performance monitoring are conducted, with key findings and outcomes reported to the IT Committee to support oversight, effectiveness reviews, and continuous improvement.

4. Accuracy, Transparency, Fairness, and Human Oversight

The Company requires AI-generated outputs to be appropriately identified or labeled where applicable, enabling users to recognize when information has been generated or assisted by AI. Users are expected to exercise professional judgment and verify the accuracy of AI-generated outputs before relying upon them.

Prior to deployment, AI-enabled solutions undergo testing and refinement processes to improve output quality and reliability. Relevant teams establish performance criteria relating to accuracy, completeness, and suitability according to business objectives and continue monitoring operational effectiveness after deployment.

The Company views AI as a decision-support tool rather than a replacement for human accountability or decision-making authority. Users remain responsible for assessing the appropriateness, accuracy, and potential impact of AI-generated outputs used in business processes or decision-making.

The Company places strong emphasis on fairness and the avoidance of bias in AI usage. AI outputs are periodically reviewed, taking into consideration the suitability of data, assumptions, and usage contexts to reduce the risk of unfair outcomes, discrimination, or unintended adverse impacts on stakeholders.

Where AI is used to support analysis, recommendations, or decision-making processes that may affect individuals or stakeholders, the related initiative must undergo review of system architecture, risk considerations, and control measures by the Solution Architecture Committee and responsible project stakeholders. These reviews help mitigate risks arising from inaccurate data, biased assumptions, or inappropriate outcomes.

The Company promotes the use of high-quality and fit-for-purpose information while periodically reviewing AI outcomes to ensure fairness, transparency, and alignment with organizational ethical standards.

For decisions that may significantly affect individuals, business operations, or stakeholders, clearly defined responsible parties are designated to review and oversee AI-assisted outcomes, ensuring accountability, fairness, and appropriate governance over important decisions.

5. AI Incident Reporting, Appeals, and Escalation Process

The Company has established dedicated communication and reporting channels related to AI usage to support the reporting of incidents, ethical concerns, information security issues, data privacy concerns, inaccuracies in AI-generated outputs, or other impacts arising from AI use.

Reported matters are subject to investigation and managed in accordance with the Company's governance and risk management processes to facilitate appropriate corrective actions, improvements, and risk mitigation measures.

Employees, users, or affected stakeholders may raise concerns, submit complaints, or request reviews of AI-generated outcomes through designated channels. Such cases are reviewed, assessed, and monitored according to their significance and potential impact to support transparency, fairness, and continuous improvement.

The IT Committee oversees significant AI-related issues, including output inaccuracies, inappropriate data usage, ethical concerns, and incidents that may affect stakeholders. Appropriate root-cause analysis, corrective actions, and follow-up measures are implemented with clearly assigned responsibilities.

For AI applications that may directly affect individuals, stakeholders, safety-critical environments, or business-critical operations, the Company applies a precautionary approach. Such applications must undergo formal risk assessments, control reviews, and approval by relevant governance authorities before deployment.

The Company does not permit AI to replace human decision-making in high-risk scenarios or processes that may have significant impacts on individuals, safety, or critical business operations. This approach supports transparency, fairness, trust, and responsible AI adoption throughout the organization.

6. AI Knowledge Development and Responsible AI Culture

The Company continuously promotes AI literacy and capability development for employees through training programs, workshops, and knowledge-sharing activities. These initiatives cover topics including ethical AI use, cybersecurity, personal data protection, and the effective and responsible application of AI technologies.

The Company also collaborates with external experts and technology partners to enhance awareness and understanding of responsible AI practices and emerging technologies.

In addition, the Company has established an internal Community of Practice for AI (CoP AI), providing a platform for experience sharing, best-practice exchange, AI use case discussions, and innovation development. This initiative supports the growth of organizational capabilities and fosters a culture of responsible AI adoption across the enterprise.

7. International Standards and Continuous Improvement

The Company's AI governance practices operate within its broader information security, cybersecurity, and privacy management framework, which aligns with internationally recognized standards including:

- **ISO/IEC 27001** – Information Security Management System (ISMS)
- **ISO/IEC 27018** – Protection of Personally Identifiable Information (PII) in Public Cloud Services
- **ISO/IEC 27032** – Cybersecurity Guidelines

These standards support systematic risk management and control mechanisms relevant to the secure and responsible use of AI.

To further strengthen AI governance in line with emerging global best practices, the Company is preparing for the adoption of ISO/IEC 42001 Artificial Intelligence Management System (AIMS) and intends to pursue certification in the future, taking into consideration the readiness of qualified certification bodies and auditors within Thailand.

The Company believes that responsible, transparent, secure, and ethical use of AI is a key enabler of long-term value creation, supporting sustainable business growth while contributing positively to society and the environment.