

Artificial Intelligence Risk Taxonomy and Assessment Framework

July 7, 2026

Executive Summary

This Framework is formulated by the Ministry of Digital Affairs pursuant to Article 16 of the Artificial Intelligence Basic Act. It serves as the basis upon which the competent authorities for the relevant industries establish their Artificial Intelligence (AI) risk management regulations, as well as a reference document for both industries and the general public to understand the AI risk governance system in Taiwan.

■ Competent Authorities for the Relevant Industries

Primary Subjects of this Document. When formulating management regulations pursuant to Paragraph 2, Article 16 of the Artificial Intelligence Basic Act, each competent authority for the relevant industries shall follow the operational workflow prescribed in this Framework to carry out risk identification, assessment, and response. Such authorities are advised to read the document in its entirety. The key points are as follows: Part 1 establishes the overall positioning of the governance architecture; Part 2 explains the applicable scope and core principles of this Framework; Part 3 sets out the main operational workflow for risk management. “AI Risk Management Measures Checklist” in the Appendix serves as an operational tool, for which cross-referencing each step in Part 3 is recommended for use.

■ Industries

Industry stakeholders are advised to begin with “Table of AI Risk Types” (Table 1) under Section 3.2 (Risk Identification) in order to grasp the risk types covered by this Framework. This should be followed with Section 3.3 (Risk Assessment), together with item (a) Measures in Support of Development, under Section 3.4.1, so as to understand the government's supportive stance

toward innovative AI use cases and its expectations regarding industry self-regulatory mechanisms.

■ General Public

The general public is advised to read Part 1 (The AI Risk Governance System) to gain an understanding of the overall objectives of the government's efforts to advance AI risk governance. The explanation of risk types under Section 3.2 (Risk Identification) shed light on the potential societal impacts of AI use cases. Those wishing to learn more about how the government protects the public from the harms posed by high-risk AI use cases may refer to the relevant content in Sections 3.3 and 3.4. The stakeholder engagement mechanism described in Part 2 may be used to communicate concerns regarding the use of AI.

Document Map

The table below sets out the overall structure of this document, together with the core function of each section and its recommended readership, enabling readers to orient themselves before reading.

Section	Core Function	Recommended Readership
Reading Guide	Audience-specific reading recommendations that assist readers of differing backgrounds in quickly locating the content most relevant to them.	● All readers
1. The AI Risk Governance System	Describes the statutory basis, positioning, and overall governance architecture of this Framework to provide the background context required for reading Part 2.	● All readers
2. Positioning and Scope of Application of the AI Risk Taxonomy and Assessment Framework	Explains the positioning, statutory basis, scope of application and exclusions, principle of proportionality, and collaborative governance model of this Framework. These are necessary contexts before proceeding to the operational workflow.	● All readers
3. Operational Workflow of the AI Risk Taxonomy and Assessment Framework	The principal workflow to be followed by each competent authority for the relevant industries in conducting AI risk management, comprising the following four sequential steps: (1) AI Use Case Inventory — compiling basic information on AI use cases within the authority's purview in	● Competent authorities for the relevant industries (primary); ● Industries

	preparation for risk assessment; (2) Risk Identification — identifying potential risk types by referring to Table of AI Risk Types (Table 1); (3) Risk Assessment — assessing the degree of risk impact and determining whether an AI use case involves high risks (Appendix 2); (4) Risk Response — planning Innovation Support or management measures according to the risk circumstances and reviewing applicable laws and regulations on a regular basis.	
Appendices 1. AI Risk Management Measures Checklist; 2. Table of Illustrative Examples of Serious Harm	Operational tools for use in each step of Part 3, comprising: (i) Table of AI Use Case Inventory (corresponding to Step 1); (ii) Table of Risk Identification, Assessments, and Responses (corresponding to Steps 2 to 4); and (iii) Supplementary Notes on the Inventory of Management Regulations.	● Competent authorities for the relevant industries

Table of Contents

1. THE AI RISK GOVERNANCE SYSTEM	1
2. POSITIONING AND SCOPE OF APPLICATION OF THE AI RISK TAXONOMY AND ASSESSMENT FRAMEWORK	3
3. OPERATIONAL WORKFLOW OF THE AI RISK TAXONOMY AND ASSESSMENT FRAMEWORK.....	4
3.1 AI USE CASE INVENTORY.....	5
3.2 RISK IDENTIFICATION.....	5
3.3 RISK ASSESSMENT.....	9
3.4 RISK RESPONSE.....	11
APPENDICES.....	17
APPENDIX 1: AI RISK MANAGEMENT MEASURES CHECKLIST FOR THE COMPETENT AUTHORITIES FOR THE RELEVANT INDUSTRIES.....	17
APPENDIX 2: TABLE OF ILLUSTRATIVE EXAMPLES OF SERIOUS HARM.....	22

1. The AI Risk Governance System

- 1.1 The AI risk governance system in Taiwan is grounded in the development principles set out in Article 4 and centered on the AI risk taxonomy and assessment framework established under Article 16 of the Artificial Intelligence Basic Act. By defining the substance of the risk types, the system assists the competent authorities for the relevant industries in advancing different governance measures of risk mitigation and to strengthen the foundations of trustworthy AI (see Figure 1).

AI Risk Governance System

AI Development Principles

(Article 4 of the AI Basic Act)



AI Risk Taxonomy and Assessment Framework

(Paragraph 1 of Article 16)



Governance Work Across Ministries and Agencies

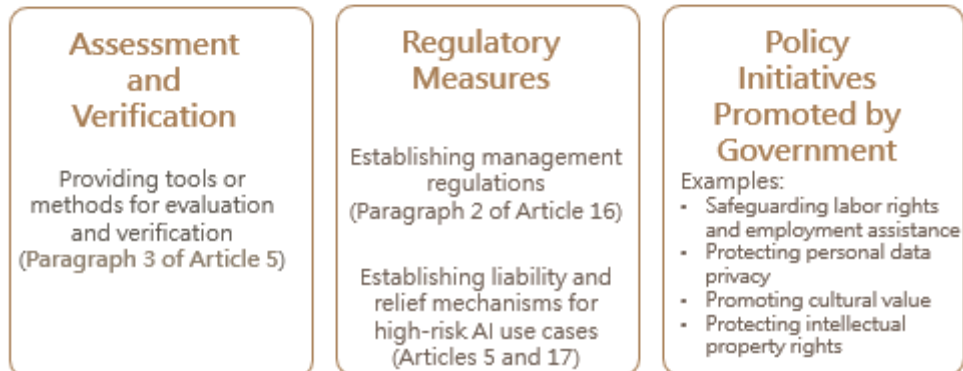


Figure 1: AI Risk Governance System

2. Positioning and Scope of Application of the AI Risk Taxonomy and Assessment Framework

- 2.1 This Framework is formulated by the Ministry of Digital Affairs pursuant to Article 16 of the Artificial Intelligence Basic Act. It centers on the promotion of technological innovation and industrial development, with priority placed on support and guidance. When establishing risk-based management regulations in accordance with Paragraph 2, Article 16 of the Act, each competent authority for the relevant industries shall follow the operational workflow prescribed in this Framework for the identification, assessment, and response to risks, so as to progressively build a common understanding of risks and a shared governance capacity. The principal function of this Framework is to assist each competent authority for the relevant industries in identifying and assessing the scope of impacts and the degree of risks that a given AI use case may entail, so as to construct the procedural basis of formulating risk-based management regulations and adopting accompanying innovation support measures. The purpose is to achieve the balance between management and promotion.
- 2.2 With respect to its scope of application, this Framework applies primarily to AI use cases in the civilian domain and expressly excludes military uses. The research and development, deployment, and application of AI systems for national defense or military purposes shall be handled in accordance with the relevant laws and regulations governing national defense.
- 2.3 In advancing AI governance in Taiwan, the enactment of laws by the government and the legislature is not the sole consideration. Of even greater importance is the incorporation of recommendations from all sectors, including industry and academia, and the fostering of cross-

disciplinary collaboration across fields such as education, law, health education, and children and youth affairs. A collaborative governance model constructed through an inclusive working structure serves to promote the Artificial Intelligence Basic Act

3. Operational Workflow of the AI Risk Taxonomy and Assessment Framework

Each competent authority for the relevant industries follows the operational workflow of "AI Use Case Inventory → Risk Identification → Risk Assessment → Risk Response" and on that basis establishes risk-based management regulations.

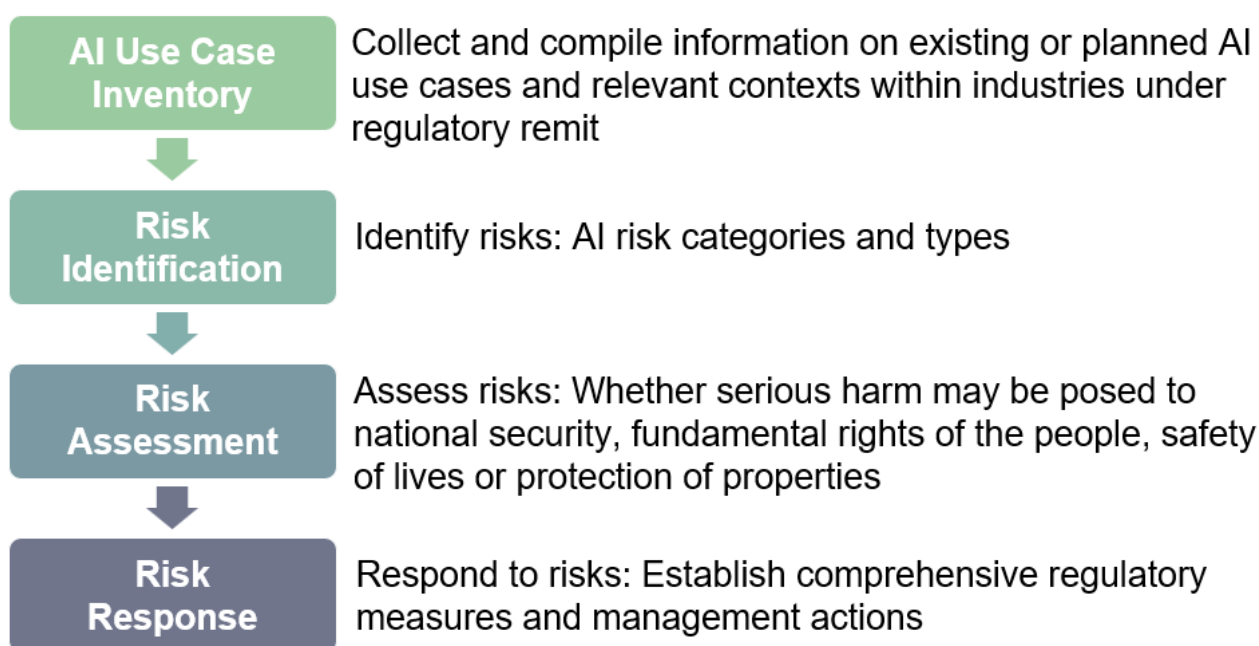


Figure 2: Operational Workflow of the AI Risk Taxonomy and Assessment Framework

3.1 AI Use Case Inventory

★This step is to be completed by reference to “Table of AI Use Case Inventory”, Item 1 of Appendix 1

3.1.1 Collect and compile background information on the AI use cases currently in use or planned for deployment within industries within the authority's purview, including the AI technologies involved, the relevant stakeholders, and the fields of application or industries concerned, as preparation for the identification, assessment, and response to risks.

3.2 Risk Identification

★This step is to be completed by reference to the "Risk Identification" column in "Risk Identification, Assessment, and Response Table", Item 2 of Appendix 1

3.2.1 Please refer to Table of AI Risk Types (Table 1) to examine whether the AI use cases in a given industry involve any of the risk subtypes listed in the table.

3.2.2 This stage aims to identify the potential risk types. Any risk identified as present shall then proceed to the next stage for assessment of its degree of risk.

3.2.3 The risk subtypes listed in Table of AI Risk Types represent a summary of the risks known at the present stage. The Ministry of Digital Affairs will conduct periodic and systematic reviews, taking into account the latest domestic and overseas research on AI risks and developments and the revision of international standards and regulations, in order to assess whether particular risk types should

be added, amended, or deleted. The competent authorities for the relevant industries shall likewise remain attentive to emerging risks not listed in this table.

Table 1: Table of AI Risk Types

Risk Type	Description
Risk Type A: Technical Design Flaws in AI Systems	
(A1) AI System Security Vulnerabilities and Attacks	An AI system may, owing to flaws in algorithmic design, contamination of training data, or hardware vulnerabilities, suffer security risks such as unauthorized access, data theft, or system manipulation.
(A2) Lack of Transparency or Explainability	The decision-making process of an AI system is difficult to understand or explain. As a result, users distrust the system and it is challenging to implement compliance and monitoring, ensure liability or rectify errors.
(A3) AI Behavior Misaligned with Human Intent and Social Values	During the design or operation of an AI system, any deviation in objective-setting may cause its behavior to drift from the developer's original intent or from societal and ethical values, producing outcomes below expectations and difficult to correct effectively.
(A4) Dangerous Capabilities of AI Systems	An AI system may possess capabilities sufficient to cause significant harm, such as deception, manipulation, autonomous acquisition of resources, or development of new capabilities. Such capabilities may be conferred by design, developed autonomously, or acquired through learning from the environment. Once misused or beyond control, they may cause large-scale damage difficult to anticipate.
(A5) Privacy Infringement and Non-Compliance with Personal	Where AI training data contain personal data, attention must be paid to whether they comply with the relevant provisions of the Personal Data Protection Act, including the lawful grounds for the collection, processing, or use of personal data and whether such use is consistent with the

Risk Type	Description
Data Protection Laws	specified purpose, and appropriate security measures must be adopted, so as to reduce the risk of leakage or improper use of personal data.
(A6) Risk of Intellectual Property Infringement	The training data of an AI system may contain content protected by intellectual property rights. Any unauthorized use may constitute infringement. An assessment shall be conducted on whether the generated work is substantially similar to another person's work, so as to avoid infringing the interests of the rights holder.
(A7) Unfair Discrimination or Bias	Owing to biases in training data or problems in algorithmic design, an AI system may produce unfair outcomes for particular groups (of certain ethnicity, gender or age for example), giving rise to systemic discrimination and subjecting these groups to unequal treatment.
(A8) False or Misleading Information	AI systems, large language models (LLMs) in particular, may at times generate content that is factually inaccurate, misleading, insufficiently researched, or difficult to comprehend. Such content may affect users' judgment and decision-making and, once disseminated on a large scale, may have an adverse impact on public understanding.
Risk Type B: Issues in Post-Deployment Operation and Human-AI Interaction	
(B1) Overreliance and Unsafe Use	Users may place excessive trust in the judgment of an AI system and adopt its recommendations without verification in critical contexts such as healthcare, law, and finance, or may develop an emotional dependence on an AI system, leading to the deterioration of their capacity for independent judgment and to the risk of improper decision-making.
(B2) Loss of Human Autonomy	When humans delegate important decisions to an AI system, or when an AI system itself makes decisions that affect human control, humans may lose their capacity for autonomous judgment.
(B3) Generation of Unlawful Content	An AI system may generate content that contravene prevailing laws and regulations. Such content includes sexual exploitation of children and youth; hate speech; incitement to violence; false advertising; or other unlawful information, potentially in violation of the Protection of Children and Youths Welfare and Rights Act, the Sexual Harassment Prevention

Risk Type	Description
	Act, the Fair Trade Act, the Consumer Protection Act, and the Personal Data Protection Act.
(B4) Fraud and Misuse of Deepfake Technology	AI technology has rendered tools for voice imitation, deepfake imagery, and automated content generation increasingly prevalent. Malicious actors may exploit such tools to impersonate identities and fabricate false images for the purpose of scams, extortion, manipulation of public opinion, and other unlawful acts.
(B5) Facilitation of Cyberattacks	AI technology can automate cyberattack activities and lower the technical threshold required to carry out attacks, such that those without a professional information technology background are also able to launch cyberattacks, thereby increasing cybersecurity risks.
(B6) Unauthorized Actions by Autonomous AI Agents	AI agent systems possess the capacity to plan autonomously, to invoke external tools, and to execute complex tasks on a continuing basis. Possibly due to incomplete objective-setting or changes in the environment, their behavior may gradually depart from the original instructions, to the extent of autonomously acquiring system access privileges beyond the authorized scope. Multi-agent systems may even cause chain reactions unforeseen by either developers or users because agents trigger each other. As a result, human intervention and correction on a timely basis will be difficult.
Risk Type C: Societal-Structural and Environmental Impacts	
(C1) Competitive Imbalance Among Companies and Countries	In the race for an advantage in AI technology, companies or countries may deploy systems hastily before these systems have been adequately tested. This heightens the risk of unsafe AI products entering the market and jeopardizes social safety and economic stability.
(C2) Power Centralization and Unfair Distribution of Benefits	The development of advanced AI technology requires substantial computing resources, expertise, and capital. Hence, highly influential technologies may be controlled by a small number of entities, exacerbating the unequal distribution of social resources.

Risk Type	Description
(C3) Increased Inequality and Decline in Employment Quality	The widespread adoption of AI systems may deepen socioeconomic inequality, including problems such as large-scale automation of work, a decline in employment quality, and an imbalance in labor-management relations.
(C4) Economic and Cultural Devaluation of Human Creativity	AI systems can reproduce and imitate the fruits of human creativity on a large scale, potentially undermining the economic foundation of creative industries such as journalism, art, and music, leading to a homogenization of cultural content and diminishing the social value of human creative work.
(C5) Environmental Harm	The development and operation of AI systems may have adverse effects on the environment. For example, generative AI models, particularly those employing deep learning techniques, require substantial energy during training, testing, and deployment, resulting in high electricity consumption and emission of greenhouse gases by data centers.
(C6) Cognitive Warfare and Information Sovereignty	AI technology significantly lowers the cost of large-scale information manipulation, enabling malicious actors to employ generative AI, bot accounts to manufacture fake public opinion, and deepfake content to systematically interfere with public discourse, electoral processes, and policy formation in democratic societies, thereby eroding the basis of social consensus and the people's trust in democracy.

3.3 Risk Assessment

★This step is to be completed by reference to the "Risk Assessment" column in "Risk Identification, Assessment, and Response Table", Item 2 of Appendix 1. The illustrative examples in Appendix 2 may also be consulted.

3.3.1 Upon completing the risk identification, the competent authority for the relevant industries shall assess the degree of risk impacts.

3.3.2 Assessment based on objective and inherent risks

The assessment on the degree of risks shall be based on the objective characteristics of the AI use case itself. The mere possibility of harm suffices, and the actual occurrence of harm is not required. In conducting the assessment, the mitigating effects of existing laws and regulations, administrative measures, or technical means shall not be taken into account. Whether such measures are adequate to reduce the risk is to be determined separately at the risk response stage.

3.3.3 Assessment is conducted on the existence of high-risk AI use cases in accordance with the legislative explanation under Paragraph 1, Article 17, of the Artificial Intelligence Basic Act and with reference to Paragraph 1, Article 5 of the Act.

- (a) The high riskiness of AI use cases is determined by reference to the potential of the risk and the degree of impacts.
- (b) For example, if the risk impact of an AI use case may cause serious harm to national security, the fundamental rights of the people (including but not limited to the bodily freedom, personal liberty, right to privacy, right to equality and other rights guaranteed under the Constitution, as well as rights protected under the international human rights covenants that have been incorporated into domestic laws, e.g., CEDAW), rights to life and safety, property protection, social order, or the ecological environment, the competent authority for the relevant industries shall determine that the use case in question constitutes a high-risk AI application as provided in Paragraph 1, Article 17 of the Artificial Intelligence Basic Act.

3.4 Risk Response

★This step is to be completed by reference to the "Response Measures" column in "Risk Identification, Assessment, and Response Table", Item 2 of Appendix 1 and to Item 3, "Supplementary Notes on the Inventory of Management Regulations and Future Planning."

3.4.1 Establishment of an inventory of existing measures and supplementing of management regulations in a timely manner

After completing the risk assessment, each competent authority for the relevant industries then proceeds to establish an inventory of the coverage of existing measures and to examine whether the prevailing laws, regulations, and administrative measures are sufficient to address the risks identified.

For risks inadequately covered or without measures in place yet, an authority shall, in accordance with the principle of proportionality, plan appropriate management or advocacy measures from the tools below, and shall periodically assess the AI use cases within its purview and actively review the formulation and amendment of applicable laws, regulations, and standards, so as to complete the relevant mechanisms and measures.

(a) Measures in Support of Development

When assessing the risks of AI use cases, the competent authority for the relevant industries shall also evaluate the opportunity costs and the risk of weakened service capacity arising from not adopting AI, so as to ensure that the risk assessment framework gives due regard to both the promotion of use cases and the prevention of harm. This approach precludes

overly conservatism that causes policies or services to fall behind societal needs and ensures the regulatory design consistent with the principle of proportionality.

i. Promoting the Development of Industry Standards and Self-Regulatory Codes

Pursuant to Paragraph 2, Article 16 of the Artificial Intelligence Basic Act, the competent authority for the relevant industries assists industries in formulating their own industry guidelines and codes of conduct by adopting the following measures:

- Convening industry representatives and academic experts to jointly formulate industry self-regulatory codes;
- Encouraging industry associations to establish self-regulatory mechanisms for their members;
- Publicly commending industry players that adopt self-regulatory codes, or giving them priority consideration in government procurement;
- Establishing mechanisms for periodical communication between industry self-regulatory organizations and the government.

International resources in common use (such as ISO 23894 and ISO 42001) may be consulted. Industries are encouraged to adopt different risk management measures, data quality controls, model transparency and monitoring mechanisms according to their respective roles.

ii. Providing Advisory Resources and Compliance Support

To assist industries in implementing this Framework smoothly, the competent authority may take into account the nature of

the relevant industries within its purview and provide reference resources such as AI risk self-assessment checklists, operational guidelines, or best-practice examples. This will enable industry participants to cross-reference the abstract concept about risk types with real-world scenarios. For small and medium-sized enterprises (SMEs) with relatively limited resources, the authority may also consider establishing a single consultation channel to assist SMEs in gradually developing their compliance capabilities.

iii. Conducting Education, Training, and Exchange Activities

To enhance industry's awareness and capacity for AI risk management:

- Organizing regular AI risk management workshops and seminars;
- Establishing industry exchange platforms to promote the sharing of best practices;
- Promoting cross-industry learning (for example, the exchange of risk management experience between the financial sector and the healthcare sector);
- Collaborating with academic and research institutions to develop professional training courses.

(b) Management Measures

After assessing the degree of risks of AI use cases within its purview, the competent authority for the relevant industries may, depending on the requirements, adopt various management measures below to respond to the risks. The following measures may be deployed individually or in combination according to the

risk circumstances, and need not be applied in sequence. The application of the measure in (iii) is subject to additional statutory requirements explained therein.

i. Taking Appropriate Administrative Management Actions

According to Its Powers and Duties

The competent authority for the relevant industries may, in accordance with its functions, duties and powers, require industries within its purview to submit explanatory materials, such as descriptive documents and verifiable information, to enforce transparency and disclosure (e.g., specific information to be provided during user interactions); and to ensure that content produced using AI technologies such as generative AI or deepfakes is proactively disclosed.

ii. Adopting Measures Such as Authorization and Review in Advance

Where the risk reaches a certain degree, the competent authority for the relevant industries may consider mandatory reviews of AI use cases in advance by requiring an impact assessment before the system formally comes online; a risk assessment and audit by an accredited third-party institution to ensure the system is transparent; adoption of risk management and data governance measures by developers, deployers, or users; retention of technical documents to demonstrate compliance with obligations; and recordation of incidents to ensure the traceability of the AI system's operation. In particular, where the deployment of an AI system involves the rights and interests of certain groups

(e.g., indigenous peoples, new immigrants, or other cultural communities), industries may be required to implement grievance and relief mechanisms, so as to provide affected persons with appropriate channels for complaints and avenues for the redress of their rights and interests.

iii. Imposing Restrictions or Prohibitions in Accordance with the Law

Where an AI use case, upon assessment, is found to give rise to the circumstances enumerated in Paragraph 1, Article 5 of the Artificial Intelligence Basic Act (such as infringement of the people's lives or of national security) and it is assessed that the risk of the use case in question still cannot be effectively managed or reduced by existing technical means, the competent authority for the relevant industries shall, in accordance with Paragraph 2, Article 16 of the same Act, restrict or prohibit the use case pursuant to the laws and regulations under its competence (including existing substantive administrative action laws or regulations subsequently formulated in connection with AI use cases).

iv. Necessary Measures Required by Law for High-Risk AI Use Cases

Where the competent authority for the relevant industries assesses an AI use case within its purview as high-risk, it shall, in accordance with Paragraph 2, Article 5 of the Artificial Intelligence Basic Act, require industries within its purview to clearly label caution notes or warnings; and shall, in accordance with Paragraph 1, Article 17 of the same Act,

clarify the attribution of liability and the conditions for such liability, and shall establish relief, compensation, or insurance mechanisms.

3.4.2 Responding to Risk in a Timely Manner and Supporting Innovative Development

When using this Framework to plan governance measures, the competent authority for the relevant industries shall give careful consideration to the possibility that excessive regulation may impede technological innovation and development. Furthermore, for AI use cases currently difficult to determine whether harm will be caused, it is advised to establish a continuous monitoring mechanism, regularly collecting application data and stakeholder feedback so as to identify signs of risk escalation at an early stage, instead of initiating the relevant procedures only after harm has been confirmed. For AI use cases already determined to be high-risk but whose patterns of use continue to evolve (e.g., new forms of misuse of deepfake technology), the authorities for the relevant industries are also advised to establish a mechanism for collecting real-world cases, so as to facilitate subsequent regulatory adaptation and the enhancement of management regulations.

Appendices

Appendix 1: AI Risk Management Measures Checklist for Competent Authorities for the Relevant Industries

AI Risk Management Measures Checklist

1. Table of AI Use Case Inventory

This table corresponds to the step "AI Use Case Inventory", with the principal purpose of organizing background information on AI use cases within the competent authority's purview, as the basis for subsequent risk assessment.

Item		Description	Remarks
1.1	Use Case Description		Please describe in detail the specific purpose, process, and objectives of this AI system.
1.2	AI Technology		Describe the core AI technologies that may be used (e.g., large language models (LLMs), computer vision, speech recognition, or generative AI). Key information such as data sources, data types, and deployment locations may also be added.
1.3	Stakeholders		Who will develop, deploy, use, or be affected by this system? Developers, deployers (enterprises or government agencies), end users, affected third parties, and the like.

2. Risk Identification, Assessments, and Responses Table: _____ AI Use Case

Conduct risk identification regarding the use case described above and assess whether serious harm may be caused. Next, establish an inventory on the coverage of existing measures in order to determine the priority of handling. It is recommended to consult stakeholders throughout the process.

Instructions for Completion:

- Please note that risk management (in the table below) is directed at a specific use case. If multiple use cases fall under an authority's competence, a complete risk identification, assessment, and response procedure must be carried out for each use case.
- Principles for marking the "Applicable" column at the Risk Identification stage: "No" means that the risk type bears a potential connection to the use case within the authority's purview but, upon assessment, the risk does not reach the identification threshold. "Not Applicable" means that the use case inherently lacks any source of risk of that risk type.
- For all risks marked "Applicable = Yes", "Coverage of Existing Measures" shall be assessed.
- "Whether Serious Harm May Be Caused" is an independent determination, unaffected by existing measures.
- Where "Serious Harm May Be Caused = Yes", the AI use case constitutes high risk and shall comply with the obligations under Article 17 of the Artificial Intelligence Basic Act. Where "Difficult to Determine at This Stage" is marked, a continuous monitoring mechanism shall be established in accordance with Section 3.4.2 of this Framework. Application data and stakeholder feedback shall be collected on a regular basis and a re-assessment shall be conducted at the next inventory.
- The "New Measures" column corresponds principally to the "Risk Response" step, and is to be completed for risk

types that are marked "Yes" for Risk Identification and for which the coverage of existing measures is "Insufficient or None" or "Partial".

Risk Identification			Risk Assessment			Risk Response Measures				
			Whether Serious Harm May Cause			Coverage of Existing Measures			New Measures	
Risk Item ¹	Applicable	Specific Risk Scenario Description ²	Whether Serious Harm May Be Caused (regardless of existing measures) ³	Criteria for Determined Harm	Assessment Notes	Name of Existing Laws or Administrative Measures (multiple entries permitted)	Type of Existing Measures	Coverage of Existing Measures	Measure Description	Notes
	☐ Yes ☐ No ☐ Not Applicable		☐ Yes ☐ No ☐ Difficult to Determine at This	☐ National Security ☐ Fundamental Rights of the People			Measures in Support of Innovation: ☐ Industry Standards and	☐ Sufficient (able to effectively prevent harm)		Type of measure: Risk mitigation

¹ Please conduct assessments on each of the 20 risk subtypes in Table 1 and add subsequent forms accordingly.

² The identification process and details may be recorded below.

³ By reference to the illustrative examples listed in Appendix 2, assess whether serious harm to national security, the fundamental rights of the people, life safety, property protection, social order, or the ecological environment is involved. In assessing the degree of harm, the vulnerability of disadvantaged groups and of their intersecting identities shall be taken into account. Where such a determination proves difficult, it is recommended to convene experts from different disciplines for discussion, as suggested in the main text.

Risk Identification			Risk Assessment			Risk Response Measures				
			Whether Serious Harm May Cause			Coverage of Existing Measures			New Measures	
Risk Item ¹	Applicable	Specific Risk Scenario Description ²	Whether Serious Harm May Be Caused (regardless of existing measures) ³	Criteria for Determined Harm	Assessment Notes	Name of Existing Laws or Administrative Measures (multiple entries permitted)	Type of Existing Measures	Coverage of Existing Measures	Measure Description	Notes
			Stage	☐ Life Safety ☐ Property Protection ☐ Social Order ☐ Ecological Environment			Self-Regulations ☐ Advisory and Compliance Resources ☐ Other Management Measures: ☐ Restriction or Prohibition ☐ Prior Authorization ☐ Transparency Requirements ☐ Other	☐ Partial (able to reduce but not eliminate harm) ☐ Insufficient (low relevance of existing measures to this risk) ☐ No relevant measures yet. Notes:		status: Whether it may affect AI innovation, development, and application:

Risk Identification			Risk Assessment			Risk Response Measures				
			Whether Serious Harm May Cause			Coverage of Existing Measures			New Measures	
Risk Item ¹	Applicable	Specific Risk Scenario Description ²	Whether Serious Harm May Be Caused (regardless of existing measures) ³	Criteria for Determined Harm	Assessment Notes	Name of Existing Laws or Administrative Measures (multiple entries permitted)	Type of Existing Measures	Coverage of Existing Measures	Measure Description	Notes
							Notes:			

3. Supplementary Notes on Inventory of Management Regulations and Future Plans

Where relevant practices and measures cannot be entered in the table above, general supplementary notes may be provided here.

Issue	Description

Appendix 2: Table of Illustrative Examples of Serious Harm

The contents are illustrative only and do not represent actual scenarios in Taiwan. This table serves as a reference tool for the competent authorities in assessing serious harm, and does not replace the comprehensive case-by-case determination made in accordance with the assessment principle set out in Section 3.3.2 of this Framework and the protected interests listed in Section 3.3.3.

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
(A1) AI System Security Vulnerabilities and Attacks	Disruption of the functioning of control systems of critical infrastructure	An AI system influences a user's decision-making without the user's knowledge, impairing the right to informational self-determination	Unlawful interference with smart medical devices or automated navigation	Unlawful tampering with a core financial database, resulting in loss of assets		Failure of a pollution-monitoring system and the resulting leakage of hazardous substances
(A2) Lack of Transparency or Explainability	Algorithms in major defense decision-making systems lack auditability	Fair procedures lacking for administrative	Unspecific attribution of liability for errors in AI-	No mechanisms for complaints or explanations regarding declines		

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
		dispositions or review of relief payments	assisted medical diagnosis	based on credit scores		
(A3) AI Behavior Misaligned with Human Intent and Social Values		Patterns of conduct that systematically infringe upon individual dignity			Public ethical standards adversely affected by algorithms	
(A4) Dangerous Capabilities of AI Systems	An AI system is used to generate or to assist in obtaining information related to controlled hazardous substances		Loss of control over an automated system and the resulting industrial or traffic accidents	Paralysis of the national financial settlement and clearing system	Triggering large-scale social unrest or regional conflict	
(A5) Privacy Infringement and Non-Compliance with Personal	Leakage of the nation's critical classified data or sensitive information	Abuse of biometric technology and mass surveillance		Misappropriation of personal data to conduct unlawful financial transactions	Digital governance models that improperly steer	

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
Data Protection Laws					behavior and cause biases	
(A6) Risk of Intellectual Property Infringement					Loss of a company's core trade secrets or patent assets	
(A7) Unfair Discrimination or Bias		Systematic differential treatment in workplace recruitment and admissions review	Harm to rights and interests due to unfair medical-resource-allocation algorithms	Specific groups excluded from basic financial services	Aggravation of conflict and antagonism between social groups	
(A8) False or Misleading Information	Public opinion improperly steered by automated tools	Depriving the public of the right to obtain accurate public information	Casualties caused by erroneous disaster-prevention guidance or public-health information	Substantial losses to investors due to misleading information	Damage to the credibility of the media and to the function of social communication	
(B1) Overreliance	Strategic failures due to blind		Automated medical decision-making		Deterioration of the social	

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
and Unsafe Use	adherence to AI recommendations for critical decisions		delaying the window for emergency care		system's capacity for emergencies and autonomous responses	
(B2) Loss of Human Autonomy		Without the user's informed consent, an AI system induces or constrains, through algorithms, the user's political voting intentions, medical choices, or major financial decisions				
(B3) Generation of Unlawful Content	Broadcasting of instructions for unlawful violence or extremist behavior	Dissemination of child and youth sexual exploitation content or seriously discriminatory speech	Distribution of content inciting self-harm or unlawful medical guidance		Systematic incitement of unlawful criminal conduct	

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
(B4) Fraud and Misuse of Deepfake Technology	Fabricating the speech of public officials to affect the functioning of government	Dissemination of non-consensual intimate images		Use of deepfake technology to fraudulently obtain substantial assets	Erosion of the foundation of social trust through disinformation	
(B5) Facilitation of Cyberattacks	Paralysis of government, military, or critical communications networks		Unlawful remote control of life-sustaining medical-equipment systems	Ransomware locking up operational assets of companies on a large scale		
(B6) Unauthorized Actions by Autonomous AI Agents	An AI agent exceeds the scope of its authorization and autonomously accesses government or critical infrastructure systems		Due to behavioral deviation, an AI agent autonomously executes instructions that endanger personnel's safety in medical or industrial settings	An AI agent autonomously executes unauthorized financial transactions and causes substantial property loss	Multi-agent systems trigger each other and cause chain service disruptions that are difficult to halt. As a result, society's normal	

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
					functioning is affected	
(C1) Competitive Imbalance Among Companies and Countries	Core AI technology is highly concentrated among a small number of suppliers, creating the risk of technology supply-chain disruption				Supply disruptions result in operational losses due to excessive concentration of and overreliance on the technologies in key industries	
(C2) Power Centralization and Unfair Distribution of Benefits					AI compute resources and technological capacities are concentrated among a small number of entities. As a	

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
					result, small and medium-sized enterprises, academic and research institutions, and remote communities are unable to gain equal access, structurally widening the digital divide	
(C3) Increased Inequality and Decline in Employment Quality	Upheaval in the labor market triggers political and economic instability			Mass unemployment leads to the depletion of individuals' assets	Excessive burden on the social safety net leads to institutional collapse	

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
(C4) Economic and Cultural Devaluation of Human Creativity		Generative AI systems use unauthorized content for training on a large scale and systematically. As a result, creators in particular industries (such as journalism and art) lose their core source of income				
(C5) Environmental Harm						Large-scale AI training significantly increases the energy consumption of data centers, affecting the nation's overall

Risk Type	Causing Serious Harm (with reference to the legislative explanation of Article 17 and to Paragraph 1, Article 5, of the Artificial Intelligence Basic Act)					
	National Security	Fundamental Rights of the People	Life Safety	Property Protection	Social Order	Ecological Environment
						carbon emission targets
(C6) Cognitive Warfare and Information Sovereignty	AI disseminates disinformation on a large scale, systematically affecting public-policy discussions and democratic decision-making processes		AI disseminates erroneous emergency-response guidance or public-health information, causing casualties among the public	AI-assisted manipulation of public opinion triggers market panic or misguided investment momentums, causing substantial losses to investors	The combined use of deepfake technology and automated manipulation of public opinion systematically erodes the foundation of social trust	