



CHAROEN POKPHAND FOODS PUBLIC COMPANY LIMITED
บริษัท เจริญโภคภัณฑ์อาหาร จำกัด (มหาชน)

Cybersecurity Risk Management Policy for AI

นโยบายการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ สำหรับระบบ AI





**นโยบายการบริหารความเสี่ยง
ด้านความมั่นคงปลอดภัยไซเบอร์ สำหรับระบบ AI**
บริษัท เจริญโภคภัณฑ์อาหาร จำกัด (มหาชน) และบริษัทย่อย

Cybersecurity Risk Management Policy for AI
Charoen Pokphand Foods Public Company Limited and Subsidiaries

วัตถุประสงค์ (Objectives)

นโยบายการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ สำหรับระบบ AI ของ บริษัท เจริญโภคภัณฑ์อาหาร จำกัด (มหาชน) ฉบับนี้ออกภายใต้นโยบายการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ (Cyber Risk Management Policy) ของบริษัท โดยมีวัตถุประสงค์เพื่อกำหนดมาตรฐานขั้นต่ำในการกำกับการบริหารจัดการความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ในการพัฒนาและการทำงานของระบบ AI อย่างมีประสิทธิภาพ ประสิทธิภาพ สอดคล้องกับเป้าหมายทางธุรกิจของบริษัท และ/หรือบริษัทย่อย ให้มีความปลอดภัยด้านไซเบอร์ และสามารถจำกัดความเสี่ยงและผลกระทบจากภัยคุกคามไซเบอร์ในรูปแบบต่างๆ ที่อาจส่งผลกระทบต่อความต่อเนื่องในการดำเนินธุรกิจ ด้านการเงิน ความเชื่อมั่น และ ชื่อเสียงของบริษัท และ/หรือบริษัทย่อยได้ รวมทั้งเพื่อเป็นแนวทางเดียวกันให้ผู้บริหารและพนักงานยึดถือและปฏิบัติ โดยนโยบายฉบับนี้ยังสอดคล้องกับหลักการกำกับดูแลระบบ AI ที่ได้รับการยอมรับในระดับสากลได้แก่ หลักการใช้ระบบ AI อย่างรับผิดชอบ มีจริยธรรม โปร่งใส และ เป็นธรรม (FEAT: Fairness, Ethics, Accountability and Transparent) นอกจากนี้นโยบายฉบับนี้มีการกำหนดบทบาทหน้าที่ที่ชัดเจนในการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์สำหรับระบบ AI โดยครอบคลุมบทบาทหน้าที่ของ แนวป้องกัน 3 ชั้น (Three Lines of Defense) และครอบคลุมถึงการสื่อสารให้ผู้บริหารและพนักงานทุกระดับได้ตระหนักถึงเรื่องการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์สำหรับระบบ AI ภายใต้หน้าที่และความรับผิดชอบของตนเองอีกด้วย

The Cybersecurity Risk Management Policy for AI of Charoen Pokphand Foods Public Company Limited (CPF) is established under the Company's Cyber Risk Management Policy. The purpose of this policy is to prescribe minimum standards for governing cybersecurity risk management in the development and use of AI systems, ensuring that such activities are conducted efficiently and effectively and are aligned with the business objectives of the Company and its subsidiaries. This policy aims to ensure adequate cybersecurity protection and to mitigate risks and impacts arising from various forms of cyber threats that may affect business continuity, financial performance, stakeholder confidence, and the reputation of the Company and its subsidiaries. It also serves as a common framework and guideline for executives and employees to consistently adopt and implement across the organization. Furthermore, this policy is aligned with internationally recognized AI governance principles and frameworks, including the principles of Fairness, Ethics, Accountability, and Transparency (FEAT), promoting the responsible, ethical, transparent, and trustworthy use of AI technologies. This policy also clearly defines the roles and responsibilities for managing cybersecurity risks associated with AI systems. It encompasses the responsibilities of the Three Lines of Defense and includes communication and awareness initiatives to ensure

that executives and employees at all levels understand the importance of cybersecurity risk management for AI systems and are aware of their respective duties and responsibilities in this regard.

ขอบเขต (Scope)

นโยบายฉบับนี้มีขอบเขตครอบคลุมถึงบริษัทฯ และ/หรือบริษัทย่อยที่ไม่ได้จดทะเบียนในตลาดหลักทรัพย์ โดยครอบคลุมทุกกระบวนการ ทุกกลุ่มธุรกิจ/หน่วยธุรกิจ ทั้งในและต่างประเทศ ตลอดจนครอบคลุมการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ที่สำคัญสำหรับระบบ AI ทั้งนี้บริษัทย่อยที่จดทะเบียนในตลาดหลักทรัพย์และบริษัทย่อยของบริษัทย่อยที่จดทะเบียนในตลาดหลักทรัพย์สามารถนำนโยบายฉบับนี้ไปปรับใช้ให้เข้ากับบริบททางธุรกิจและกระบวนการบริหารภายในของแต่ละบริษัทบนพื้นฐานของภูมิสังคมของแต่ละประเทศให้มีความสอดคล้องกัน

This Policy applies to the Company and/or its non-listed subsidiaries. It covers all business groups, business units and processes, both domestic and international, as well as significant cybersecurity risk management activities related to AI systems. Listed subsidiaries and subsidiaries of listed subsidiaries may adapt this Policy to align with their specific business contexts and internal management processes, taking into consideration the socio-cultural environments of the countries in which they operate, while maintaining consistency with the principles and requirements set forth in this Policy.

นิยาม (Definitions)

บริษัทฯ หมายถึง บริษัท เจริญโภคภัณฑ์อาหาร จำกัด (มหาชน)

The Company means Charoen Pokphand Foods Public Company Limited.

บริษัทย่อย หมายถึง บริษัทย่อยตามกฎหมายว่าด้วยหลักทรัพย์และตลาดหลักทรัพย์ และเป็นบริษัทย่อยที่ปรากฏในงบการเงินของบริษัท แต่ไม่รวมถึง

(ก) บริษัทย่อยที่มีหุ้นสามัญจดทะเบียนในตลาดหลักทรัพย์ใดๆ

(ข) บริษัทย่อยของบริษัทตาม (ก)

Subsidiary means a subsidiary as defined under the securities and exchange laws and included in the Company's financial statements, excluding:

(a) any subsidiary whose ordinary shares are listed on any stock exchange; and

(b) any subsidiary of a subsidiary referred to in paragraph (a).

ผู้บริหาร หมายถึง ผู้บริหารของบริษัท เจริญโภคภัณฑ์อาหาร จำกัด (มหาชน) และ/หรือบริษัทย่อย

Management means the Management of Charoen Pokphand Foods Public Company Limited and/or subsidiaries.

พนักงาน หมายถึง พนักงานที่ได้รับค่าจ้างเป็นรายเดือนหรือรายวันของบริษัทฯ และ/หรือบริษัทย่อย ไม่ว่าจะเป็นการว่าจ้างประจำหรือชั่วคราว หรือการว่าจ้างด้วยสัญญาพิเศษ

Employee means an employee who is paid on a monthly or daily basis by the Company and/or its subsidiaries, whether employed on a permanent, temporary, or special contractual basis.

ปัญญาประดิษฐ์ หมายถึง ระบบที่เลียนแบบปัญญาของมนุษย์ที่สามารถทำงานต่างๆ เช่น เรียนรู้ จัดจำ ตัดสินใจ ปฏิบัติงาน และสร้างสรรค์เนื้อหาใหม่ได้ ผ่านกลไกการเรียนรู้จากข้อมูลการสร้างเงื่อนไขอัตโนมัติ โดยครอบคลุมถึงระบบที่มีส่วนประกอบของโมเดลประเภท Machine Learning, Deep Learning ตลอดจน Generative AI เช่น Large Language Model (LLM) และระบบขั้นสูง เช่น Agentic AI ทั้งนี้ ไม่ครอบคลุมระบบอัตโนมัติที่มนุษย์เป็นผู้กำหนดขั้นตอนการทำงานเช่น Robotic Process Automation (RPA) หรือระบบที่ทำงานอัตโนมัติเมื่อเข้าเงื่อนไขที่กำหนดล่วงหน้า (Condition Matching) เป็นต้น

Artificial Intelligence (AI) means a system that simulates human intelligence and capable of performing tasks such as learning, memorization, decision-making, task execution, and the generation of new content through data-driven learning and automated rule-generation mechanisms.

AI includes systems incorporating Machine Learning, Deep Learning, and Generative AI models, such as Large Language Models (LLMs), as well as advanced AI systems, such as Agentic AI.

AI does not include automation systems in which operational processes are explicitly predefined by humans, such as Robotic Process Automation (RPA), or systems that operate automatically based solely on predefined conditions (Condition Matching).

ทรัพย์สินของระบบ AI หมายถึง ทรัพย์สินที่เกี่ยวข้องกับการพัฒนาและใช้งาน AI ซึ่งสามารถระบุเจ้าของและลงทะเบียนได้ เช่น โมเดล AI, ข้อมูลฝึกสอน, โค้ด, ระบบ และสิทธิการใช้งาน

AI Asset Inventory means a collection of assets related to the development and use of AI that can be identified, assigned to an owner, and registered, including AI models, training datasets, source code, systems, and software licenses.

ความเสี่ยง หมายถึง เหตุการณ์ที่มีความไม่แน่นอนใดๆ ซึ่งเมื่อเกิดขึ้นจะส่งผลกระทบในเชิงลบหรือเป็นอุปสรรคต่อการบรรลุ วิสัยทัศน์ พันธกิจ เป้าหมาย และกลยุทธ์ทางธุรกิจของบริษัท หรือกลุ่มธุรกิจ/หน่วยธุรกิจ

Risk means any event involving uncertainty that, if it occurs, may adversely affect or impede the achievement of the Company's or a business group's/business unit's vision, mission, objectives, and business strategies.

ความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์สำหรับระบบ AI หมายถึง ความเสี่ยงจากเหตุการณ์ที่อาจเกิดขึ้นแล้ว ส่งผลกระทบต่อข้อมูลที่สำคัญ การเงิน หรือการหยุดชะงักต่อการดำเนินธุรกิจ อันเนื่องมาจากการพัฒนา หรือการใช้งาน ปัญญาประดิษฐ์ (Artificial Intelligence: AI) ที่ไม่มีความปลอดภัยเพียงพอ หรือถูกคุกคามจากภัยไซเบอร์ซึ่งอาจส่งผลให้เกิด การละเมิดข้อมูล การขโมยข้อมูล หรือการทำลายข้อมูลเพื่อให้บริษัทไม่สามารถดำเนินธุรกิจหรือให้บริการได้

Cybersecurity Risk for AI means the risk arising from events that may adversely affect critical information, financial performance, or business operations due to the insecure development or use of Artificial Intelligence (AI) systems, or as a result of cyber threats targeting such systems. Such risks may lead to data breaches, theft, loss, or destruction, potentially preventing the Company from conducting its business operations or delivering its services.

ความเสี่ยงของการใช้งานระบบ AI สามารถเกิดขึ้นได้ในกระบวนการเรียนรู้ตั้งแต่ขั้นตอนการเตรียมข้อมูลและการพัฒนาโมเดล ตลอดจนการนำระบบ AI ไปใช้งาน โดยจำแนกความเสี่ยงออกเป็น 3 ด้าน ดังนี้

Risks associated with the use of AI systems may arise throughout the AI lifecycle, from data preparation and model development to the deployment and operation of AI systems. Such risks can be categorized into the following three areas:

1. **ความเสี่ยงด้านข้อมูล** หมายถึง ความเสี่ยงที่ข้อมูลสำหรับการเรียนรู้ของโมเดลไม่มีคุณภาพ เช่น ไม่เป็นปัจจุบัน ไม่ถูกต้อง ไม่หลากหลายเพียงพอ ส่งผลให้โมเดลสร้างผลลัพธ์ที่ไม่น่าเชื่อถือ เชิงลบ โน้มเอียง หรือเป็นเท็จ นอกจากนี้ การรักษาความลับ และความมั่นคงปลอดภัยของข้อมูลในขั้นตอนการเรียนรู้หรือการใช้งานระบบ AI ที่ไม่รัดกุมเพียงพอ อาจส่งผลให้ข้อมูลสำคัญของบริษัทหรือข้อมูลส่วนบุคคลของบริษัท หรือ ลูกค้า รั่วไหลไปภายนอกได้

Data Risk means the risk arising from data quality deficiencies used for model training, such as data that is outdated, inaccurate, incomplete, or insufficiently diverse, which may result in unreliable, biased, misleading, or inaccurate outputs generated by AI models. In addition, inadequate data confidentiality and security controls during the training, development, or use of AI systems may lead to the unauthorized disclosure or leakage of the Company's confidential information, personal data, or customer information.

2. **ความเสี่ยงด้านการพัฒนาโมเดล** หมายถึง ความเสี่ยงที่โมเดลอาจเรียนรู้และสร้างเงื่อนไขที่ไม่สามารถอธิบายที่มาของผลลัพธ์ได้ รวมทั้งอาจสร้างผลลัพธ์ที่ไม่ถูกต้องไม่น่าเชื่อถือ เช่น ผลลัพธ์ที่มีความโน้มน้าวใจ ไม่แน่นอน ไม่แม่นยำหรือให้ข้อมูลที่เกินจริง (Hallucination) จนส่งผลให้เกิดความเสี่ยงต่อการดำเนินธุรกิจ ความเสี่ยงด้านการปฏิบัติงาน และความเสี่ยงที่การใช้งานหรือให้บริการลูกค้าอาจไม่สอดคล้องตามแนวทางการให้ดำเนินการของบริษัทอย่างเป็นธรรม

Model Risk means the risk that an AI model may learn patterns and generate decision rules in a manner that does not allow the basis of its outputs to be clearly explained. Such models may also produce incorrect or unreliable outputs, including biased, inconsistent, inaccurate, or false information (hallucinations). These risks may result in adverse impacts on business operations, operational risk exposure, and the risk that AI-enabled services or customer interactions may not be conducted in accordance with the Company's principles of fairness and responsible business practices.

3. **ความเสี่ยงที่เกิดจากภัยคุกคามทางไซเบอร์** หมายถึง ความเสี่ยงจากการโจมตีรูปแบบใหม่ที่เจาะจงกับระบบ AI เพื่อเข้าถึงและให้ระบบแสดงข้อมูลสำคัญโดยไม่ได้รับ อนุญาต หรือมุ่งประสงค์ให้ระบบประมวลผลผิดพลาดหรือบิดเบือนผลลัพธ์ รวมทั้ง เพื่อค้นหาช่องโหว่ด้านความปลอดภัย โดยใช้รูปแบบการโจมตีต่าง ๆ เช่น การโจมตี ระบบผ่านคำสั่ง (Prompt Injection) การเข้าถึงโครงสร้างการทำงานของโมเดล (Model Inversion) การแทรกหรือแก้ไขข้อมูลสำหรับเรียนรู้ระหว่างพัฒนาโมเดล (Data Poisoning) การป้อนข้อมูลปรับแต่งที่มีวัตถุประสงค์ให้โมเดล AI ทำงาน ผิดพลาดระหว่างใช้งาน (Adversarial Attack) เป็นต้น

Cybersecurity Risk means the risk arising from cyber-attacks specifically targeting AI systems, which may result in unauthorized access to sensitive information, the disclosure of confidential information, manipulation of system processing, distortion of outputs, or the exploitation of security vulnerabilities. Such attacks may include Prompt Injection, Model Inversion, Data Poisoning, and Adversarial Attacks, among others.

หน้าที่และความรับผิดชอบ (Roles and responsibilities)

คณะกรรมการบริษัท

ดูแลและสนับสนุนให้มีการดำเนินการบริหารความเสี่ยงด้านการใช้เทคโนโลยีและความมั่นคงปลอดภัยไซเบอร์ ผ่านทางคณะกรรมการเทคโนโลยีและความมั่นคงปลอดภัยทางไซเบอร์ (Technology and Cyber Security Committee)

Board of Directors

Ensure the governance and support of technology and cybersecurity risk management initiatives through the designated Technology and Cyber Security Committee.

คณะกรรมการเทคโนโลยีและความมั่นคงปลอดภัยทางไซเบอร์

กำกับดูแลด้านการใช้เทคโนโลยีและความมั่นคงปลอดภัยไซเบอร์ อย่างมีประสิทธิภาพและประสิทธิผล สอดคล้องกับเป้าหมายทางธุรกิจ และสนับสนุนคณะกรรมการบริษัทในการกำหนดนโยบาย และกลยุทธ์ และการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ของกลุ่มบริษัท

Technology and Cyber Security Committee

To develop the risk management policy, define acceptable levels of risk, and prepare risk management guidelines and manuals; to promote awareness and understanding of risk management across the organization; to provide advice and support to all departments in implementing risk management practices; and to monitor and follow up on the progress of risk management activities.

สำนักบริหารความเสี่ยงองค์กร

พัฒนานโยบายการบริหารความเสี่ยง ความเสี่ยงที่ยอมรับได้ และคู่มือการบริหารความเสี่ยง รวมทั้งส่งเสริมและเผยแพร่ความรู้เกี่ยวกับการบริหารความเสี่ยง ตลอดจนให้คำแนะนำแก่หน่วยงานต่างๆ ในการดำเนินการบริหารความเสี่ยง พร้อมทั้งติดตามความคืบหน้าของการบริหารความเสี่ยง

Risk Management Office

To develop risk management policies, determine acceptable risk levels, and establish comprehensive risk management manuals. To actively promote and disseminate risk management knowledge across the organization. To provide recommendations and support to relevant departments in executing risk management activities, and to monitor and follow up on the implementation progress to ensure effectiveness and alignment with organizational objectives.

ผู้บริหารหน่วยงานบริหารความเสี่ยงด้านความมั่นคงปลอดภัยสารสนเทศ

ดูแลและกำกับให้มีการดำเนินการเป็นไปตามนโยบายและมาตรฐานการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ มีหน้าที่บริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ตามหลักความรับผิดชอบซึ่งเทียบเท่า ผู้บริหารระดับสูงด้านการ

รักษาความปลอดภัยสารสนเทศ (CISO: Chief Information Security Officer) ตามแนวทางมาตรฐานในการจัดการด้านความมั่นคงปลอดภัยไซเบอร์สากล

Information Security Risk Management

To oversee and ensure the implementation of cybersecurity risk management in accordance with the established policy. The role shall include responsibility for managing cybersecurity risks at a level equivalent to those of the Chief Information Security Officer (CISO), in accordance with international standards and best practices for cybersecurity management.

หน่วยงานบริหารความเสี่ยงด้านความมั่นคงปลอดภัยสารสนเทศ

พัฒนานโยบายและมาตรฐานการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ รวมทั้งส่งเสริมและเผยแพร่ความรู้และความตระหนักเกี่ยวกับการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ ตลอดจนให้คำแนะนำแก่หน่วยงานต่างๆ ในการดำเนินการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ พร้อมทั้งติดตามความคืบหน้าของการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์และรายงานต่อคณะกรรมการเทคโนโลยีและความมั่นคงปลอดภัยทางไซเบอร์ และรายงานต่อคณะกรรมการบริหารความเสี่ยงตามรอบที่กำหนด หรือเมื่อมีเหตุการณ์ความเสี่ยงด้านความมั่นคงปลอดภัยทางไซเบอร์ในระดับสำคัญ

Information Security Risk Office

To develop policies on cybersecurity risk management and to promote and disseminate knowledge and awareness regarding cybersecurity risks throughout the organization. To provide advice and support to relevant departments in implementing cybersecurity risk management practices, to monitor the progress of risk mitigation efforts, and report to the Technology and Cybersecurity Committee, as well as to the Risk Management Sub-Committee in accordance with the defined reporting schedule or in response to significant cybersecurity risk events.

หน่วยงานเทคโนโลยีสารสนเทศด้านโครงสร้างพื้นฐานระบบและความปลอดภัยทางไอที

พัฒนานโยบายและมาตรฐานด้านการรักษาความปลอดภัยข้อมูลสารสนเทศ และบริหารจัดการด้านความมั่นคงปลอดภัยสารสนเทศของบริษัทฯ ให้ดำเนินไปอย่างมีประสิทธิภาพ สอดคล้องกับเป้าหมายทางธุรกิจในอนาคต และสอดคล้องกับระดับความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ในปัจจุบัน และรายงานการประเมินผลการปฏิบัติงานตามนโยบายต่อคณะกรรมการเทคโนโลยีและความมั่นคงปลอดภัยทางไซเบอร์ (Technology and Cyber Security Committee)

Infrastructure & Cyber Security– CPF IT Center

To develop information security policies and standards and to ensure the effective management of the Company's information security operations in alignment with future business objectives and the current cybersecurity risk

landscape. The performance in accordance with the established policies shall be evaluated and reported to the Technology and Cyber Security Committee.

หน่วยงานงานเทคโนโลยีสารสนเทศด้านการส่งมอบผลิตภัณฑ์ IoT เพื่อธุรกิจ

พัฒนาหรือการนำเอาระบบ AI มาใช้เพิ่มประสิทธิภาพการดำเนินธุรกิจ ใช้ในการปฏิบัติงาน หรือให้บริการลูกค้าภายใต้กรอบการกำกับดูแลที่เหมาะสม ตลอดจนให้คำแนะนำแก่หน่วยงานต่างๆ ในการดำเนินการพัฒนาหรือนำเอาระบบ AI มาใช้ พร้อมทั้งดูแล วางแผน จัดการ ทรัพยากรและทรัพยากรที่เกี่ยวข้องของระบบ AI ให้สอดคล้องกับโครงสร้างและแผนการใช้งาน AI ให้ปลอดภัยของบริษัท และสอดคล้องกับระดับความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ในปัจจุบัน

IoT Business Solution– CPF IT Center

To develop and adopt AI systems to enhance business efficiency, support operational activities, and deliver customer services under an appropriate governance framework. Provide advice and guidance to relevant functions on the development and adoption of AI systems. Furthermore, oversee, plan, and manage AI-related assets and resources to ensure alignment with the Company's AI architecture and secure AI adoption plan, while maintaining consistency with the current level of cybersecurity risk.

คณะกรรมการกำกับดูแลการพัฒนาปัญญาประดิษฐ์

กำหนดแนวทาง กำกับดูแล การพัฒนา สำหรับการนำเอาระบบ AI มาใช้เพิ่มประสิทธิภาพการดำเนินธุรกิจ ใช้ในการปฏิบัติงาน หรือให้บริการลูกค้าภายใต้กรอบการกำกับดูแลที่เหมาะสม สอดคล้องกับหลักการใช้งาน AI อย่างรับผิดชอบ มีจริยธรรม โปร่งใส และ เป็นธรรม (FEAT: Fairness, Ethics, Accountability and Transparent) โดยครอบคลุมถึงการบริหารจัดการความเสี่ยงของระบบ AI ครอบคลุมถึงกระบวนการพัฒนาและการใช้งานระบบ AI (AI Lifecycle) อย่างรัดกุมเหมาะสม เพื่อให้แน่ใจว่าระบบมีความมั่นคงปลอดภัย ถูกต้อง เชื่อถือได้ สามารถควบคุมความเสี่ยงและอธิบายที่มาของผลลัพธ์จากระบบ AI ได้ รวมทั้งมีแนวทางในการดูแลและคุ้มครองผู้ใช้บริการอย่างโปร่งใสและเป็นธรรม ซึ่งผลการประเมินความเสี่ยง การกำหนดมาตรการจัดการความเสี่ยง การนำไปปฏิบัติ และการติดตามและรายงานความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ รวมถึงรายงานระดับและสถานะของความเสี่ยงต่อ คณะกรรมการเทคโนโลยีและความมั่นคงปลอดภัยทางไซเบอร์ (Technology and Cyber Security Committee)

AI Committee

To define guidelines and oversee the development and adoption of AI systems to enhance business efficiency, support operational activities, and deliver customer services under an appropriate governance framework, in alignment with the principles of responsible, ethical, transparent, and fair AI use (FEAT: Fairness, Ethics, Accountability and Transparency). This includes AI risk management throughout the AI lifecycle, covering the development and use of AI systems in a prudent and appropriate manner to ensure that AI systems are secure, accurate, reliable, capable of effectively managing risks, and able to explain the basis of their outputs. It also includes establishing guidelines for the transparent and fair treatment and protection of users. Risk assessment

results, the establishment and implementation of risk mitigation measures, and the monitoring and reporting of cybersecurity risks, including the reporting of risk levels and risk status, shall be reported to the Technology and Cyber Security Committee.

ผู้บริหารของกลุ่มธุรกิจ/หน่วยธุรกิจ

รับผิดชอบโดยตรงต่อการบริหารความเสี่ยงต่างๆในขอบเขตงานที่รับผิดชอบ โดยครอบคลุมถึงการระบุและประเมินความเสี่ยง การกำหนดมาตรการจัดการความเสี่ยง การนำไปปฏิบัติ และการติดตามการดำเนินการและผลจากการตัดสินใจของระบบ AI และรายงานความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ รวมถึงรายงานระดับและสถานะของความเสี่ยงต่อหน่วยงานบริหารความเสี่ยงด้านความมั่นคงปลอดภัยสารสนเทศ (Information Security Risk Office)

Business Group/Business Unit Management

Directly responsible for managing risks within their scope of responsibility, including the identification and assessment of risks, the establishment of risk mitigation measures, implementation, monitoring, and reporting of cybersecurity risks. This includes reporting the level and status of risks to the Information Security Risk Office.

สำนักตรวจสอบภายใน

ประเมินกระบวนการปฏิบัติงานและให้คำแนะนำเพื่อพัฒนาระบบการบริหารความเสี่ยง

Internal Audit Office

Evaluate operational processes and offer recommendations to enhance the effectiveness of the risk management framework.

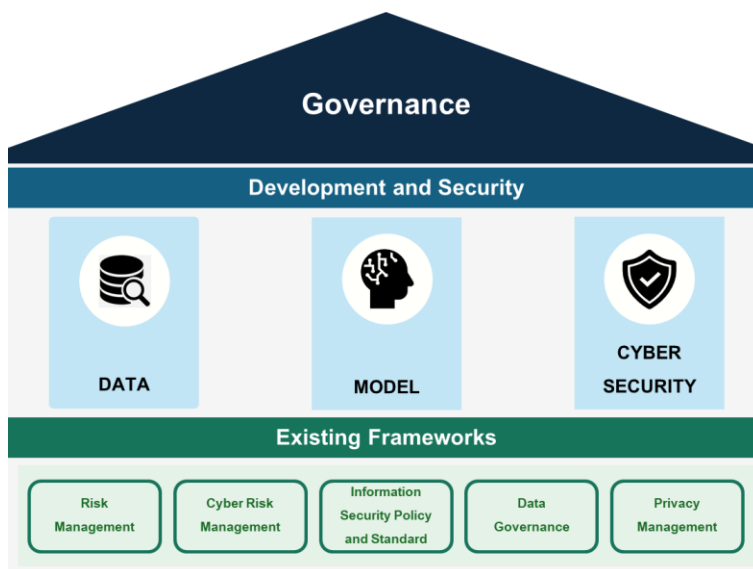
พนักงาน

ให้ความร่วมมือในการบริหารความเสี่ยง และดำเนินการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ โดยถือเป็นส่วนหนึ่งของการปฏิบัติงานที่รับผิดชอบ ปฏิบัติตนให้สอดคล้องกับวัฒนธรรมการบริหารความเสี่ยง รวมทั้งรายงานความเสี่ยงที่พบตามช่องทางที่กำหนดอย่างทันเวลา

Employees

All employees shall cooperate in the execution of risk management activities and incorporate cybersecurity risk management into their respective areas of responsibility. They are expected to demonstrate behavior consistent with the organization's risk-aware culture and are required to report any identified cybersecurity risks promptly through the appropriate reporting channels as specified by the company.

แนวปฏิบัติ (Guidelines)



นโยบายการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ สำหรับ AI ประกอบด้วย ธรรมชาติของการนำระบบ AI มาใช้งาน (Governance) และ การพัฒนาและการรักษาความมั่นคงปลอดภัยของการใช้งานระบบ AI (Development and Security)

The Cybersecurity Risk Management Policy for AI consists of AI Governance and AI Development and Security.

1. ธรรมชาติของการนำระบบ AI มาใช้งาน (AI Governance)

เพื่อให้มีการกำกับดูแลระบบ AI สอดคล้องตามหลักการใช้งานระบบ AI อย่างรับผิดชอบ มีจริยธรรม โปร่งใส และ เป็นธรรม (FEAT: Fairness, Ethics, Accountability and Transparent) รวมถึงมีการบริหารจัดการความเสี่ยงจากการพัฒนาและใช้งานระบบ AI อย่างรัดกุมเหมาะสม ผู้บริหารและคณะทำงานกำกับดูแลการพัฒนาปัญญาประดิษฐ์ (AI Steering Committee) ต้องระบุนิติทางกลยุทธ์เกี่ยวกับการนำ AI มาใช้ และแนวทางการกำกับดูแลการพัฒนาและใช้งานระบบ AI เพื่อให้เท่าทันกับพัฒนาการของเทคโนโลยีรูปแบบลักษณะการใช้งานหรือความเสี่ยงที่เปลี่ยนแปลงไป โดยการบริหารจัดการความเสี่ยงจากการพัฒนาหรือใช้งาน AI ให้พิจารณาดำเนินการดังนี้

To ensure that AI systems are governed in accordance with the principles of responsible, ethical, transparent, and fair AI use (FEAT: Fairness, Ethics, Accountability, and Transparency), and that risks arising from the development and use of AI systems are managed in a prudent and appropriate manner, management and the AI Steering Committee shall define strategic directions for AI adoption and establish governance guidelines for the development and use of AI systems to keep pace with evolving technologies, emerging use cases, and changing risk landscapes.

In managing risks arising from the development and use of AI systems, the following practices shall be implemented

- จัดทำทะเบียนการใช้งาน AI (AI Asset Inventory)

Establish and maintain an AI Asset Inventory

- ระบุความเสี่ยงจากการใช้งานระบบ AI และบริหารจัดการภายใต้กรอบนโยบายการบริหารจัดการความเสี่ยงที่ยอมรับได้ (Risk Appetite) ของบริษัท รวมทั้งจัดให้มีการประเมินและติดตาม ความเสี่ยงของการใช้งานระบบ AI อย่างต่อเนื่อง เพื่อให้มีการควบคุมความเสี่ยงเหมาะสมกับลักษณะการใช้งาน

Identify risks associated with the use of AI systems and manage such risks within the Company's approved risk appetite framework. In addition, conduct ongoing risk assessments and monitoring of AI systems to ensure that risk controls remain appropriate to the nature and extent of AI usage.

- กรณีที่มีการนำระบบ AI มาใช้กับงานที่เกี่ยวข้องกับการตัดสินใจเชิงกลยุทธ์ (Strategic Function) หรืองานที่โต้ตอบกับลูกค้า (Customer Interaction) ให้พิจารณาให้มนุษย์มีส่วนร่วมในการตัดสินใจ โดยอาจเลือกใช้วิธี Human in the Loop หรือ Human over the Loop ขึ้นกับการพิจารณาความเสี่ยงและผลกระทบของการใช้งานระบบ AI ที่มีต่อการดำเนินธุรกิจและลูกค้า

In cases where AI systems are used for Strategic Functions or Customer Interactions, consideration shall be given to involving humans in the decision-making process through Human-in-the-Loop (HITL) or Human-over-the-Loop (HOTL) approaches, depending on the risks and potential impacts of the use of AI systems on business operations and customers.

- กรณีใช้งานระบบ AI กับงานที่โต้ตอบกับผู้ใช้หรือลูกค้า ผู้พัฒนาต้องแจ้งและเปิดเผยเกี่ยวกับการให้บริการด้วยระบบ AI ต่อผู้ใช้หรือลูกค้าก่อนให้บริการ และมีตัวเลือกให้ลูกค้าสามารถปิด (Disable) หรือข้าม (Bypass) บริการที่ใช้งานร่วมกับระบบ AI ได้ รวมทั้งจัดให้มีการให้ความรู้เกี่ยวกับวิธีการใช้งานอย่างปลอดภัยแก่ผู้ใช้หรือลูกค้า

In cases where AI systems are used for interactions with users or customers, developers shall notify them and disclose the use of AI systems prior to service delivery. Such users or customers shall be provided with the option to disable or bypass services that utilize AI systems. In addition, education and guidance on the safe use of these services should be provided to them.

2. การพัฒนาและการรักษาความมั่นคงปลอดภัยของการใช้งานระบบ (AI Development and Security)

เพื่อให้มีการควบคุมความเสี่ยงข้อมูลที่นำมาใช้ในการเรียนรู้ การพัฒนาโมเดล ระบบ AI ให้มีความถูกต้อง เชื่อถือได้ โปร่งใส และมั่นคงปลอดภัยจากภัยคุกคามไซเบอร์รูปแบบใหม่ที่เกิดขึ้น ให้ผู้พัฒนาดำเนินการดังนี้

To ensure appropriate controls over data used for training and model development, and to ensure that AI systems are accurate, reliable, transparent, and secure against emerging cyber threats, developers shall implement the following measures:

2.1. การควบคุมความเสี่ยงด้านข้อมูล (Data Risk Controls)

- มีแนวทางควบคุมและประเมินคุณภาพข้อมูลที่ใช้ในการเรียนรู้ของโมเดลโดยจัดทำหลักเกณฑ์เพื่อใช้ในการประเมินคุณภาพข้อมูลที่นำมาใช้ในการเรียนรู้ของโมเดลอย่างชัดเจน และครอบคลุมมิติด้านความถูกต้อง ความเป็นปัจจุบัน ปริมาณและความหลากหลายของข้อมูล เพื่อลดความเสี่ยงที่โมเดลอาจสร้างผลลัพธ์ที่ไม่น่าเชื่อถือ เชิงลบ ไน้มเอียง หรือเป็นเท็จ

Establish guidelines for the control and assessment of the quality of data used for model training. The guidelines shall define clear criteria covering data accuracy, timeliness, volume, and diversity to reduce the risk of AI models generating unreliable, biased, or inaccurate outputs.

- ดำเนินการตรวจสอบและประเมินคุณภาพข้อมูล รวมทั้งมีแนวทางแก้ไขข้อมูลให้เป็นไปตามหลักเกณฑ์ที่กำหนด ก่อนนำมาใช้ในกระบวนการเรียนรู้ของโมเดล

Conduct the data quality reviews and assessments and apply appropriate data remediation measures in accordance with the established guidelines before data is used in the model training process.

- มีแนวทางป้องกันข้อมูลรั่วไหลจากระบบ AI อย่างน้อยครอบคลุม

Establish guidelines for preventing data leakage from AI systems, covering at a minimum:

- การจำกัดสิทธิ์ของระบบ AI ในการเข้าถึงข้อมูล (Access Control) ของบริษัท เพื่อป้องกันการเข้าถึงข้อมูลสำคัญโดยไม่จำเป็น ตลอดจน ควบคุมสิทธิการเข้าถึงข้อมูลสำหรับเรียนรู้และโครงสร้างการทำงานของโมเดล

Restrict AI systems' access to Company data through appropriate access controls to prevent unnecessary access of sensitive information and strictly manage access rights to training data and model architecture.

- การปกป้องข้อมูลส่วนบุคคลหรือข้อมูลอ่อนไหว ตั้งแต่การเตรียม ข้อมูล การพัฒนาโมเดล และ การใช้งานระบบ AI เช่น ทำ Data Masking, Data Hashing, Input Sanitization, เพื่อป้องกันการรั่วไหลของข้อมูลดังกล่าว

Protect personal and sensitive data throughout the data preparation, model development, and AI system usage phases by implementing measures such as data masking, data hashing, and input sanitization to prevent data leakage.

- การแยกข้อมูลบริษัทออกจากข้อมูลสาธารณะ กรณีใช้ ระบบ AI ของบุคคลภายนอก เช่น การทำ Data Boundary เพื่อแยก ข้อมูลของบริษัทในการใช้งาน Generative AI

Segregate Company data from public data when utilizing third-party AI systems; for example, by establishing logical or physical data boundaries to isolate Company data when using Generative AI services.

2.2. การควบคุมความเสี่ยงด้านการพัฒนาโมเดล (Model Risk Controls)

- มีเกณฑ์การควบคุมและประเมินประสิทธิภาพของโมเดล โดยจัดทำตัวชี้วัดประสิทธิภาพ (Evaluation Metrics) อย่างชัดเจนและครอบคลุมด้านความถูกต้องแม่นยำและด้านความน่าเชื่อถือ เหมาะสมกับความเสี่ยงและผลกระทบของลักษณะการนำไปใช้งาน

Establish controls and criteria for assessing model performance by defining clear and comprehensive evaluation metrics that address accuracy and reliability, commensurate with the risks and potential impacts associated with the intended use of the model.

- การทดสอบและติดตามประสิทธิภาพและความน่าเชื่อถือของผลลัพธ์ทั้งก่อนและหลังใช้งานอย่างต่อเนื่อง เพื่อให้มั่นใจว่าโมเดลมีความสม่ำเสมอในการให้ผลลัพธ์ที่ถูกต้องแม่นยำและน่าเชื่อถือ เป็นไปตามเกณฑ์ที่กำหนด โดยอาจทดสอบด้วยข้อมูลหลากหลายรูปแบบ เช่น ข้อมูลชุดใหม่ (Unseen Data) ข้อมูลที่มีลักษณะคลาดเคลื่อนจากข้อมูลเรียนรู้ ข้อมูลกรณีพิเศษ (Edge Case) เป็นต้น รวมทั้งอาจพิจารณาให้มีกระบวนการ Feedback loop เพื่อนำผลลัพธ์มาปรับปรุงโมเดลให้มีความถูกต้องแม่นยำและน่าเชื่อถือต่อเนื่อง

Conduct ongoing testing and monitoring of model performance and output reliability before and after deployment to ensure that the model consistently delivers accurate and reliable results in accordance with established criteria. Testing should be performed using various types of data,

including unseen data, out-of-distribution data, and edge cases. In addition, a feedback loop mechanism should be implemented to continuously improve the model's accuracy and reliability based on operational outcomes

- กรณีใช้ Generative AI ให้กำหนดแนวทางการลดความเสี่ยงจากการที่ระบบให้ข้อมูลที่เป็นเท็จ (Hallucination) เช่น ใช้เทคนิค Retrieval-Augmented Generation เพื่อช่วยให้โมเดลมีแหล่งข้อมูลสำหรับอ้างอิง หรือปรับลักษณะคำตอบให้เหมาะสม เป็นต้น

In cases where Generative AI is used, establish guidelines to reduce the risk of hallucinations, defined as instances where the system may generate false information. Such guidelines should include the use of Retrieval-Augmented Generation (RAG) techniques to provide the model with reference data sources or adjusting response characteristics as appropriate.

- กำหนดแนวทางให้ผู้พัฒนา ผู้ใช้งาน และผู้ที่เกี่ยวข้องสามารถเข้าใจและอธิบายผลลัพธ์ของระบบ AI (Explainability) ได้ เช่น การจัดทำเอกสารเพื่ออธิบายพารามิเตอร์ที่โมเดลใช้ การให้โมเดลประเภท Generative AI แสดงกระบวนการคิด (Reasoning) หรือระบุแหล่งอ้างอิง (Referencing) ควบคู่กับการแสดงผลลัพธ์ เป็นต้น

Establish guidelines to ensure that developers, users, and other relevant stakeholders are able to comprehend and explain the outputs of AI systems (Explainability). Such guidelines should include preparing documentation that explains the decision criteria or underlying logic used by the model, enabling Generative AI models to display their reasoning process, or providing references alongside the generated outputs.

- กำหนดให้มีมนุษย์คอยกำกับ ดูแล หรือแทรกแซงการทำงานของระบบ AI เพื่อให้แน่ใจว่า AI ทำงานอย่างถูกต้อง ปลอดภัย และมีจริยธรรม (Human Oversight)

Establish human oversight measures designed to maintain human supervision, monitoring, and intervention, where appropriate, to ensure that AI systems operate correctly, safely, and ethically.

2.3. การควบคุมความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์ (Cybersecurity Risk Controls)

กำหนดให้การพัฒนา ระบบ AI ให้มีแนวทางป้องกันและตรวจจับการโจมตีรูปแบบใหม่ที่เจาะจงกับระบบ AI รวมทั้งยกระดับแนวทางรักษาความมั่นคงปลอดภัยทางไซเบอร์ที่มีอยู่ให้ครอบคลุมการรับมือกับภัยดังกล่าว เช่น การทำ Content Filtering กับระบบ AI ทั้งขานำเข้าข้อมูล (Prompt Filtering) และขาแสดงผลลัพธ์ (Response Filtering) เพื่อป้องกันไม่จาระบบแสดงข้อมูลที่ไม่พึงเปิดเผยหรือข้อมูลเท็จลบ การทดสอบด้วยรูปแบบการโจมตีต่าง ๆ อย่างต่อเนื่อง เช่น ทดสอบการทำ Data Poisoning หรือ Adversarial Testing เพื่อหาข้อจำกัดและปิดช่องโหว่ของโมเดลที่ตรวจพบ นอกจากนี้ ผู้ให้บริการทางการเงินควรติดตามและประเมินภัยคุกคามใหม่ ๆ อย่างต่อเนื่อง พร้อมปรับปรุงแนวทางป้องกันและตรวจจับให้ทันสมัยและเหมาะสมกับสถานการณ์โดยอาจอ้างอิงจากมาตรฐาน ได้แก่

- OWASP (Open Worldwide Application Security Project) Top 10
- OWASP Machine Learning Security Top 10
- OWASP Top 10 for Large Language Model Applications

Require AI system development to incorporate measures for preventing and detecting emerging attacks specifically targeting AI systems, as well as enhancing existing cybersecurity controls to address such threats. This may include implementing Content Filtering for AI systems at both the input stage (Prompt Filtering) and the output stage (Response Filtering) to prevent the disclosure of sensitive information or the generation of inappropriate content. Conduct continuous testing against various attack scenarios, such as Data Poisoning and Adversarial Testing, to identify limitations and remediate vulnerabilities discovered within AI models. In addition, continuously monitor and assess emerging threats, and update prevention and detection measures to ensure that they remain current and appropriate to the prevailing threat landscape. Reference may be made to relevant standards, including:

- OWASP (Open Worldwide Application Security Project) Top 10
- OWASP Machine Learning Security Top 10
- OWASP Top 10 for Large Language Model Applications

การยอมรับความเสี่ยง (Risk Acceptance)

เพื่อให้การบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์สำหรับระบบ AI มีประสิทธิภาพและประสิทธิผล สอดคล้องกับเป้าหมาย ดังนั้นหากไม่สามารถปฏิบัติตามนโยบายหรือแนวทางการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์สำหรับระบบ AI ให้ผู้บริหารจัดการโครงการ ผู้พัฒนา ผู้ดำเนินการ ของระบบ AI หรือตัวแทนของหน่วยงานที่ใช้งานต้องทำการประเมินความเสี่ยงและผลกระทบที่อาจเกิดขึ้น ทบทวนและยอมรับความเสี่ยงโดยผู้บริหารของกลุ่มธุรกิจ/หน่วยธุรกิจเป็นลายลักษณ์อักษรถึงเหตุผลและผลกระทบที่อาจเกิดขึ้น และนำเสนอต่อคณะกรรมการกำกับดูแลการพัฒนาปัญญาประดิษฐ์ (AI Committee) เพื่อทราบ

To ensure that cybersecurity risk management for AI is carried out effectively and efficiently and remains aligned with the Company's objectives, where compliance with this Policy or the AI Cybersecurity Risk Management Guidelines is not feasible, the AI project manager, developer, operator, or representative of the business unit utilizing the AI system shall assess the risks and potential impacts. Such risks shall be reviewed and formally accepted in writing by the management of the relevant business group or business unit, including the justification for the non-compliance and the associated impacts. The risk assessment and risk acceptance shall then be reported to the AI Committee for acknowledgment.

การทบทวนนโยบาย (Policy Review)

ให้หน่วยงานบริหารความเสี่ยงด้านความมั่นคงปลอดภัยสารสนเทศ (Information Security Risk Office) ทบทวนนโยบายฉบับนี้เป็นประจำทุกปี หรือก่อนหน้านั้นตามสมควร หากพบว่านโยบายฉบับนี้ ไม่เหมาะสมกับสภาพการดำเนินธุรกิจของบริษัทฯ และ/หรือบริษัทย่อย หรือไม่สอดคล้องกับนโยบายการบริหารความเสี่ยง โดยให้หน่วยงานบริหารความเสี่ยงด้านความมั่นคงปลอดภัยสารสนเทศนำเสนอแนะนโยบายฉบับใหม่ ที่ได้มีการทบทวนและปรับปรุงแก้ไขดังกล่าวต่อคณะกรรมการเทคโนโลยีและความมั่นคงปลอดภัยทางไซเบอร์ (Technology and Cyber Security Committee) เพื่อพิจารณาก่อนนำเสนอต่อประธานคณะผู้บริหารเพื่ออนุมัติและนำไปใช้บังคับต่อไป

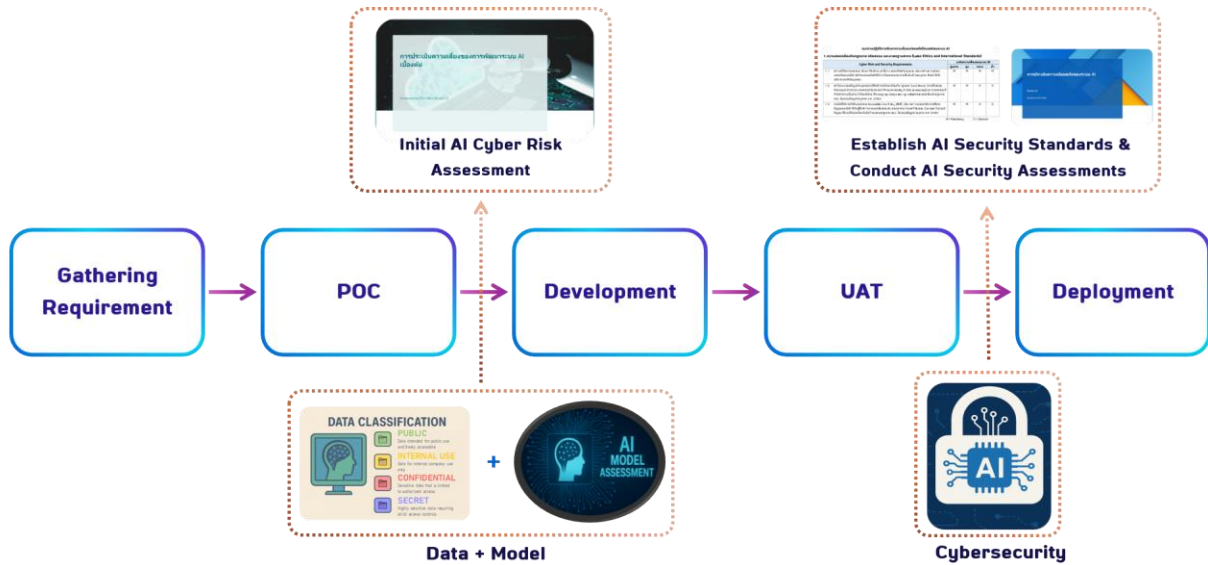
The Information Security Risk Office shall review this Policy at least annually, or more frequently as deemed necessary, to ensure ongoing alignment with the Risk Management Policy and appropriateness for the business operations of the Company and its subsidiaries. Any proposed amendments shall be submitted to the Technology and Cyber Security Committee for review and recommendation, and subsequently to the Chief Executive Officer for approval prior to implementation.

นโยบายการบริหารความเสี่ยงด้านความมั่นคงปลอดภัยไซเบอร์สำหรับระบบ AI (Cyber Risk Management Policy for AI) ฉบับนี้ผ่านการอนุมัติโดยประธานคณะผู้บริหาร เมื่อวันที่ 15 ตุลาคม 2568

The Cyber Risk Management Policy for AI was approved by the Chief Executive Officer on 15 October 2025.

ภาคผนวก (Appendix)

ภาพรวมกระบวนการประเมิน ความคุม และติดตามความเสี่ยงด้าน AI (High Level of AI Cyber Risk Assessment Flow)



* In process of continuous improvement