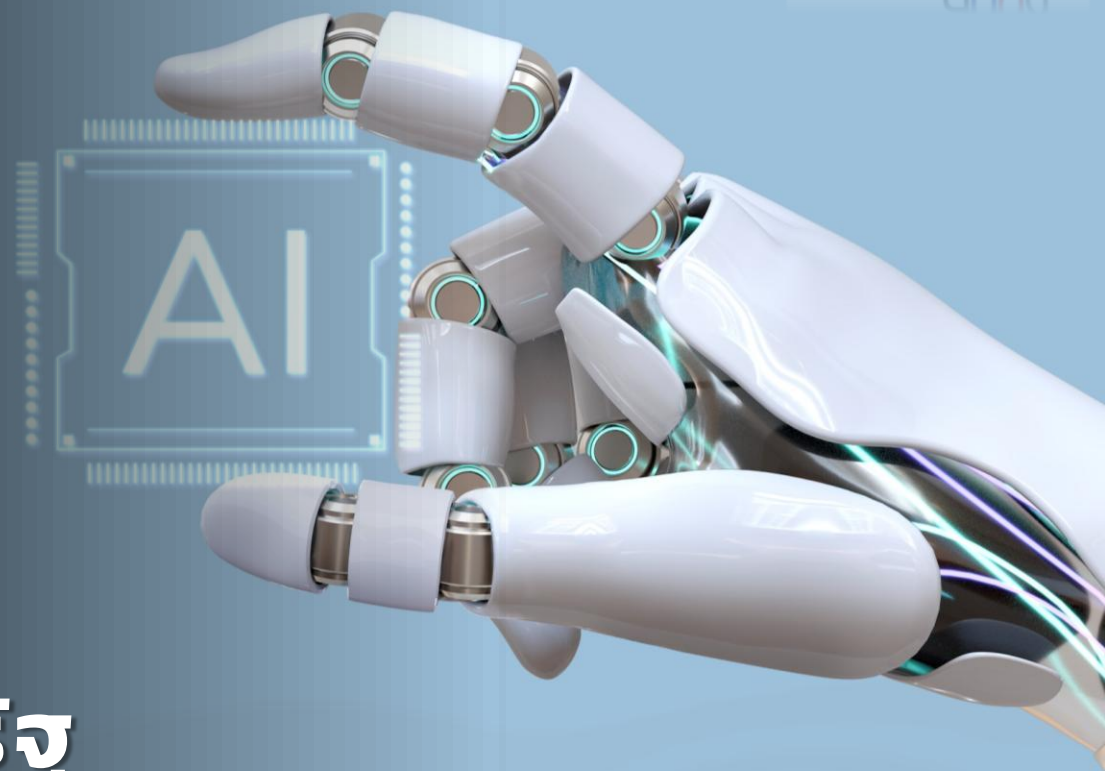


Leveraging AI to Detect and Prevent Fraud in Public Sector

AI จะช่วยให้การบริหารจัดการ
ป้องกัน ตรวจจับ ความไม่ชอบ
มาพากล ช่องโหว่ รอยร้าว ใน
ระดับการทำงานชั้นต่างๆ ภาครัฐ



ว่าที่ รต. ภูวิช ชัยกรเรีงเดช

ผู้อำนวยการฝ่ายกำหนดมาตรฐาน

สำนักงานคณะกรรมการการรักษาความมั่นคงปลอดภัยไซเบอร์แห่งชาติ



ประสบการณ์การทำงาน

• สกมช.

ผู้อำนวยการฝ่ายสืบสวนและตรวจพิสูจน์หลักฐานทางไซเบอร์

• ธนาคารกรุงไทย

รองผู้อำนวยการฝ่ายความมั่นคงปลอดภัยเทคโนโลยีสารสนเทศ

CSIRT Head (Cyber Security Incident Response Team)

พยานผู้เชี่ยวชาญทางด้านมาตรการความปลอดภัย

คดีมีงาช้างพลกหลวง ลักลอบทำธุรกรรมทางการเงิน

• บริษัท กรุงไทยคอมพิวเตอร์เซอร์วิส

ผู้จัดการฝ่ายบริหารจัดการเทคนิค (Security Team)

• ธนาคารออมสิน

พนักงานความปลอดภัยไอที (ที่ปรึกษาและตรวจพิสูจน์หลักฐาน)

ประกาศนียบัตรและการอบรมทางไซเบอร์

- CertiProf Cyber Security Foundation Professional Certificate (CSFPC)
- EC-Council Computer Hacking Forensic Investigator v9 (CHFI)
- CompTIA Security+
- Cybersecurity for Managers: A Playbook (MIT Sloan School of Management)
- SANS FOR508: Advanced Digital Forensics, Incident Response, and Threat Hunting
- SANS SEC504: Hacker Tools, Techniques, Exploits, and Incident Handling
- SANS SEC699: Purple Team Tactics - Adversary Emulation for Breach Prevention & Detection
- SANS FOR578: Cyber Threat Intelligence
- SANS FOR498: Digital Acquisition and Rapid Triage
- SANS FOR608: Enterprise-Class Incident Response and Threat Hunting



ว่าที่ ร.ต. ภูวิช ชัยกรเรีงเดช
Acting Sub Lt. PHUWICH CHAIYAKORNROENGDET
ผู้อำนวยการฝ่ายกำหนดมาตรฐาน
สำนักบริการโครงสร้างพื้นฐาน (สกมช.)
การศึกษา วิทยาศาสตร์มหาบัณฑิต (คอมพิวเตอร์)
ครุศาสตร์อุตสาหกรรมบัณฑิต (วิศวกรรมโทรคมนาคม)



สำนักงานคณะกรรมการการรักษาความมั่นคงปลอดภัยไซเบอร์แห่งชาติ



มีตัวย่อว่า **สทศ**. และมีชื่อภาษาอังกฤษ
"The National Cyber Security Agency (NCSA)"

เป็นองค์การมหาชนที่จัดตั้งขึ้นตามพระราชบัญญัติการรักษาความมั่นคงปลอดภัยไซเบอร์ พ.ศ. 2562

มีนายกรัฐมนตรีรักษาการตามพระราชบัญญัติ และมีอำนาจกำกับดูแล
 โดยทั่วไปซึ่งกิจการของสำนักงานให้เป็นไปตามหน้าที่และอำนาจ

มีหน้าที่และอำนาจตาม พ.ร.บ. ไซเบอร์ฯ ที่กำหนดให้ **มีมาตรการในการ ป้องกัน รับมือ และลดความเสี่ยง** จากภัยคุกคามทางไซเบอร์ รวมทั้งสร้างเครือข่ายความร่วมมือ และการสร้างความรู้ความเข้าใจให้แก่ทุกภาคส่วน

ขอบเขตการดำเนินการ

สำนักงานคณะกรรมการการรักษา
ความมั่นคงปลอดภัยไซเบอร์แห่งชาติ
(สทช.)

IDENTIFY | PROTECT | **DETECT** | RESPONSE | RECOVER

- ฝ้าระวัง วิเคราะห์และตรวจจับภัยคุกคามไซเบอร์
- รับมือ ประสานงาน และจัดการกับสถานการณ์ภัยคุกคามไซเบอร์
- วิเคราะห์ภัยคุกคามไซเบอร์ที่ซับซ้อนและมัลแวร์
- ตรวจสอบประเมินความมั่นคงปลอดภัยของระบบ
- แลกเปลี่ยนและเผยแพร่ข้อมูลภัยคุกคามทางไซเบอร์
- ประเมินความเสี่ยงและวางแผนความต่อเนื่องของบริการสำคัญ
- พัฒนา อบรม บุคลากรทางไซเบอร์
- สร้างความตระหนักรู้ด้านไซเบอร์



Critical Information Infrastructure: CII

ด้านความมั่นคง
ของรัฐ

ด้านเทคโนโลยีสารสนเทศ
และโทรคมนาคม

ด้านพลังงาน
และสาธารณูปโภค

ด้านสาธารณสุข

ด้านการเงิน
การธนาคาร

ด้านบริการภาครัฐที่
สำคัญ

ด้านการขนส่ง
และโลจิสติกส์

ด้านอื่นตามที่คณะกรรมการ
ประกาศกำหนดเพิ่มเติม

AI in Cybersecurity





การตรวจจับ (Detection) และ มอนิเตอร์ (Monitoring) ภัยคุกคามทางไซเบอร์



1



2



3



4

ตรวจจับพฤติกรรมผิดปกติ (Anomaly Detection)

AI ใช้อัลกอริทึม Machine Learning เพื่อวิเคราะห์พฤติกรรมของระบบหรือผู้ใช้ เช่น ปริมาณทราฟฟิก การล็อกอิน การเรียกใช้งานระบบ

หากตรวจพบพฤติกรรมที่ "เบี่ยงเบน" จากรูปแบบปกติ เช่น การเข้าถึงไฟล์สำคัญนอกเวลาทำการ หรือการส่งข้อมูลจำนวนมากผิดปกติ ระบบสามารถแจ้งเตือนทันที

วิเคราะห์ข้อมูลแบบเรียลไทม์ (Real-Time Threat Monitoring)

AI สามารถสแกน Log หรือ Data Stream จากไฟร์วอลล์, เซิร์ฟเวอร์, หรือ endpoint ต่าง ๆ ได้ตลอดเวลา

ระบบสามารถ ระบุสัญญาณของภัยคุกคาม (Indicators of Compromise - IoC) เช่น การโจมตีแบบ DDoS, ransomware หรือการสแกนพอร์ตผิดปกติ

ความสามารถด้าน Threat Intelligence

AI เชื่อมโยงข้อมูลภัยคุกคามจากหลายแหล่งทั่วโลก (Threat Feed) เพื่อตรวจสอบว่าเหตุการณ์ในระบบมีความเชื่อมโยงกับภัยที่รู้จักหรือไม่

ระบบสามารถ วิเคราะห์ ความเชื่อมโยงของภัย (correlation) ระหว่างเหตุการณ์ในระบบและภัยคุกคามจากแหล่งภายนอก

ความสามารถในการเรียนรู้ (Self-Learning)

AI พัฒนาความแม่นยำในการตรวจจับด้วย Deep Learning / Reinforcement Learning

ยิ่งระบบมีข้อมูลมากขึ้นเท่าใด ก็ยังสามารถแยกแยะภัยคุกคามที่แท้จริงออกจาก "false positive" ได้ดีขึ้น



ความสามารถของ AI	ผลลัพธ์ที่ได้รับจริง
การวิเคราะห์แบบ Real-time	ตรวจจับภัยคุกคามได้ภายในวินาที
ตรวจจับพฤติกรรมเบี่ยงเบน (Anomaly Detection)	ตรวจเจอภัยคุกคามรูปแบบใหม่ (Zero-day)
การเรียนรู้จากข้อมูล (Machine Learning)	ปรับตัวได้กับภัยคุกคามที่เปลี่ยนแปลง
การเชื่อมต่อข้อมูล Threat Intelligence	รับรู้ภัยคุกคามจากแหล่งภายนอกทันที
การทำ Automation Response	ลดภาระทีม SOC และเวลาการตอบสนอง
ความแม่นยำ (Accuracy)	สูงถึง 90-98% เมื่อเทรนโมเดลจากข้อมูลจริง
False Positive Rate	ลดลงอย่างชัดเจนด้วย Anomaly-based AI
เวลาตรวจจับ (Detection Time)	เร็วกว่ามนุษย์หลายเท่า (มิลลิวินาที - วินาที)
ความสามารถในการตอบสนองแบบอัตโนมัติ	สูงมาก เมื่อใช้ร่วมกับ SOAR
การปรับตัวกับภัยใหม่ (Zero-day)	ดีกว่า signature-based มาก

ข้อควรระวัง การนำ AI มาใช้งานการปฏิบัติงาน

● ข้อมูลตัวอย่างไม่ครอบคลุม

AI เรียนรู้จากข้อมูลที่ถูกป้อนเข้าไป หากข้อมูลนั้นมีอคติ เช่น ไม่ครอบคลุมทุกสถานการณ์หรือมีความไม่สมดุล AI อาจตัดสินใจผิดพลาดหรือมีอคติในการตรวจจับภัยคุกคาม

● ภัยคุกคามที่ไม่เคยพบมาก่อน

AI อาจตรวจจับภัยคุกคามใหม่ ๆ ที่ไม่เคยมีในข้อมูลฝึกสอนได้ไม่ดี เช่น การโจมตีแบบ zero-day

● การแจ้งเตือนผิดพลาด

AI อาจระบุกิจกรรมที่ไม่เป็นอันตรายว่าเป็นภัยคุกคาม (false positives) หรือมองข้ามภัยคุกคามจริง (false negatives) ซึ่งอาจสร้างปัญหาในการใช้งานจริง

เราสามารถเชื่อมั่นใน AI ได้ในระดับหนึ่งสำหรับการตรวจจับและป้องกันภัยคุกคามทางไซเบอร์ เนื่องจากมันมีศักยภาพสูงในการวิเคราะห์ข้อมูลและตอบสนองอย่างรวดเร็ว อย่างไรก็ตาม AI ไม่สามารถทำงานได้อย่างสมบูรณ์แบบโดยปราศจากอคติหรือข้อผิดพลาด เพราะมันขึ้นอยู่กับคุณภาพของข้อมูลและอาจมีปัญหากับภัยคุกคามที่ไม่เคยพบมาก่อน การทำงานร่วมกันระหว่าง AI และมนุษย์จึงเป็นสิ่งจำเป็นเพื่อให้ได้ผลลัพธ์ที่แม่นยำ ตรงไปตรงมา และปราศจากอคติมากที่สุดเท่าที่จะเป็นไปได้



Template นโยบายการใช้เทคโนโลยี Generative AI

← ↻ 🔒 https://www.ncsa.or.th/news/c4fcecd03463638c8c00036f?category=news

Font Size A A A 🔍 Search

หน้าแรก เกี่ยวกับสำนักงาน บริการของเรา เอกสารเผยแพร่ และคลังความรู้ ข่าวสาร กิจกรรม และประกาศ ติดต่อสำนักงาน

🔊 Template นโยบายการใช้เทคโนโลยี Generative AI ที่ยอมรับได้ (Acceptable Use Policy: Generative AI)

📅 วันที่ 14 มีนาคม 2568 เวลา 15:08:49

แชร์   

Template นโยบายการใช้เทคโนโลยี Generative AI ที่ยอมรับได้ (Acceptable Use Policy: Generative AI)

Download ได้ที่ : <https://www.ncsa.or.th/docpdf/view/ba37afc6636334401e000142>

นโยบายนี้จัดทำขึ้นโดยมีเนื้อหาสาระสำคัญ ที่สามารถนำไปเป็นแนวทางการจัดกิจกรรมการให้ความรู้ การทราบถึงโทษของการไม่ปฏิบัติตามนโยบาย และตัวอย่างพฤติกรรมที่ควรทำและไม่ควรทำ (Do's & Don'ts) ของการใช้งานเทคโนโลยี Generative AI รวมถึงแอปพลิเคชันหรือบริการ Generative AI ที่สำนักงานแนะนำ เพื่อให้พนักงาน ลูกจ้าง ผู้ปฏิบัติงานและผู้รับจ้างของสำนักงานสามารถนำเทคโนโลยี Generative AI ไปใช้งานได้อย่างถูกต้องเหมาะสม มีความรับผิดชอบ และสามารถนำเทคโนโลยีดังกล่าวไป รังสรรค์งานได้อย่างปลอดภัย มีประสิทธิภาพ ตลอดจนป้องกันความเสี่ยงที่อาจเกิดขึ้นจากการใช้งานที่ไม่เหมาะสม ในกรณีนี้สำนักงานจึงได้กำหนดวัตถุประสงค์ของนโยบายดังกล่าวไว้

เอกสาร Template รูปแบบ MS word นโยบายการใช้เทคโนโลยี Generative AI ที่ยอมรับได้ (Acceptable Use Policy: Generative AI)

– <https://ncsa.or.th/docpdf/view/bffc88ab366261c51300035d>

#teamthailandfor cybersecurity #NCSA

AutoSave Off    generative_ai_template_... Saved 🔍 PC

File Home Insert Design Layout References Mailings Review View Help Comments   

1 2 3 4 5 6

(ตัวอย่าง)

นโยบายการใช้เทคโนโลยี Generative AI ที่ยอมรับได้
(Acceptable Use Policy: Generative AI)
หน่วยงาน.....(โปรดระบุ).....

๑. บทนำ

ปัจจุบันเทคโนโลยี Generative AI มีความสามารถในการสร้างสรรค์เนื้อหาได้หลากหลายรูปแบบ เช่น ข้อความ ภาพ วิดีโอ เสียง ซอร์สโค้ด เป็นต้น โดยการสั่งการผ่านข้อความ หรือคำสั่ง (Prompt) ที่ผู้ใช้งานเป็นผู้กำหนด อีกทั้งยังสามารถสร้างเนื้อหาได้อย่างสมจริงและโต้ตอบกับผู้ใช้งานได้คล้ายคลึงกับมนุษย์ อันสามารถช่วยประหยัดเวลาและเพิ่มประสิทธิภาพในการดำเนินงานได้ตั้งแต่การเขียนบันทึกข้อความ การเขียนบทความ การเขียนโปรแกรม การแก้ไขปัญหาด้านเทคนิค การสร้างเอกสารนำเสนอประกอบการประชุม การสร้างสื่อประชาสัมพันธ์ หรือการสร้างเนื้อหาอื่นใดตามแต่ผู้ใช้งานจะสร้างสรรค์

อย่างไรก็ตาม เมื่อมีการนำ เทคโนโลยี Generative AI มาใช้ในการสร้างสรรค์และ เกิด ข้อมูลเท็จที่ไม่ได้ตั้งใจให้เกิดความเข้าใจผิด (Misinformation) และข้อมูลเท็จที่จงใจ ให้เกิดความเข้าใจผิด (Disinformation) กรณีเช่นนี้จะถูกจัดเป็นหนึ่งในความเสี่ยงระดับโลกที่รุนแรงที่สุด ตามรายงาน Global Risks Report 2023 - 2024 ของ World Economic Forum ที่เนื้อหาที่ได้จากการสร้างสรรค์จากเทคโนโลยี Generative AI ที่อาจทำให้ผู้ใช้งานเชื่อได้ว่าข้อมูลนั้นถูกต้อง มีความน่าเชื่อถือและสามารถนำไปใช้งานได้ทันที โดยไม่จำเป็นต้องพิจารณาถึงความเหมาะสมและข้อจำกัดของเทคโนโลยี โดยเฉพาะอย่างยิ่งข้อจำกัดทั้งด้านภาษาไทย ความเข้าใจในบริบทหรือวัฒนธรรมของประเทศไทย หรือบริบทของหน่วยงาน ซึ่งอาจก่อให้เกิดผลกระทบในเชิงลบต่อบุคคล องค์กร สังคมและประเทศชาติ เช่น

- ข้อมูลเท็จที่ไม่ได้ตั้งใจให้เกิดความเข้าใจผิด (Misinformation) และข้อมูลเท็จที่จงใจให้เกิดความเข้าใจผิด (Disinformation)
- เนื้อหาไม่ถูกต้องตามข้อเท็จจริง (Confabulation)
- เนื้อหาที่มีความอันตรายหรือรุนแรง (Dangerous or Violent Recommendations)
- เนื้อหาส่งผลกระทบต่อความเป็นส่วนตัวของข้อมูล (Data Privacy)
- เนื้อหาส่งผลกระทบต่อประเด็นที่มีความอ่อนไหว และมีเนื้อหาที่ขัดต่อหลักจริยธรรม (Sensitive and Ethics Context)
- เนื้อหาที่มีข้อมูลที่ไม่ถูกต้องครบถ้วนหรือถูกทำให้บิดเบือน (Information Integrity)
- เนื้อหาที่สร้างก่อให้เกิดอคติ แบ่งแยกและการเลือกปฏิบัติ (Bias and Discrimination)
- ข้อมูลส่วนบุคคลรั่วไหล (Personal Data Leakage)
- การละเมิดทรัพย์สินทางปัญญา (Intellectual Property Infringement)

Page 1 of 6 Thai Accessibility: Good to go Focus 90%

กรณีศึกษาจากหน่วยงานภาครัฐทั่วโลกที่นำ AI มาใช้ในการรักษาความมั่นคงปลอดภัยทางไซเบอร์

หน่วยงานความมั่นคงของสหรัฐอเมริกา (CISA)

ระบบ AI วิเคราะห์ข้อมูลเครือข่ายแบบเรียลไทม์ เช่น การโจมตีแบบ DDoS หรือการบุกรุกระบบ โดยใช้ระบบ Automated Indicator Sharing (AIS) เพื่อประเมินข้อมูลภัยคุกคามจากทั้งภาครัฐและเอกชน

หน่วยงานความมั่นคงของสหรัฐอเมริกา (CISA)

พัฒนาแดชบอร์ดแบบโต้ตอบที่ใช้ AI ในการวิเคราะห์ข้อมูลภัยคุกคามจากแหล่งต่างๆ โดยใช้เทคนิคการตรวจจับความผิดปกติ การประมาณค่าอย่างต่อเนื่อง และการสร้างข้อความหรือโค้ดโดยอัตโนมัติ (Generative AI) เพื่อช่วยให้นักวิเคราะห์สามารถระบุและตอบสนองได้อย่างรวดเร็วและแม่นยำ

สำนักงานโครงการวิจัยขั้นสูงด้านกลาโหม (DARPA) ของสหรัฐอเมริกา

ได้จัดการแข่งขัน AI Cyber Challenge (AIxCC) เพื่อส่งเสริมการพัฒนาเครื่องมือ AI สำหรับการตรวจจับและแก้ไขช่องโหว่ด้านความปลอดภัยทางไซเบอร์โดยอัตโนมัติ ในการแข่งขันนี้ ทีมต่างๆ ได้พัฒนาโมเดล AI ที่สามารถค้นพบช่องโหว่ใหม่ 22 รายการและแก้ไขได้ 15 รายการโดยอัตโนมัติ

หน่วยงานความมั่นคงของอังกฤษ (NCSC)

ใช้กรอบการประเมินความมั่นคงปลอดภัยทางไซเบอร์ (Cyber Assessment Framework) ซึ่งรวมถึงการประเมินการใช้ AI ในการตรวจจับและตอบสนองต่อภัยคุกคาม การประเมินนี้ช่วยให้หน่วยงานสามารถระบุช่องโหว่และปรับปรุงมาตรการความปลอดภัยได้อย่างมีประสิทธิภาพ

AI กับความมั่นคงของประเทศไทย

สถานะปัจจุบันของกลยุทธ์ความมั่นคงของประเทศไทยกับ AI

ประเทศไทยกำลังให้ความสำคัญกับการใช้ AI ในการเสริมสร้างความมั่นคงของประเทศไทย มีการจัดตั้งหน่วยงานและกลยุทธ์ต่างๆ ดังนี้

• หน่วยงาน:

- ศูนย์ปฏิบัติการ AI ด้านความมั่นคง (AI Security Operation Center - AISO)
- คณะกรรมการขับเคลื่อนการพัฒนาและใช้ปัญญาประดิษฐ์เพื่อความมั่นคง
- สำนักงานคณะกรรมการดิจิทัลเพื่อเศรษฐกิจและสังคม (สดช.)

• กลยุทธ์:

- แผนปฏิบัติการด้านปัญญาประดิษฐ์แห่งชาติ (พ.ศ. 2565 - 2569)
- กลยุทธ์ชาติว่าด้วยปัญญาประดิษฐ์ (พ.ศ. 2564 - 2569)

พื้นที่ที่มีศักยภาพสำหรับการรวมและการปรับปรุง AI

จากกลยุทธ์ดังกล่าว มีพื้นที่ 5 ด้านที่มีศักยภาพสำหรับการรวมและการปรับปรุง AI ดังนี้

1. การวิเคราะห์ข้อมูลข่าวกรอง:

- พัฒนาระบบ AI วิเคราะห์ข้อมูลจากสื่อสังคมออนไลน์ โดรน ดาวเทียม เพื่อระบุภัยคุกคาม
- คาดการณ์เหตุการณ์ล่วงหน้า เตือนภัยคุกคาม ช่วยให้หน่วยงานความมั่นคงตัดสินใจได้อย่างมีประสิทธิภาพ

2. การป้องกันและต่อต้านภัยไซเบอร์:

- พัฒนาระบบ AI ตรวจสอบและป้องกันการโจมตีทางไซเบอร์ ปกป้องโครงสร้างพื้นฐานสำคัญ
- สืบสวนหาต้นตอของการโจมตีทางไซเบอร์ นำผู้กระทำผิดมาลงโทษ

3. การพัฒนาอาวุธยุทโธปกรณ์:

- พัฒนาอาวุธที่มีความแม่นยำสูง ทำงานโดยอัตโนมัติ
- พัฒนาระบบป้องกันภัยทางอากาศ ตรวจสอบและยิงสกัดโดรน ขีปนาวุธ เครื่องบิน

4. การรักษาความปลอดภัยภายในประเทศ:

- พัฒนาระบบ AI วิเคราะห์ภาพจากกล้องวงจรปิด ระบุบุคคลต้องสงสัย ป้องกันอาชญากรรม
- ติดตามบุคคล ยานพาหนะ ควบคุมฝูงชน รักษาความสงบเรียบร้อย

5. การเสริมสร้างขีดความสามารถทางทหาร:

- พัฒนาระบบจำลองการฝึกซ้อม ฝึกทหารให้พร้อมรบ
- วิเคราะห์กลยุทธ์ วางแผนปฏิบัติการทางทหาร เพิ่มประสิทธิภาพในการรบ

ตัวอย่างการใช้งาน AI ด้านความมั่นคงในไทย

- สำนักงานตำรวจแห่งชาติ ใช้ AI วิเคราะห์ภาพจากกล้องวงจรปิดเพื่อติดตามผู้ต้องหา
- กระทรวงดิจิทัลเพื่อเศรษฐกิจและสังคม ใช้ AI ป้องกันการโจมตีทางไซเบอร์ต่อเว็บไซต์ภาครัฐ
- กองทัพอากาศไทย พัฒนาโดรนตรวจการณ์ติดตั้ง AI



ผลกระทบด้านลบของ AI

1. การใช้ AI ในทางที่ผิด

เทคโนโลยี AI มีศักยภาพสูง แต่ก็แฝงไว้ด้วยอันตรายหากถูกนำไปใช้ในทางที่ผิด ตัวอย่างเช่น การโจมตีทางไซเบอร์ การสร้าง deepfakes หรือการพัฒนาอาวุธอัตโนมัติ สิ่งเหล่านี้ล้วนเป็นภัยคุกคามต่อความมั่นคงและความปลอดภัยของทั้งบุคคลและสังคม

2. ภัยคุกคามต่อความเป็นส่วนตัวและข้อมูลส่วนบุคคล

ระบบ AI จำเป็นต้องใช้ข้อมูลจำนวนมากในการเรียนรู้และทำงาน ซึ่งอาจนำไปสู่ความเสี่ยงในการรั่วไหลหรือถูกนำไปใช้ในทางมิชอบได้ ตัวอย่างเช่น การติดตามพฤติกรรม การโฆษณาสินค้าที่ตรงใจ หรือการเลือกปฏิบัติต่อบุคคล

3. ความเหลื่อมล้ำจากการแทนที่แรงงานมนุษย์

AI มีประสิทธิภาพสูงและทำงานได้รวดเร็วกว่ามนุษย์มาก สิ่งนี้อาจนำไปสู่การถูกแทนที่ในบางอาชีพ ส่งผลต่อการว่างงานและความเหลื่อมล้ำทางสังคม

4. ความเสี่ยงหากระบบ AI ล้มเหลว

ระบบ AI มีความซับซ้อนและทำงานโดยอัตโนมัติ การล้มเหลวของระบบ AI อาจสร้างความเสียหายร้ายแรง ตัวอย่างเช่น อุบัติเหตุจากรถยนต์ไร้คนขับ หรือการตัดสินใจที่ผิดพลาดของระบบ AI ในธุรกิจ

แนวทางในการจัดการกับผลกระทบทางลบของ AI

- พัฒนานโยบายและกฎหมายเพื่อควบคุมการใช้ AI
- พัฒนาเทคโนโลยี AI ที่มีความปลอดภัยและโปร่งใส
- ส่งเสริมการศึกษาและความรู้เกี่ยวกับ AI
- เตรียมพร้อมสำหรับการเปลี่ยนแปลงของตลาดแรงงาน
- พัฒนาระบบรองรับผู้ได้รับผลกระทบจาก AI



ผลกระทบด้านลบของ AI

ตัวอย่างผลกระทบด้านลบของ AI:

ด้านความปลอดภัย:

• แอ็กเตอร์ใช้ AI สร้างอีเมลหลอกลวง:

- กรณีศึกษา: การโจมตีเว็บไซต์ของ UNICEF ในปี 2022 ทำให้เว็บไซต์ล่มชั่วคราว

• มัลแวร์ AI:

- กรณีศึกษา: นักวิจัยพบรหัสแม่เรียง AI ที่สร้างขึ้นโดยนาซี ถูกฝังลงในซอฟต์แวร์ AI ปัจจุบันหลายตัว

ด้านจริยธรรม:

• การเลือกปฏิบัติ:

- AI อาจเรียนรู้และขยายอคติที่มีอยู่ในชุดข้อมูลที่ใช้ฝึกอบรม

• การล่วงละเมิดสิทธิมนุษยชน:

- AI อาจถูกใช้เพื่อติดตาม ควบคุม หรือจำกัดเสรีภาพของบุคคล

ด้านเศรษฐกิจ:

• การสูญเสียงาน:

- AI และหุ่นยนต์อัตโนมัติอาจแทนที่แรงงานในบางภาคส่วน

• ความเหลื่อมล้ำ:

- ประโยชน์จาก AI อาจกระจุกตัวอยู่กับกลุ่มคนเพียงบางกลุ่ม

ตัวอย่างเพิ่มเติมของผลกระทบด้านลบของ AI:

• ด้านสิ่งแวดล้อม:

- การใช้พลังงานของระบบ AI เพิ่มขึ้น
- การผลิตและกำจัดฮาร์ดแวร์ AI ส่งผลกระทบต่อสิ่งแวดล้อม

• ด้านสุขภาพจิต:

- การใช้ AI มากเกินไปอาจส่งผลต่อสุขภาพจิต



ใช้ deepfake
ปลอมเป็นหัวหน้า
สูญเงิน 900 ล้านบาท

อีจัน



PLEASE DON'T DO THAT



Discover



Submit



Welcome to the AIID



Discover Incidents



Spatial View



Table View



Entities



Taxonomies



Word Counts



Submit Incident Reports



Submission Leaderboard



Blog



Welcome to the

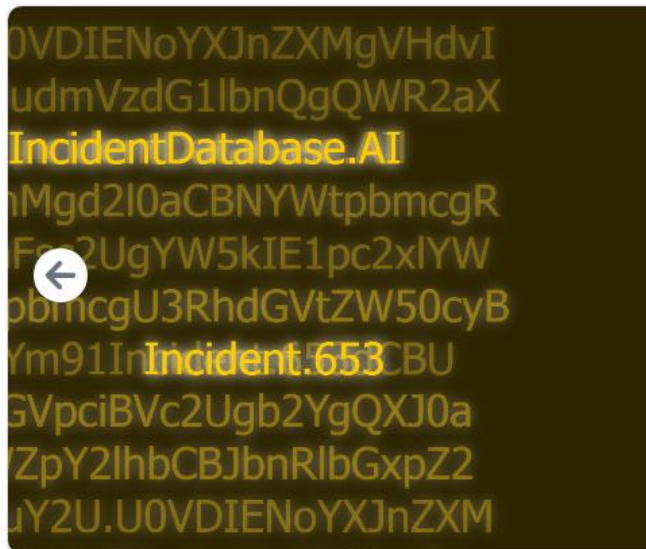
AI Incident Database



Search over 3000 reports of AI harms

Search

Discover



SEC Charges Two Investment Advisers with Making False and Misleading Statements About Their Use of Artificial Intelligence

Latest Incident Report

2024-03-18 [sec.gov](#)

Washington D.C., March 18, 2024 --- The Securities and Exchange Commission today announced settled charges against two investment advisers, Delphia (USA) Inc. and Global Predictions Inc., for making false and misleading statements about the...

Read More →

Canada lawyer under fire for submitting fake cases created by AI chatbot

Chong Ke, from Vancouver, under investigation after allegedly using ChatGPT to cite case law - but those cases did not exist



Samsung Bans ChatGPT Among Employees After Sensitive Code Leak

Siladitya Ray Forbes Staff

Siladitya Ray is a New Delhi-based Forbes news team reporter.

Follow



May 2, 2023, 07:17am EDT

Updated May 2, 2023, 07:31am EDT

TOPLINE Samsung Electronics has banned the use of ChatGPT and other AI-powered chatbots by its employees, Bloomberg reported, becoming the latest company to crack down on the workplace use of AI services amid concerns about sensitive internal information being leaked on such platforms.



การที่ AI ถูกแฮกสามารถก่อให้เกิดผลกระทบที่รุนแรงในหลายๆ ด้าน ทั้งในเรื่องของความปลอดภัยข้อมูล ความมั่นคงของระบบ และผลกระทบต่อสังคมและเศรษฐกิจ

1. ความปลอดภัยข้อมูล

- **การเข้าถึงข้อมูลส่วนบุคคล:** หากระบบ AI ที่จัดการข้อมูลส่วนบุคคลถูกแฮก อาชญากรไซเบอร์สามารถเข้าถึงข้อมูลที่สำคัญและเป็นความลับ เช่น ข้อมูลบัตรเครดิต, ข้อมูลสุขภาพ, หรือข้อมูลการทำธุรกรรมทางการเงิน
- **การเปิดเผยข้อมูลที่เป็นความลับ:** ข้อมูลที่เป็นความลับขององค์กรหรือรัฐบาล เช่น ข้อมูลทางธุรกิจ, แผนการตลาด, หรือข้อมูลที่มีความสำคัญต่อความมั่นคงของชาติ อาจถูกเปิดเผยหรือขายต่อในตลาดมืด

2. ความมั่นคงของระบบ

- **การควบคุมระบบที่สำคัญ:** แฮกเกอร์สามารถควบคุมระบบ AI ที่ถูกใช้งานในโครงสร้างพื้นฐานสำคัญ เช่น ระบบไฟฟ้า, ระบบการขนส่ง, หรือระบบการสื่อสาร ทำให้เกิดการหยุดชะงักและความเสียหายที่รุนแรง
- **การสร้างความสับสน:** การปลอมแปลงข้อมูลหรือการส่งข้อมูลเท็จเข้าสู่ระบบ AI อาจทำให้เกิดความสับสนและการตัดสินใจที่ผิดพลาด ซึ่งสามารถนำไปสู่ความเสียหายที่ยากจะกู้คืน

3. ผลกระทบต่อสังคมและเศรษฐกิจ

- **การสูญเสียความเชื่อมั่น:** การที่ระบบ AI ถูกแฮกสามารถทำให้ประชาชนและองค์กรสูญเสียความเชื่อมั่นในเทคโนโลยี AI และระบบที่ใช้ AI
- **ความเสียหายทางเศรษฐกิจ:** การแฮกระบบ AI อาจทำให้เกิดความเสียหายทางการเงิน เช่น การขโมยทรัพย์สินทางปัญญา, การโจรกรรมทางการเงิน, หรือการทำลายโครงสร้างพื้นฐานทางการค้า

4. ผลกระทบต่อประชาชน

- **ความเป็นส่วนตัวและความปลอดภัย:** การเข้าถึงข้อมูลส่วนบุคคลหรือข้อมูลที่เป็นความลับสามารถนำไปใช้ในการแอบอ้างหรือการโจมตีทางไซเบอร์ เช่น การโจรกรรมตัวตนหรือการข่มขู่
- **การเสื่อมสภาพจิตใจ:** การที่ข้อมูลส่วนตัวหรือข้อมูลที่เป็นความลับถูกเผยแพร่หรือถูกนำไปใช้ในทางที่ไม่เหมาะสมสามารถส่งผลกระทบต่อผู้ที่ได้รับผลกระทบ



AI (AI Incident Database) รวบรวมเหตุการณ์ที่เกิดจากเทคโนโลยี AI เพื่อช่วยให้เราเข้าใจและลดความเสี่ยงที่อาจเกิดขึ้นจาก AI ได้

- 1. อัลกอริทึมของ Department for Work and Pensions (DWP):** อัลกอริทึมที่ใช้โดยรัฐบาลสหราชอาณาจักรในการตรวจสอบการฉ้อโกงสิทธิประโยชน์ด้านที่อยู่อาศัยทำให้ผู้คนกว่า 200,000 รายถูกระบุผิดพลาดและถูกสอบสวนโดยไม่มีเหตุผล ทำให้เกิดความเครียดและความไม่สะดวกแก่ผู้บริสุทธิ์จำนวนมาก
- 2. การใช้ Deepfake ในทางที่ผิด:** บัญชี Deepfake ที่ชื่อ "Amandine Le Pen" ลอกลวงผู้ใช้จำนวนมากด้วยการแพร่กระจายข้อมูลเท็จแสดงให้เห็นถึงอันตรายของเทคโนโลยี Deepfake ในการบิดเบือนความคิดเห็นของสาธารณชน
- 3. ตลาดใต้ดินสำหรับ LLMs:** โมเดลภาษาขนาดใหญ่ (LLMs) เช่นที่มาจาก OpenAI ถูกนำมาใช้ในบริการที่เป็นอันตราย รวมถึงการสร้างมัลแวร์และการโจมตี ฟิชซิง ซึ่งยากต่อการตรวจจับ
- 4. การพัฒนากลยุทธ์ของนักต้มตุ๋น:** AI ถูกใช้เพื่อเพิ่มความซับซ้อนของการหลอกลวง ทำให้การระบุและจัดการกับปัญหาเหล่านี้ยากขึ้นและก่อให้เกิดปัญหาสำคัญต่อบุคคล



AVR 15 2024 Amandine et Chloé Le Pen, Lena Marechal sont des faux profils TikTok créés par IA au service de l'extrême droite.

ข้อดีและข้อเสีย ของปัญญาประดิษฐ์

ข้อดี (Advantages)

- มีประสิทธิภาพและความแม่นยำ
- ขจัดความผิดพลาดของมนุษย์
- สามารถลดต้นทุนการผลิตได้
- การปรับปรุงการตัดสินใจของมนุษย์
- การปรับปรุงเวิร์กโฟลว์ของมนุษย์
- ความได้เปรียบทางการทำงาน
- การรับและวิเคราะห์ข้อมูลที่มีประสิทธิภาพ

ข้อเสีย (Disadvantages)

- ไม่อาจสามารถควบคุมได้
- ปัญญาประดิษฐ์ไม่มีความรู้สึกนึกคิด
- ความเสียหายและสลายตัวของระบบ
- จำนวนงานที่ลดลงของมนุษย์
- มีค่าใช้จ่ายในการดำเนินการสูง
- ขาดความคิดสร้างสรรค์การคิดนอกกรอบ
- การพิจารณาทางด้านจริยธรรม

การจัดการผลกระทบของ AI

1. ออกกฎหมาย/ระเบียบควบคุมการใช้งาน AI

- กำหนดกรอบการทำงานที่ชัดเจนสำหรับการพัฒนาและใช้งาน AI
- กำหนดมาตรฐานด้านความปลอดภัยและจริยธรรม
- กำกับดูแลการใช้งาน AI ในด้านที่มีความเสี่ยงสูง เช่น อาวุธอัตโนมัติ

2. เสริมสร้างความมั่นคงปลอดภัยของโครงสร้างพื้นฐานไซเบอร์

- พัฒนาระบบป้องกันภัยคุกคามทางไซเบอร์
- ยกกระดับมาตรฐานความปลอดภัยของข้อมูล
- เพิ่มการฝึกอบรมบุคลากรด้านความปลอดภัยไซเบอร์

3. พัฒนาบุคลากรที่มีความรู้ด้าน AI

- ส่งเสริมการศึกษา STEM
- พัฒนาทักษะการคิดวิเคราะห์และแก้ปัญหา
- สร้างหลักสูตรการฝึกอบรมเฉพาะทางด้าน AI

4. ส่งเสริมการวิจัยและพัฒนาด้าน AI แบบมีจริยธรรม

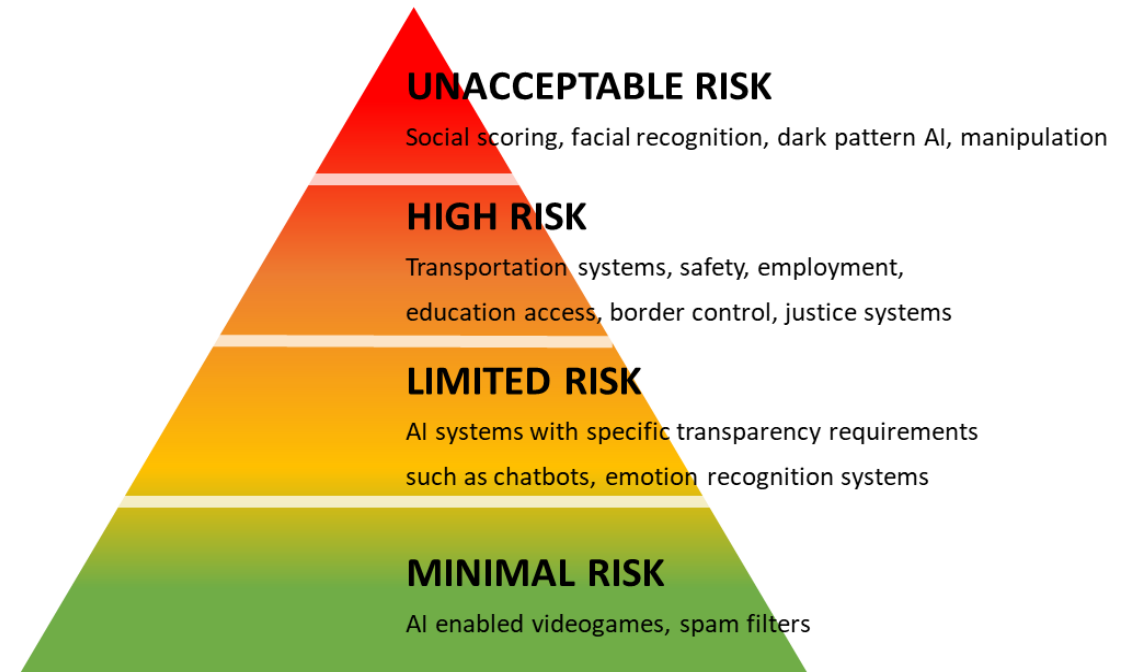
- กำหนดแนวทางการวิจัย AI ที่มีความรับผิดชอบ
- พัฒนาระบบ AI ที่โปร่งใสและตรวจสอบได้
- เน้นการวิจัย AI ที่ส่งเสริมประโยชน์ต่อสังคม

แนวทางเพิ่มเติม

- สร้างความตระหนักรู้และความเข้าใจเกี่ยวกับ AI ให้แก่ประชาชน
- ส่งเสริมการมีส่วนร่วมของภาคประชาสังคมในการกำหนดนโยบาย AI
- พัฒนาระบบรองรับผู้ได้รับผลกระทบจาก AI

สรุป

การจัดการผลกระทบของ AI จำเป็นต้องอาศัยความร่วมมือจากทุกภาคส่วน ทั้งภาครัฐ ภาคเอกชน ภาคประชาสังคม และประชาชนทั่วไป โดยต้องดำเนินการควบคู่ไปกับการพัฒนาเทคโนโลยี AI อย่างยั่งยืนและมีจริยธรรม



ตัวอย่าง**การประยุกต์ใช้ AI** เพื่อให้บริการประชาชนในด้านต่าง ๆ จากทั่วโลก



สวัสดิการสังคม

ตรวจจับรูปแบบการฉ้อโกง
เพื่อเรียกคืนสิทธิสวัสดิการ

ความมั่นคง ภายในประเทศ

เฝ้าระวังด้วยระบบจดจำใบหน้า
จากภาพ วิดีโอและข้อมูลที่บันทึก
ไว้ในกล้อง CCTV

การคมนาคมขนส่ง

รถรับส่งขับเคลื่อนตัวเอง
ด้วยความเร็วต่ำกว่า 50 กม./ชม.
ไปตามเส้นทางที่กำหนด

เหตุฉุกเฉิน

จัดประเภทการโทร.ตามความเร่งด่วน
โดยใช้เทคโนโลยีการจดจำเสียง
และ Machine Learning

สาธารณสุข

ติดตามการแพร่กระจายของเชื้อโรค
โดยใช้ Machine Learning

การศึกษา

วิเคราะห์ความก้าวหน้าของผู้เรียน
ค้นหาความไม่สอดคล้องกันระหว่าง
สิ่งที่สอนกับสิ่งที่ยังไม่เข้าใจ

ประชาสัมพันธ์

ใช้ Chatbot (ผู้ช่วยเสมือน)
ในฝ่ายบริการประชาชน

ความท้าทาย และทิศทางในอนาคต ของ AI ในภาครัฐ



Q&A

Contact us

National Cyber Security Agency - NCSA



NCSA Thailand



Line ID : NCSA Thailand



E-Mail:



saraban@ncsa.or.th



02-142-6885

shorturl.at/hkAC2

