

กรอบการกำกับดูแล
การใช้งานปัญญาประดิษฐ์
ในภาคตลาดทุน

Artificial Intelligence Governance

Framework and Handbook

ฉบับปรับปรุงครั้งที่ 1 | 1st Revised Edition
สิงหาคม 2569

สำนักงานคณะกรรมการกำกับหลักทรัพย์และตลาดหลักทรัพย์

กรอบการกำกับดูแล
การใช้งานปัญญาประดิษฐ์
ในภาคตลาดทุน

Artificial Intelligence
Governance
Framework and Handbook

ฉบับปรับปรุงครั้งที่ 1 | 1st Revised Edition
สิงหาคม 2569

สำนักงานคณะกรรมการกำกับหลักทรัพย์และตลาดหลักทรัพย์

สารบัญ

1. บทนำ.....	1
1.1 ความสำคัญและความเป็นมา	1
1.2 บทนิยาม	3
1.3 การใช้งานปัญญาประดิษฐ์ในภาคตลาดทุน	4
2. กรอบการกำกับดูแลการใช้งาน AI (AI Governance Framework)	8
2.1 การกำกับดูแล (Governance)	9
2.2 หลักการที่ควรคำนึงถึง (Principles).....	11
2.3 การบริหารจัดการความเสี่ยง (Risk Management)	15
2.4 แนวทางการบริหารจัดการระบบ AI (AIMS Controls)	20
2.5 ปัจจัยส่งเสริมให้บรรลุวัตถุประสงค์ (Enablers).....	25
ภาคผนวก.....	29
ภาคผนวก 1 จุดมุ่งหมายของการนำกรอบการกำกับดูแลไปใช้และการใช้งานภาคผนวก.....	30
ภาคผนวก 2 กระบวนการ 12 ขั้นตอน สำหรับการประเมินผลกระทบ AI (AI Impact Assessment Flow).....	31
ภาคผนวก 3 ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยง (Risk Rubrics and Scoring for AI Use Cases).....	34
ภาคผนวก 4 AI Inventory.....	36
ภาคผนวก 5 AIMS Controls	43
ภาคผนวก 6 ตัวอย่างการประยุกต์ใช้เครื่องมืออย่างง่าย Lite Example - GenAI Summary Tool.....	50
ภาคผนวก 7 ตัวอย่างการประยุกต์ใช้เครื่องมือกรณี Use Case เสี่ยงสูง - Automated Trading Support.....	58
ภาคผนวก 8 ความเสี่ยงสำคัญจากการใช้งาน AI ในภาคตลาดทุน	74
ภาคผนวก 9 คำถามที่พบบ่อย (FAQs).....	79
ภาคผนวก 10 ตัวอย่างตัวชี้วัดและข้อมูลสำหรับการกำกับติดตามการใช้งาน AI	86
ภาคผนวก 11 ตัวอย่างแบบรายงานเหตุการณ์ที่เกี่ยวข้องกับ AI	87
ภาคผนวก 12 ตัวอย่างแบบประเมินผู้ให้บริการ AI และความเสี่ยงจากการกระจุกตัว.....	88
ภาคผนวก 13 ตัวอย่างแนวทางการเปิดเผยการใช้ AI และการป้องกัน AI-washing	89
ภาคผนวก 14 ตัวอย่างคำถามสำหรับการกำกับติดตาม การทบทวนภายใน และการเตรียมความพร้อมสำหรับ supervisory dialogue ...	90
บรรณานุกรม	91

สารบัญตาราง

ตารางที่ 1	กลไกการนำกรอบการกำกับดูแลไปใช้งานจริง (Operating Mechanism).....	19
ตารางที่ 2	ตัวอย่างรายละเอียดกระบวนการ 12 ขั้นตอนสำหรับการบริหารจัดการและประเมินผลกระทบ AI ..	32
ตารางที่ 3	ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยงในการประเมินความเสี่ยงของ Use Case	34
ตารางที่ 4	ตัวอย่าง AI Use Case Summary (Register)	36
ตารางที่ 5	ตัวอย่าง Detailed AI Inventory (Use Case Detail)	38
ตารางที่ 6	ตัวอย่าง AIMS Controls: Mutual and Role-specific Controls	45
ตารางที่ 7	ตัวอย่างการประเมินกระบวนการ 12 ขั้นตอนของ GenAI Summary Tool.....	50
ตารางที่ 8	ตัวอย่างการจัดทำ AI Inventory ของ GenAI Summary Tool	51
ตารางที่ 9	ตัวอย่างการเลือกมาตรการควบคุม AIMS Controls ของ GenAI Summary Tool.....	55
ตารางที่ 10	ตัวอย่าง trigger หรือปัจจัยสำคัญที่ต้องพิจารณากระดับไปประเมินเชิงลึก.....	57
ตารางที่ 11	ตัวอย่างมิติคุณลักษณะของระบบ Agentic AI ที่ต้องพิจารณาเพิ่มเติม	58
ตารางที่ 12	ตัวอย่างการประเมินกระบวนการ 12 ขั้นตอนของ Automated Trading Support	61
ตารางที่ 13	ตัวอย่างการประเมินความเสี่ยงตาม Risk Rubrics ของ Automated Trading Support.....	63
ตารางที่ 14	ตัวอย่างการจัดทำ AI Inventory ของ Automated Trading Support.....	64
ตารางที่ 15	ตัวอย่างการเลือกมาตรการควบคุม AIMS Controls ของ Automated Trading Support	71
ตารางที่ 16	ตัวอย่าง trigger หรือปัจจัยที่ต้องระงับการใช้งานหรือทบทวนระบบ Agentic AI.....	73
ตารางที่ 17	ประเภทความเสี่ยงจากการใช้งาน AI ในภาคตลาดทุนและตัวอย่างมาตรการควบคุมความเสี่ยง	74
ตารางที่ 18	ตัวอย่างรายละเอียดการพิจารณาปรับปรุงการควบคุมตามแนวทาง Reuse-Extend-Align	82
ตารางที่ 19	ตัวอย่างตัวชี้วัดและข้อมูลสำหรับการกำกับติดตามการใช้งาน AI.....	86
ตารางที่ 20	ตัวอย่างรายละเอียดรายงานเหตุการณ์ที่เกี่ยวข้องกับ AI และข้อมูลที่ควรจัดเก็บ	87
ตารางที่ 21	ตัวอย่างรายละเอียดการประเมินผู้ให้บริการ AI และความเสี่ยงจากการกระจุกตัว.....	88
ตารางที่ 22	ตัวอย่างรายละเอียดคำถามสำหรับทบทวนความพร้อมสำหรับ supervisory dialogue	90

สารบัญภาพ

ภาพที่ 1 บทนำของการกำกับดูแลการใช้งาน AI	2
ภาพที่ 2 การใช้งานปัญญาประดิษฐ์ในภาคตลาดทุน	4
ภาพที่ 3 กรอบการกำกับดูแลการใช้งาน AI (AI Governance Framework)	8
ภาพที่ 4 การกำกับดูแล (Governance)	10
ภาพที่ 5 หลักการที่ควรคำนึงถึง (Principles)	14
ภาพที่ 6 มุมมองการบริหารจัดการความเสี่ยง AI (Risk Lens)	15
ภาพที่ 7 ความเสี่ยงจากการประยุกต์ใช้ AI ในบริบทของภาคตลาดทุน	18
ภาพที่ 8 กลไกการนำกรอบไปใช้ในทางปฏิบัติ (Operating Mechanism)	19
ภาพที่ 9 AIMS Controls	20
ภาพที่ 10 ปัจจัยส่งเสริมให้บรรลุวัตถุประสงค์ (Enablers)	25
ภาพที่ 11 กระบวนการประเมิน 12 ขั้นตอน	31
ภาพที่ 12 ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยง	34
ภาพที่ 13 การแบ่งบทบาทเพื่อเลือกใช้มาตรการควบคุม	44
ภาพที่ 14 Mutual and Role-specific Controls	45



AI GOVERNANCE FRAMEWORK NAVIGATOR

แนวทางการใช้งาน กรอบการกำกับดูแลการใช้งาน AI ในภาคตลาดทุน



กรอบ AI ฉบับนี้ ออกแบบให้ผู้ใช้มีส่วนเกี่ยวข้องทุกบทบาทสามารถศึกษาและทำความเข้าใจ “ทั้งฉบับ” Navigator นี้ช่วยจัดลำดับการอ่านและการใช้งาน และเลือกใช้งานในส่วนที่เกี่ยวข้องได้อย่างมีประสิทธิภาพ



5 GUIDES



เลือกใช้งานที่ตรงกับความต้องการ เพื่อเริ่มต้นและใช้งานกรอบได้อย่างมีประสิทธิภาพ

1

เข้าใจโครงสร้างกรอบ (Section Map)

แต่ละส่วนครอบคลุมเรื่องใด

1. Governance	การกำกับดูแลควมมีองค์ประกอบได้บ้าง ทิศทาง • รับผิดชอบ • ตรวจสอบและรายงาน เราควรยึดหลักการใดในการใช้งาน AI
2. Principles	F • E • A • T • H
3. Risk Management	บริหารจัดการความเสี่ยงจาก AI อย่างไร ประเมิน • จัดลำดับ • จัดการ
4. AIMS Controls	มาตรการควบคุม 8 หมวด ตามบทบาท Developer/Customizer/User
5. Enablers	ปัจจัยสนับสนุนความสำเร็จ กระบวนการ • เทคโนโลยี • ข้อมูล • คน
ภาคผนวก 1-14 (พ.น.)	เครื่องมือและตัวอย่างการใช้งานจริง Tools • Templates • Examples • FAQs

ช่วยให้คุณเข้าใจภาพรวมก่อน แล้วค่อยเจาะลึกตามความจำเป็น

2

คุณกำลังทำอะไร?

(Use Case Entry Point)

เลือกสถานการณ์ที่ตรงกับคุณ

▶ เริ่ม Use Case ใหม่	พ.น. 1-2-3-4
🚀 Deploy หลังประเมินเสร็จ	AIMS Controls (2.4) + พ.น. 5,10
📊 เตรียมรายงานต่อบอร์ด	Governance (2.1) + พ.น. 10,14
⚠️ เกิด AI Incident	AIMS (2.4.10) + พ.น. 11
💡 คัดเลือก/ทบทวน Vendor	AIMS (2.4.8) + พ.น. 12
📢 สื่อสารกับลูกค้า/เปิดเผย	Transparency (2.2.4) + พ.น.13
📅 ทบทวนการใช้งานประจำปี	พ.น. 2 + FAQ Q13

ช่วยให้คุณเริ่มต้นได้ถูกต้อง ตามสถานการณ์จริง

3

Use Case ของคุณเสี่ยงระดับใด? (Risk-based Path)

ประเมินระดับความเสี่ยงเพื่อเลือกเส้นทางที่เหมาะสม

ความเสี่ยงต่ำ (Low Risk)	• ใช้งานภายในองค์กร • ไม่กระทบลูกค้า/ผู้ลงทุน • ข้อมูลไม่อ่อนไหว	Lite Example (ภาคผนวก 6)
ความเสี่ยงปานกลาง (Medium Risk)	• มีผลกระทบต่อระบบงาน/ธุรกิจ/นักลงทุน/ลูกค้า • อาจมีข้อมูลอ่อนไหว	ภาคผนวก 2-5 (Think > Record > Act)
ความเสี่ยงสูง (High Risk)	• กระบวนการตัดสินใจสำคัญ • กระบตลาด/ระบบ/ผู้ลงทุน และใช้ Agentic AI ร่วมด้วย	Full Example (ภาคผนวก 7)

หมายเหตุ: เกณฑ์การแบ่งระดับความเสี่ยงและรายละเอียดการควบคุม ปรับได้ตามบริบทองค์กร

4

มีคำถาม? เริ่มต้นที่นี่ (Question Navigator)

คำถามที่พบบ่อย

• ยังไม่รู้จักเริ่มจากไหน	FAQ Q6	• ต้องเก็บเอกสารขึ้นทำอะไรบ้าง?	FAQ Q20
• ทรัพยากรจำกัด ต้องทำครบทุกข้อไหม?	FAQ Q1	• Incident ต้องรายงานเมื่อไร?	FAQ Q11
• ไม่มีข้อมูลส่วนบุคคล ต้องประเมินไหม?	FAQ Q3	• ใช้ Vendor รายเดียวกันหลายระบบ?	FAQ Q12
• ใช้ SaaS / GenAI ต้องทำอะไร?	FAQ Q10	• เชื่อมกับระบบภายนอกข้อมูล อย่างไร?	FAQ Q22
• Agentic AI ต่างจาก GenAI อย่างไร?	FAQ Q9	• เตรียมพร้อมรับการกำกับดูแลอย่างไร?	FAQ Q21

ช่วยให้คุณค้นหาคำตอบเร็วขึ้น และแก้ไขปัญหาได้ทันที

5

เครื่องมือที่ใช้บ่อย (Tools & Templates Index)

 กระบวนการ 12 ขั้นตอน ภาคผนวก 2	 Risk Rubrics ภาคผนวก 3	 AI Inventory ภาคผนวก 4	 AIMS Controls ภาคผนวก 5	 ตัวอย่าง Lite Example ภาคผนวก 6
 ตัวอย่าง Full Assessment ภาคผนวก 7	 Key AI risks in capital market ภาคผนวก 8	 FAQs คำถามที่พบบ่อย ภาคผนวก 9	 Vendor Assessment ภาคผนวก 12	 ตัวอย่าง Disclosure / การป้องกัน AI washing ภาคผนวก 13

TIER Strategic and Value Outcomes

T

Trust
สร้างความเชื่อมั่น

I

Innovation
ส่งเสริมนวัตกรรม

E

Efficiency
เพิ่มประสิทธิภาพ

R

Resilience
พร้อมรับมือและฟื้นตัว

Operating Mechanism : กลไกการนำกรอบไปใช้งานจริง



Think

พิจารณา ประเมิน
ความเสี่ยง



Record

บันทึก
การตัดสินใจ



Act

ควบคุมและติดตาม
ต่อเนื่อง

1. บทนำ

1.1 ความสำคัญและความเป็นมา

ปัจจุบันปัญญาประดิษฐ์ (Artificial Intelligence: “AI”) มีการพัฒนาขึ้นอย่างรวดเร็ว และมีการใช้งานอย่างแพร่หลาย เนื่องจากการเพิ่มขึ้นของข้อมูลอิเล็กทรอนิกส์ เครื่องคอมพิวเตอร์ที่มีประสิทธิภาพในการประมวลผลสูง การเข้าถึงที่ง่ายขึ้น สภาพแวดล้อมดังกล่าวได้สร้างโอกาสใหม่ ให้แก่ ผู้ประกอบธุรกิจในตลาดทุน (“ผู้ประกอบธุรกิจ”) ซึ่งสามารถแบ่งออกเป็น 2 ด้าน ได้แก่

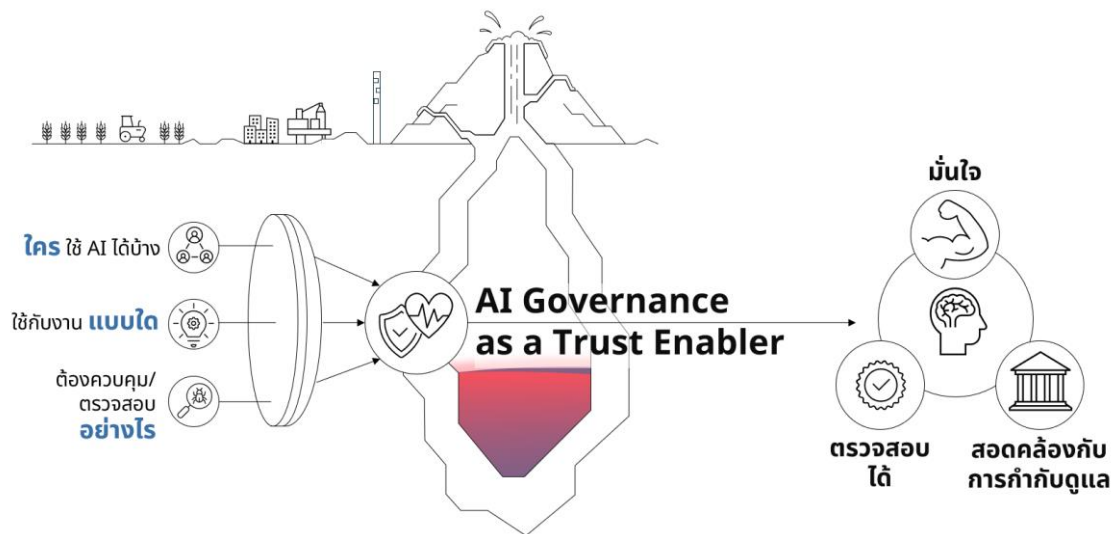
- (1) ด้านประสิทธิภาพและประสิทธิผลภายในองค์กร: AI ช่วยติดตามและบริหารจัดการความเสี่ยงในองค์กร สนับสนุนกระบวนการตัดสินใจ ตรวจสอบความผิดปกติ หรือเหตุผิดปกติอันควรสงสัย ซึ่งสามารถนำไปสู่การแก้ไขปัญหาได้อย่างทันท่วงที และยังช่วยเพิ่มประสิทธิภาพในการทำงานลดระยะเวลาการทำงานได้มากขึ้น เช่น การสรุปเนื้อหาเอกสาร การสร้าง code เป็นต้น
- (2) ด้านประสิทธิภาพและประสิทธิผลของการให้บริการลูกค้า: AI ช่วยสนับสนุนผู้ประกอบธุรกิจในการคิดค้น วางแผน และนำเสนอผลิตภัณฑ์หรือบริการให้แก่ลูกค้าของผู้ประกอบธุรกิจ และช่วยอำนวยความสะดวกให้แก่ลูกค้าในการใช้บริการของผู้ประกอบธุรกิจได้เช่นกัน

สำนักงานคณะกรรมการกำกับหลักทรัพย์และตลาดหลักทรัพย์ (“สำนักงาน ก.ล.ต.”) สนับสนุนการนำเทคโนโลยีใหม่มาใช้ในภาคตลาดทุน และเล็งเห็นถึงความเสี่ยงที่เกี่ยวข้องกับการใช้งาน AI ในปัจจุบัน จึงได้นำกรอบการกำกับดูแลการใช้งาน AI/ML ที่เผยแพร่เมื่อเดือนธันวาคม 2566 มาปรับปรุงเพิ่มเติมเป็นคู่มือฉบับนี้ เพื่อช่วยให้ผู้ประกอบธุรกิจตระหนักถึงความเสี่ยงและแนวทางการนำ AI มาใช้งานอย่างเหมาะสม และเพื่อสร้างความเชื่อมั่นให้กับผู้ใช้บริการในตลาดทุนไทย ควบคู่การส่งเสริมนวัตกรรมในภาคตลาดทุน เพิ่มประสิทธิภาพในการดำเนินงาน และพร้อมรับการเปลี่ยนแปลง ทั้งนี้ ความเสี่ยงจากการใช้งาน AI ตลอดจนหลักการพื้นฐานในคู่มือฉบับนี้ เป็นการนำเสนอผลการศึกษา การรวบรวมและวิเคราะห์ข้อมูลจากทั้งภายในประเทศและต่างประเทศ อย่างไรก็ตาม ผู้ประกอบธุรกิจยังคงต้องศึกษา ติดตาม และกำกับดูแลบริบทของการใช้งาน AI และความเสี่ยงอื่น ๆ ที่อาจจะเกิดขึ้นได้ในอนาคตจากความก้าวหน้าทางเทคโนโลยี ทั้งฮาร์ดแวร์ และซอฟต์แวร์ ที่เปลี่ยนแปลงไปอย่างรวดเร็ว การพัฒนาและประยุกต์ใช้ AI โดยไม่คำนึงถึงความเสี่ยงและผลกระทบอย่างเพียงพอ อาจก่อให้เกิดความรับผิดตามกฎหมายที่เกี่ยวข้อง ทั้งที่ใช้บังคับอยู่ในปัจจุบันและที่อาจมีการปรับปรุงในอนาคต

นอกจากนี้ การพัฒนาการใช้งาน AI ในภาคตลาดทุน มีพัฒนาการอย่างต่อเนื่องในระดับสากล ทั้งในมิติของการกำกับดูแลภายในองค์กร การกำกับดูแลบุคคลที่ 3 การเปิดเผยข้อมูลต่อผู้ลงทุน การเก็บหลักฐาน และการรายงานเหตุการณ์ ตลอดจนการติดตามผลกระทบเชิงระบบจากการพึ่งพาโครงสร้างพื้นฐาน หรือผู้ให้บริการรายเดียวกันเป็นจำนวนมาก ดังนั้น การจัดทำกรอบการกำกับดูแลการใช้งาน AI ของผู้ประกอบธุรกิจจึงควรคำนึงถึงทั้งบริบทของกฎหมายไทยและแนวปฏิบัติสากล เพื่อให้การดำเนินธุรกิจสามารถเติบโตอย่างมั่นคง โปร่งใส และสอดคล้องกับความคาดหวังของผู้ลงทุนและหน่วยงานกำกับดูแลทั้งภายในประเทศและต่างประเทศ

สำนักงาน ก.ล.ต. เห็นว่า การนำแนวทางสากลมาประยุกต์ใช้ไม่ใช่เป็นการเพิ่มภาระโดยไม่จำเป็น แต่เป็นการสนับสนุนให้ผู้ประกอบธุรกิจสามารถรับมือกับผลกระทบที่อาจเกิดขึ้นจากการใช้งาน AI ทั้งต่อการดำเนินธุรกิจและผู้ลงทุน โดยสามารถออกแบบโครงสร้างการกำกับดูแลและมาตรการควบคุมที่เหมาะสม สอดคล้องกับลักษณะและระดับความเสี่ยงของกรณีการใช้งาน (use case) แต่แต่ละประเภทได้ชัดเจนยิ่งขึ้น โดยเฉพาะการใช้งาน AI ที่มีความซับซ้อนสูงขึ้น มีการพึ่งพาบุคคลที่ 3 หรือผู้ให้บริการภายนอกมากขึ้น หรือผลลัพธ์ที่ได้จากการใช้งาน AI ส่งผลกระทบโดยตรงต่อผู้ลงทุน การตัดสินใจลงทุน สภาพคล่องของตลาด และเสถียรภาพของระบบโดยรวม อย่างไรก็ตาม ผู้ประกอบธุรกิจยังคงมีหน้าที่ในการติดตาม ทบทวน และปรับปรุงกระบวนการทำงานที่เกี่ยวข้องกับ AI ให้เป็นไปตามพัฒนาการทางด้านกฎหมายหรือหลักเกณฑ์ที่เกี่ยวข้องกับ AI ของไทย ทั้งที่มีอยู่ในปัจจุบันและในอนาคตต่อไป

ด้วยเหตุนี้ คู่มือฉบับนี้จึงมุ่งส่งเสริมให้ผู้ประกอบธุรกิจนำหลักการกำกับดูแลที่อิงความเสี่ยง (risk-based approach) หลักความสมเหตุสมผล (“Proportionality”) และการติดตามอย่างต่อเนื่องตลอดวงจรชีวิตของระบบ AI มาปรับใช้ควบคู่กัน โดยให้ความสำคัญ ทั้งการสร้างนวัตกรรมอย่างรับผิดชอบ และการลดความเสี่ยงต่อผู้ใช้บริการ ผู้ลงทุน และตลาดทุนโดยรวมในระยะยาว



ภาพที่ 1 บทนำของการกำกับดูแลการใช้งาน AI

1.2 บทนิยาม

คำศัพท์ที่ใช้ในคู่มือฉบับนี้ มีความหมายดังต่อไปนี้

อัลกอริทึม (Algorithm) หมายถึง ขั้นตอนและวิธีการในการประมวลผลเพื่อหาผลลัพธ์ที่ชัดเจน ด้วยระบบคอมพิวเตอร์ หรือวิธีการอื่นใดในทำนองเดียวกัน

โมเดล (Model) หมายถึง แบบจำลองข้อมูลทางคณิตศาสตร์ ตัวแบบหรือโปรแกรมที่ถูกสร้างขึ้น จากอัลกอริทึมและถูกสอนโดยอาศัยชุดข้อมูลเพื่อให้มีความสามารถในการทำงาน ตามวัตถุประสงค์ที่กำหนด เช่น การพยากรณ์ ผลการดำเนินธุรกิจ การแยกแยะชนิดของวัตถุ และการสร้างรูปภาพ เป็นต้น

ปัญญาประดิษฐ์ (Artificial Intelligence: AI) หมายถึง เทคโนโลยีที่ถูกพัฒนาขึ้นเพื่อให้ คอมพิวเตอร์ มีคุณสมบัติหรือพฤติกรรมใกล้เคียงมนุษย์ เช่น การเรียนรู้ การรับรู้และตอบสนองต่อสภาพแวดล้อม การให้เหตุผล และการแก้ไขปัญหา เป็นต้น ตามวัตถุประสงค์ที่มนุษย์กำหนด

การเรียนรู้ของเครื่อง (Machine Learning: “ML”) หมายถึง AI ประเภทหนึ่งที่มีความสามารถในการเรียนรู้และปรับปรุงประสิทธิภาพการทำงานของตน โดยใช้อัลกอริทึมในการวิเคราะห์และ เรียนรู้จากข้อมูล ที่ได้รับการสอน หรือสภาพแวดล้อม

ปัญญาประดิษฐ์แบบรู้สร้าง (Generative AI) หมายถึง AI ที่มีความสามารถในการสร้าง (Generate) ชุดข้อมูลใหม่ขึ้นมาในหลากหลายรูปแบบตามข้อความหรือคำสั่ง (Prompt) ที่กำหนด ไม่ว่าจะเป็นข้อความ ภาพ เสียง Code และสารสนเทศอื่น ๆ จากการเรียนรู้ ของโมเดลโดยมีตัวอย่างของ Generative AI ซึ่งเป็นที่รู้จัก ในวงกว้างคือ ChatGPT, Claude หรือ Gemini เป็นต้น

ปัญญาประดิษฐ์แบบตัวแทน (Agentic AI) หมายถึง AI ขั้นสูงประเภทหนึ่งที่มีความสามารถในการรับรู้ การใช้เหตุผล การวางแผนและการดำเนินการอย่างอิสระ สามารถเข้าใจเป้าหมายและริเริ่มดำเนินการได้เอง โดยใช้เครื่องมือและความรู้ตามบริบทต่าง ๆ

ความเป็นธรรม (Fairness) หมายถึง การตัดสินใจที่อยู่บนพื้นฐานของความเท่าเทียม เป็นธรรม โดยไม่เลือกปฏิบัติกับบุคคล กลุ่มบุคคลหรือสิ่งใด ๆ และหลีกเลี่ยงการใช้ข้อมูลส่วนบุคคลที่มีความอ่อนไหว (Sensitive Personal Data) เช่น เพศ เชื้อชาติ ศาสนา และความเห็นทางการเมือง เป็นต้น ในกระบวนการตัดสินใจ เว้นแต่ การกระทำดังกล่าวจะเป็นไปอย่างสมเหตุสมผล

จริยธรรม (Ethics) หมายถึง การออกแบบ พัฒนาและใช้งาน AI ที่พฤติกรรมหรือผลจากการทำงาน ของ AI สอดคล้องกับกฎระเบียบ แนวปฏิบัติ หลักจริยธรรม หรือหลักการอันดีที่พึงปฏิบัติตามบริบทของการนำ AI ไปประยุกต์ใช้

ความรับผิดชอบ (Accountability) หมายถึง ต้องมีผู้รับผิดชอบต่อผลลัพธ์ของ AI มีการกำกับดูแล การทดสอบ และการตรวจสอบ (auditability) กำหนดโครงสร้างความรับผิดชอบในองค์กรอย่างชัดเจน

ความโปร่งใส (Transparency) หมายถึง ความสามารถในการอธิบายเหตุการณ์ การกระทำ กระบวนการทำงาน และกิจกรรมต่าง ๆ เกี่ยวกับ AI รวมถึงตรวจสอบย้อนหลังเหตุการณ์ที่เกิดขึ้นได้

การกำกับดูแลโดยมนุษย์ (Human Oversight) หมายถึง กลไกที่ทำให้บุคคลที่ได้รับมอบหมาย สามารถกำกับดูแลการทำงานของระบบ AI ได้อย่างมีประสิทธิภาพ โดยมีวัตถุประสงค์เพื่อป้องกันหรือลดความเสี่ยง ที่อาจเกิดขึ้นจากการใช้งาน AI

แนวทางการบริหารจัดการระบบ AI (AI Management System: “AIMS”) หมายถึง ขั้นตอนและวิธีการในการบริหารจัดการระบบ AI ที่ครอบคลุมตลอดวงจรชีวิต AI (AI Lifecycle) เริ่มตั้งแต่การกำหนดนโยบาย AI ไปจนถึงการยุบหรือยกเลิก AI ที่คำนึงถึงความปลอดภัยและผลกระทบต่อผู้เกี่ยวข้องตลอดวงจรชีวิต

วงจรชีวิต AI (AI Lifecycle) หมายถึง ชุดกระบวนการที่ครอบคลุมตั้งแต่การกำหนดแนวคิด ออกแบบ พัฒนา ทดสอบ นำไปใช้งาน ฝ้าติดตาม ไปจนถึงการเลิกใช้งานระบบ AI โดยต้องมีการบริหารความเสี่ยงและควบคุมคุณภาพ และฝังความปลอดภัยตลอดทุกขั้นตอน เพื่อให้มั่นใจในความมั่นคงปลอดภัย ความโปร่งใส และความรับผิดชอบต่อผลกระทบที่เกิดขึ้นกับธุรกิจและผู้ใช้

Developer (ผู้พัฒนา) หมายถึง บุคคลหรือทีมที่รับผิดชอบการออกแบบ พัฒนา ทดสอบ ตรวจสอบ ปรับปรุง หรือควบคุมองค์ประกอบหลักของโมเดลหรือระบบ AI เพื่อวัตถุประสงค์เฉพาะขององค์กร รวมถึงการกำหนดสถาปัตยกรรมระบบ การจัดการข้อมูล การตรวจสอบความถูกต้องของโมเดล และการบริหารวงจรชีวิตของระบบ AI ให้มีความปลอดภัย เชื่อถือได้ และสอดคล้องกับข้อกำหนดขององค์กร

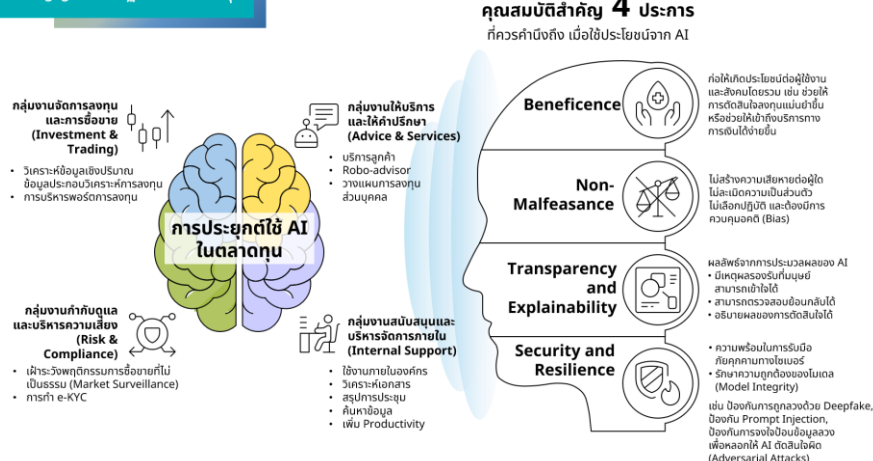
Customizer (ผู้ปรับแต่ง) หมายถึง บุคคลหรือทีมที่นำโมเดล ระบบ หรือบริการ AI ที่มีอยู่แล้วมาปรับแต่ง เชื่อมต่อ กำหนดค่า หรือออกแบบ workflow ให้เหมาะสมกับบริบทการใช้งานขององค์กร โดยไม่ได้รับผิดชอบต่อการพัฒนาหรือควบคุมองค์ประกอบหลักของโมเดล แต่สามารถกำหนดพฤติกรรม ขอบเขตการทำงาน ความสามารถ หรือระดับความเป็นอิสระของระบบ AI ผ่านการออกแบบการใช้งานและการเชื่อมต่อระบบ

User (ผู้ใช้งาน/ผู้ให้บริการต่อ) หมายถึง หน่วยงานหรือบุคคลที่นำระบบ AI ไปใช้งานจริงในกระบวนการทำงาน เช่น ฝ่ายธุรกิจ ฝ่ายปฏิบัติการ บริการลูกค้า การวิจัย การกำกับดูแล หรือใช้งานภายในองค์กร รวมถึงผู้ที่รับผิดชอบต่อการใช้งานให้เป็นไปตามนโยบายและมาตรการขององค์กร

1.3 การใช้งานปัญญาประดิษฐ์ในภาคตลาดทุน

ปัจจุบัน AI ได้รับการพัฒนาอย่างรวดเร็วและสามารถนำมาประยุกต์ใช้ได้หลากหลายรูปแบบ เพื่อให้การใช้งาน AI ประสบความสำเร็จ ผู้ประกอบธุรกิจจึงควรเริ่มจากการกำหนดวัตถุประสงค์ของการใช้งานที่ชัดเจน และเหมาะสม โดยคำนึงถึงโอกาสหรือประโยชน์ที่องค์กรจะได้รับ รวมถึงความเป็นไปได้ของการใช้งานซึ่งสามารถพิจารณาได้จากปัจจัยต่าง ๆ เช่น ต้นทุนในการลงทุนและดำเนินงาน ความสามารถทางเทคโนโลยี ความพร้อมของข้อมูล ความพร้อมของบุคลากรและทรัพยากรที่เกี่ยวข้อง เป็นต้น

การใช้งานปัญญาประดิษฐ์ในภาคตลาดทุน



ภาพที่ 2 การใช้งานปัญญาประดิษฐ์ในภาคตลาดทุน

จากการศึกษาข้อมูลทั้งภายในประเทศและต่างประเทศพบว่า ในปัจจุบัน ผู้ประกอบธุรกิจมีการใช้ AI เพื่อวัตถุประสงค์ในการเพิ่มประสิทธิภาพและประสิทธิผลของงาน โดยสามารถจัดกลุ่มแบ่งตามกิจกรรมสำคัญ ในภาคตลาดทุน ได้ 4 กิจกรรมหลัก ดังนี้

(1) กลุ่มงานจัดการลงทุนและการซื้อขาย (Investment & Trading)

หัวใจสำคัญของการขับเคลื่อนสภาพคล่องในตลาดทุน โดยมีการนำเทคโนโลยี ML มาใช้ในการวิเคราะห์ข้อมูลเชิงปริมาณและข้อมูลอื่นใดที่นำมาใช้ประกอบการวิเคราะห์การลงทุน (Alternative Data) เพื่อกำหนดกลยุทธ์การลงทุนอัตโนมัติและการทำหน้าที่เป็นผู้ดูแลสภาพคล่อง (Market Making) วัตถุประสงค์หลักเพื่อการสร้างผลตอบแทนส่วนเกิน (Alpha) และการบริหารจัดการพอร์ตการลงทุนให้มีความแม่นยำและรวดเร็ว

อย่างไรก็ตาม ความเสี่ยงที่ต้องระมัดระวังเป็นพิเศษคือ ความเสี่ยงเชิงระบบ (Systemic Risk) ที่เกิดจากพฤติกรรมแห่ตามกันของอัลกอริทึม (Herding Behavior) จากผู้ประกอบธุรกิจจำนวนมาก ใช้โมเดล AI ที่มีโครงสร้างและเทรนด้วยข้อมูลชุดเดียวกัน ส่งผลให้ AI เหล่านี้จะตัดสินใจ "เทขาย" หรือ "ซื้อ" ไปในทางเดียวกันเมื่อเกิดเหตุการณ์บางอย่าง ส่งผลให้สภาพคล่องในตลาดหายไปทันทีและเกิดความผันผวนรุนแรง (Pro-cyclicality) ซึ่งอาจนำไปสู่สถานะตลาดผันผวนรุนแรงอย่างกะทันหัน หรือ Flash Crash นอกจากนี้ ความซับซ้อนของโมเดลที่อธิบายได้ยาก (Black-box) ยังเป็นอุปสรรคสำคัญต่อการตรวจสอบย้อนหลังเมื่อเกิดคำสั่งซื้อขายที่ผิดปกติ ผู้ประกอบธุรกิจจึงต้องให้ความสำคัญกับหลักการ Explainability และมี Human-in-the-loop ในขั้นตอนสุดท้าย เพื่อให้ยังคงสามารถควบคุมและยับยั้งเหตุการณ์ที่ไม่พึงประสงค์ได้อย่างทันการณ์ และสร้างความโปร่งใสต่อหน่วยงานกำกับดูแล

(2) กลุ่มงานให้บริการและให้คำปรึกษา (Advice & Services)

มิติที่กระทบต่อผู้ลงทุนรายย่อยโดยตรง ผ่านการใช้งาน Robo-advisor และการสร้าง AI-generated research เพื่อตอบโจทย์การวางแผนการลงทุนรายบุคคล วัตถุประสงค์ของการใช้งานในส่วนนี้คือการสร้างความเท่าเทียมในการเข้าถึงคำปรึกษาที่มีคุณภาพ (Accessibility) และการลดต้นทุนการให้บริการเพื่อให้เข้าถึงกลุ่มนักลงทุนในวงกว้างได้มากขึ้น

ทว่าในมิติดังกล่าวนี้อาจมีความเสี่ยงที่ควรให้ความสำคัญอย่างยิ่งกับความเหมาะสมในการให้คำแนะนำ (Suitability) เนื่องจาก AI อาจนำเสนอผลิตภัณฑ์ที่ไม่สอดคล้องกับความเสี่ยงที่แท้จริงของลูกค้า หรือเกิดความลำเอียงของอัลกอริทึม (Algorithmic Bias) ที่ส่งเสริมผลิตภัณฑ์ของบริษัทตนเองมากกว่าผลประโยชน์ของลูกค้า และบริษัทอาจต้องรับผิดชอบต่อบทวิเคราะห์การลงทุนที่สร้างจาก AI เสมือนพนักงานเขียนเอง หลักการ Non-Malfeasance และ Explainability จึงต้องถูกนำมาใช้อย่างเข้มข้น เพื่อให้มั่นใจว่าคำแนะนำนั้นเป็นธรรมและสามารถชี้แจงที่มาของคำแนะนำให้แก่ลูกค้าได้อย่างชัดเจน

(3) กลุ่มงานกำกับดูแลและบริหารความเสี่ยง (Risk & Compliance)

เครื่องมือสำคัญขององค์กรที่ใช้ AI ในการเฝ้าระวังพฤติกรรมการซื้อขายที่ไม่เป็นธรรม (Market Surveillance) และการพิสูจน์ตัวตนทางดิจิทัล (e-KYC) โดยมีวัตถุประสงค์เพื่อยกระดับความเชื่อมั่นในโครงสร้างพื้นฐานของตลาดทุนและเพิ่มประสิทธิภาพในการตรวจจับทุจริตแบบเรียลไทม์

ความเสี่ยงหลักในกลุ่มนี้ไม่ได้มาจากตัวโมเดลเพียงอย่างเดียว แต่รวมถึงภัยคุกคามภายนอกในรูปแบบ Adversarial Attacks เช่น การใช้ Deepfake เพื่อปลอมแปลงตัวตน หรือการจงใจป้อนข้อมูลลงเพื่อหลอกระบบตรวจจับธุรกรรมฟอกเงิน (AML) ผู้ประกอบธุรกิจจึงต้องมุ่งเน้นที่หลักการ Security เป็นลำดับแรกเพื่อรักษาความถูกต้องของข้อมูล และหลักการ Beneficence เพื่อให้มั่นใจว่าระบบมีความแม่นยำสูงพอที่จะไม่สร้างภาระให้แก่ลูกค้าผู้บริสุทธิ์จากความผิดพลาดของระบบ (False Positives)

(4) กลุ่มงานสนับสนุนและบริหารจัดการภายใน (Internal Support)

กิจกรรมสนับสนุนและการใช้งาน AI ภายในองค์กร สำหรับงานวิเคราะห์เอกสาร สรุปการประชุมหรือตอบคำถามทั่วไป แต่มีความสำคัญอย่างยิ่งในการเพิ่มศักยภาพการแข่งขัน โดยเฉพาะการนำ AI มาใช้ช่วยนักวิเคราะห์ในการสรุปข้อมูลเศรษฐกิจและร่างบทวิเคราะห์หลักทรัพย์ (Research Automation) รวมถึงการตรวจสอบความถูกต้องของสัญญาทางกฎหมาย วัตถุประสงค์หลักคือการเพิ่มผลิตภาพ (Productivity) และลดข้อผิดพลาดจากการปฏิบัติงานของมนุษย์

อย่างไรก็ตาม ความเสี่ยงที่สำคัญที่สุดคือ การรั่วไหลของข้อมูลความลับ (Data Privacy & Leakage) หากพนักงานนำข้อมูลภายในไปประมวลผลผ่านโมเดลสาธารณะ รวมถึงความเสี่ยงจากข้อมูลที่บิดเบือนหรือข้อมูลเท็จ (Hallucination) ที่ AI อาจสร้างขึ้นในบทวิเคราะห์ ซึ่งอาจส่งผลกระทบต่อชื่อเสียงขององค์กรและการตัดสินใจของนักลงทุนในระยะถัดมา หลักการ Beneficence, Security และมี Human-in-the-loop (HITL) ในการตรวจสอบผลลัพธ์ จึงต้องถูกนำมาใช้เพื่อให้การใช้งานภายในองค์กรเกิดประโยชน์สูงสุดภายใต้การควบคุมความปลอดภัยของข้อมูลที่รัดกุม

เพื่อช่วยให้องค์กรสามารถเชื่อมโยงการใช้งาน AI กับเป้าหมายเชิงยุทธศาสตร์ขององค์กรได้อย่างชัดเจน คู่มือฉบับนี้ได้จัดทำเพื่อให้ผู้ประกอบธุรกิจคำนึงถึงหลักคิด การส่งมอบคุณค่าที่พึงประสงค์ (value propositions) จากการใช้งาน AI ภายใต้กรอบการกำกับดูแลที่เหมาะสมผ่าน 4 มิติหลัก ได้แก่ Trust, Innovation, Efficiency และ Resilience (TIER) ทั้งนี้ กรอบ TIER ถูกพัฒนาขึ้นในฐานะกรอบสังเคราะห์เชิงผลลัพธ์ที่ต่อยอดจากแนวทางการใช้งาน AI/ML ในตลาดทุน การบริหารจัดการความเสี่ยงด้าน IT และแนวทางการสร้างความพร้อมในการดำเนินงานในสถานการณ์ไม่ปกติ (operational resilience) เพื่อช่วยให้องค์กรสามารถประเมินได้ว่าการใช้งาน AI ที่จัดให้มันนำไปสู่การสร้างคุณค่า การเพิ่มประสิทธิภาพ และการบริหารความเสี่ยงอย่างสมดุลได้จริงหรือไม่

- **Trust** การกำกับดูแลที่มุ่งสร้างความเชื่อมั่นต่อการใช้งาน AI ผ่านความโปร่งใส ความเท่าเทียม ความรับผิดชอบ และความสามารถในการตรวจสอบย้อนหลัง ซึ่งสอดคล้องกับแนวคิด AI Trustworthiness โดยในบริบทตลาดทุน มิติดังนี้มุ่งเสริมสร้างความเชื่อมั่นของผู้ลงทุน เพิ่มความน่าเชื่อถือของกระบวนการตัดสินใจ และสนับสนุนความโปร่งใสของระบบตลาดทุนโดยรวม (IOSCO, 2021; European Union, 2024; ISO/IEC, 2023)

- **Innovation** การบริหารความเสี่ยงเพื่อสนับสนุนให้เกิดนวัตกรรมอย่างสร้างสรรค์และรับผิดชอบต่อตามแนวทางของ NIST AI RMF และหน่วยงานกำกับดูแลที่ส่งเสริมการทดลองใช้งานอย่างปลอดภัย โดยมิตินี้มุ่งเปิดพื้นที่ให้องค์กรสามารถพัฒนาและทดลองใช้ AI ได้อย่างมั่นใจ ลดอุปสรรคด้านการกำกับดูแล และเร่งการนำนวัตกรรมไปใช้จริงในภาคตลาดทุน (NIST, 2023; MAS, 2019, 2023)

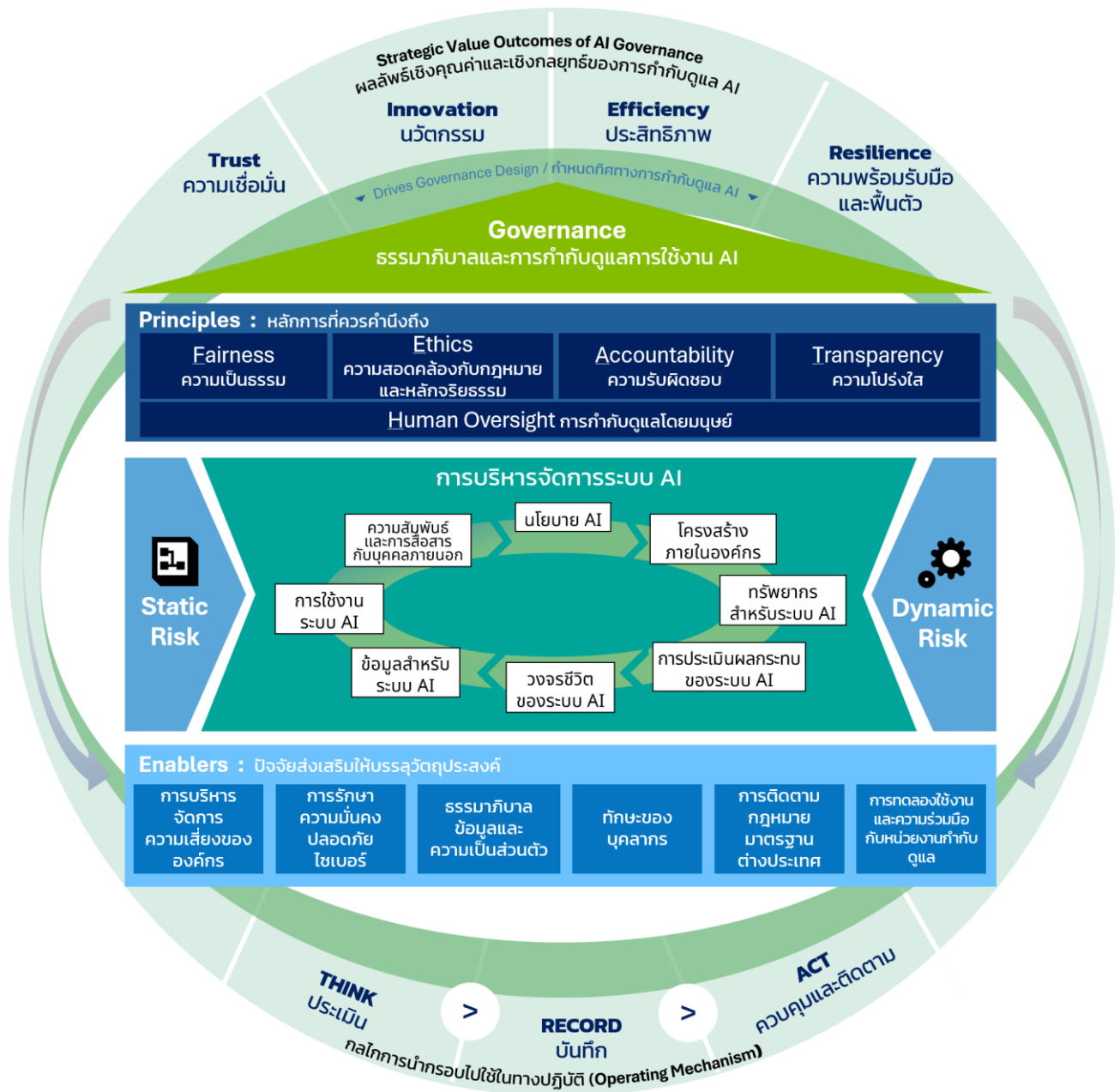
- **Efficiency** การใช้ AI เพื่อเพิ่มผลิตภาพและใช้ทรัพยากรอย่างคุ้มค่า ตามหลัก OECD AI Principles และแนวทางการนำเทคโนโลยีมาเพิ่มประสิทธิภาพของกระบวนการในภาคตลาดทุน มิตินี้ช่วยลดภาระการดำเนินงานที่ซ้ำซ้อน ยกกระดับคุณภาพข้อมูลและการตัดสินใจ และเพิ่มความคล่องตัวในการดำเนินธุรกิจ (OECD, 2019; IOSCO, 2021) และ

- **Resilience** ความสามารถของระบบตลาดทุนในการรองรับและฟื้นตัวจากความเสี่ยงที่เกิดจากการใช้งาน AI โดยอ้างอิงแนวคิดด้าน systemic และ operational resilience จาก BIS และ IOSCO มิตินี้มุ่งเสริมความมั่นคงของระบบตลาดทุน ลดผลกระทบเชิงระบบ และสร้างความพร้อมในการรับมือกับเหตุการณ์ไม่คาดคิดหรือความล้มเหลวของระบบ AI และดำเนินการได้อย่างต่อเนื่อง (BIS, 2021; BIS, 2024; IOSCO, 2024)

การบรรลุผลลัพธ์เชิงยุทธศาสตร์ดังกล่าวขึ้นอยู่กับสิ่งที่องค์กรสามารถกำหนดระดับของการกำกับดูแล มาตรการควบคุม และการติดตามผลให้เหมาะสมกับลักษณะการใช้งาน ระดับความเสี่ยง และนัยสำคัญของแต่ละ use case ของ AI ได้อย่างเหมาะสมตามหลัก Proportionality มากน้อยเพียงใด

หลัก Proportionality ในการกำกับดูแลการใช้งาน AI คือการกำหนดระดับความเข้มงวดของการกำกับดูแลให้สอดคล้องกับระดับความเสี่ยง บริบทการใช้งาน และผลกระทบที่อาจเกิดขึ้นจริง รวมถึงระดับอิทธิพลของ AI ต่อการตัดสินใจของผู้ใช้และผู้ร่วมตลาด (Risk-based, Context-driven and Outcome-oriented Governance) โดยมุ่งสร้างสมดุลระหว่างการคุ้มครองผู้ลงทุน การรักษาความเชื่อมั่นของตลาดทุน และการสนับสนุนนวัตกรรมอย่างยั่งยืน

2. กรอบการกำกับดูแลการใช้งาน AI (AI Governance Framework)



ภาพที่ 3 กรอบการกำกับดูแลการใช้งาน AI (AI Governance Framework)

กรอบการกำกับดูแลการใช้งาน AI (AI Governance Framework) ประกอบด้วย 5 ส่วน ในลักษณะของ layered governance framework ดังนี้

- (1) การกำกับดูแล (Governance Layer)
- (2) หลักการที่ควรคำนึงถึง (Principles Layer)
- (3) การบริหารจัดการความเสี่ยง (Risk Management Layer)
- (4) แนวทางบริหารจัดการระบบ AI (AIMS Controls Layer)
- (5) ปัจจัยสนับสนุน (Enablers Layer)

ทั้งนี้ ผู้ประกอบธุรกิจสามารถกำหนดวัตถุประสงค์เชิงคุณค่าและเชิงกลยุทธ์ของการกำกับดูแล AI และกลไกการนำกรอบไปใช้ เพื่อสนับสนุนให้เกิดการใช้งาน AI อย่างรับผิดชอบ สมเหตุสมผลกับความเสี่ยง และสร้างคุณค่าต่อภาคตลาดทุนใน 4 มิติ ได้แก่ Trust, Innovation, Efficiency และ Resilience (TIER) ซึ่งเป็นผลลัพธ์เชิงคุณค่าและเชิงกลยุทธ์ที่กรอบการกำกับดูแลมุ่งส่งเสริม โดยโครงสร้างกรอบการกำกับดูแลทั้ง 5 ส่วนมีความเชื่อมโยงและถูกนำไปใช้ผ่านกลไกการนำไปใช้งานจริง (Operating Mechanism: Think > Record > Act) กล่าวคือ (1) Think คือการพิจารณา use case และจัดระดับความเสี่ยงผ่านกระบวนการ 12 ขั้นตอนและเกณฑ์การประเมินความเสี่ยง (Risk Rubrics), (2) Record คือการบันทึกข้อมูลและเหตุผลใน AI Inventory เพื่อสร้าง visibility, accountability และ traceability, และ (3) Act คือการเลือกและดำเนินการมาตรการควบคุมตามระบบการบริหารจัดการ AI (AI Management System: AIMS) ให้เหมาะสมกับบริบทและความเสี่ยงของการใช้งาน

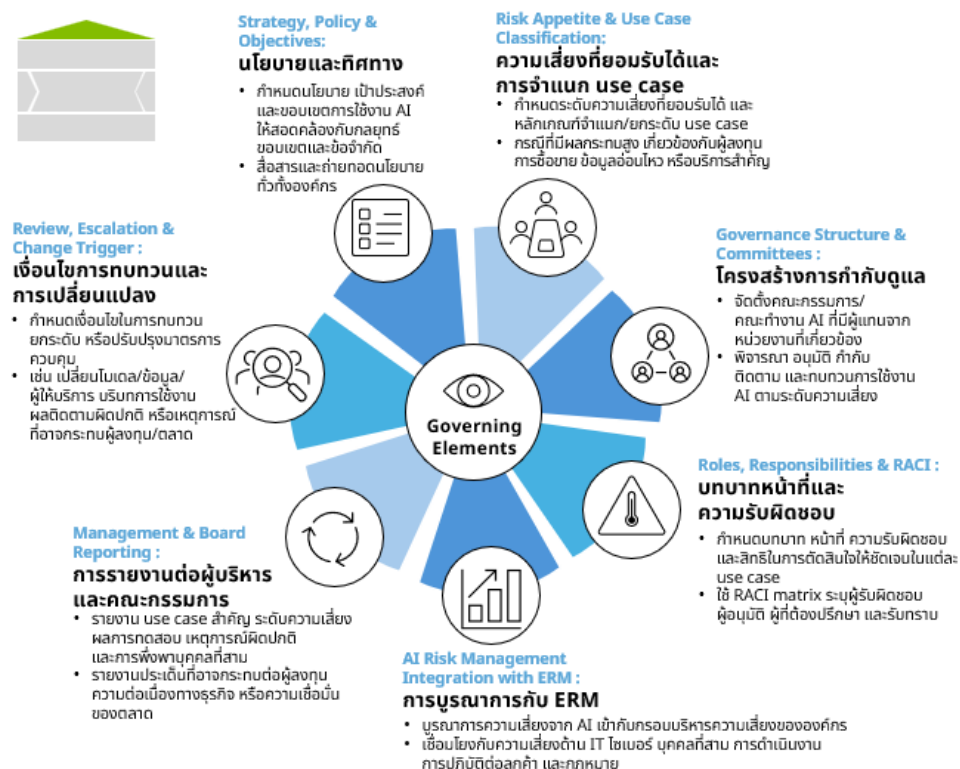
2.1 การกำกับดูแล (Governance)

การกำกับดูแลเป็นองค์ประกอบสำคัญที่ทำให้การนำ AI มาใช้งานในองค์กรเกิดความโปร่งใส มีจริยธรรม มีความรับผิดชอบและปลอดภัย ผ่านการจัดตั้งโครงสร้าง นโยบาย กระบวนการควบคุม oversight มาตรฐาน และการตรวจสอบที่ครอบคลุมตลอดวงจรชีวิตของระบบ AI ผู้ประกอบธุรกิจควรจัดให้มีกรอบการกำกับดูแลการใช้งาน AI ที่ดี (AI Governance Framework) โดยมุ่งเน้นการมีส่วนร่วมของผู้บริหาร การกำหนดบทบาทหน้าที่ที่ชัดเจน มีการติดตาม ประเมินผลและรายงานการประยุกต์ใช้ AI รวมถึงการปฏิบัติตามกรอบของกฎหมายและหน่วยงานกำกับดูแลในส่วนที่เกี่ยวข้อง และควรจัดให้มีกรอบการกำกับดูแลความเสี่ยงจากการใช้งาน AI ที่ครอบคลุมอย่างน้อยในประเด็นต่อไปนี้

- (1) การกำหนดทิศทาง เป้าประสงค์ นโยบาย และขอบเขตการใช้งาน AI (Strategy, Policy & Objectives) ให้สอดคล้องกับกลยุทธ์ขององค์กร ลักษณะธุรกิจ ความพร้อมของบุคลากร ระบบงาน ข้อมูล และระดับความเสี่ยงที่เกี่ยวข้อง
- (2) การกำหนดระดับความเสี่ยงที่ยอมรับได้ในการใช้งาน AI รวมถึงหลักเกณฑ์ในการจำแนกและยกระดับการพิจารณากรณีการใช้งาน (Risk Appetite & Use Case Classification) โดยเฉพาะกรณีที่มีผลกระทบสูง มีความซับซ้อนสูง เกี่ยวข้องกับผู้ลงทุน หรือบริการสำคัญ
- (3) การกำหนดกลไกกำกับดูแลที่เหมาะสม (Governance Structure & Committees) เช่น คณะกรรมการ คณะทำงานที่มีผู้แทนจากหน่วยงานที่เกี่ยวข้อง เพื่อพิจารณา อนุมัติ กำกับ ติดตาม และทบทวนการใช้งาน AI ตามระดับความเสี่ยง
- (4) การกำหนดบทบาท หน้าที่ ความรับผิดชอบ และสิทธิในการตัดสินใจของผู้เกี่ยวข้องในแต่ละ AI use case ให้ชัดเจน (Roles, Responsibilities & RACI) โดยอาจใช้ RACI matrix หรือเอกสารเทียบเท่า เพื่อระบุผู้รับผิดชอบหลัก ผู้มีอำนาจอนุมัติ ผู้ที่ต้องได้รับการปรึกษาหารือ และผู้ที่ต้องได้รับทราบ
- (5) การบูรณาการความเสี่ยงจากการใช้งาน AI เข้ากับกรอบการบริหารความเสี่ยงขององค์กร (AI Risk Management Integration with Enterprise Risk Management (“ERM”)) โดยเชื่อมโยงกับความเสี่ยงด้านเทคโนโลยีสารสนเทศ ความมั่นคงปลอดภัยไซเบอร์ บุคคลที่ 3 การดำเนินงาน การปฏิบัติต่อลูกค้า และความเสี่ยงด้านกฎหมายหรือกฎเกณฑ์ที่เกี่ยวข้อง

- (6) การจัดให้มีกลไกรายงานต่อผู้บริหารระดับสูงหรือคณะกรรมการตามความเหมาะสม (Management & Board Reporting) โดยครอบคลุม use case ที่สำคัญ ระดับความเสี่ยง ผลการทดสอบหรือทวนสอบ เหตุการณ์ผิดปกติ การพึ่งพาผู้ให้บริการภายนอก และประเด็นที่อาจกระทบต่อผู้ลงทุน ความต่อเนื่องทางธุรกิจ หรือความเชื่อมั่นของตลาด
- (7) การกำหนดเงื่อนไขในการทบทวน ยกระดับ หรือปรับปรุงมาตรการควบคุมเมื่อมีการเปลี่ยนแปลงที่มีนัยสำคัญ (Review, Escalation & Change Trigger) เช่น การเปลี่ยนโมเดล แหล่งข้อมูล ผู้ให้บริการ ระดับความเป็นอิสระของระบบ บริบทการใช้งาน ผลการติดตามที่ผิดปกติ หรือเหตุการณ์ที่อาจกระทบต่อผู้ลงทุนหรือตลาด

ทั้งนี้ คณะกรรมการหรือที่ปรึกษาของคณะกรรมการของผู้ประกอบธุรกิจควรมีความรู้ ความเชี่ยวชาญ หรือประสบการณ์ที่เกี่ยวข้องกับการประยุกต์ใช้ AI อย่างเพียงพอ และพิจารณาให้การใช้งาน AI เชื่อมโยงกับกลยุทธ์และเป้าหมายทางธุรกิจขององค์กรอย่างชัดเจน โดยมีการสนับสนุนจากคณะกรรมการหรือผู้บริหารระดับสูง มีผู้รับผิดชอบหลักหรือกลไกกำกับดูแลที่เหมาะสม มีการจัดสรรทรัพยากรตามระดับความสำคัญ และความเสี่ยงของ use case และองค์กรควรติดตามพัฒนาการของ AI และทบทวนแผนการนำ AI ไปใช้ อย่างต่อเนื่อง เพื่อให้การกำกับดูแล AI ไม่จำกัดอยู่เพียงการควบคุมความเสี่ยง แต่สนับสนุนการนำ AI ไปใช้จริง อย่างรับผิดชอบและสร้างคุณค่าได้ นอกจากนี้ ควรสร้างการมีส่วนร่วมในการบริหารจัดการจากความร่วมมือ ในหลายภาคส่วน และให้สอดคล้องกับบริบทขององค์กร องค์กรสามารถศึกษาตัวอย่างแผนภูมิแสดงความรับผิดชอบ ตามหลักการ RACI ได้จากแนวปฏิบัติการใช้ปัญญาประดิษฐ์อย่างมั่นคงปลอดภัย สำนักงานคณะกรรมการ รักษาความมั่นคงปลอดภัยไซเบอร์แห่งชาติ (“สกมช.”)



ภาพที่ 4 การกำกับดูแล (Governance)

2.2 หลักการที่ควรคำนึงถึง (Principles)

หลักการพื้นฐานที่ควรคำนึงถึงเพื่อให้การใช้งาน AI บรรลุวัตถุประสงค์ที่กำหนดไว้ ได้แก่

2.2.1 ความเป็นธรรม (Fairness)

AI ควรถูกออกแบบและพัฒนาโดยคำนึงถึงความเป็นธรรม ความยุติธรรม ความเท่าเทียม และความหลากหลายทางสังคม เพื่อไม่ให้บุคคลหรือกลุ่มบุคคลบางส่วนได้รับการเลือกปฏิบัติอย่างไม่สมเหตุสมผล รวมถึงให้โอกาสประชาชนทุกคน รวมถึงกลุ่มคนด้อยโอกาสและผู้ทุพพลภาพ ได้รับประโยชน์จาก AI ได้อย่างทั่วถึง และเท่าเทียม การใช้งาน AI อาจก่อให้เกิดผลที่ไม่เป็นธรรมได้จากสาเหตุหลัก 2 สาเหตุ ได้แก่

(1) อคติที่มาจากข้อมูล (Data Bias) เนื่องจากข้อมูลที่ใช้ในการสอนโมเดลเอนเอียง ไม่ครอบคลุมทุกกลุ่มประชากร หรือมีขนาดกลุ่มตัวอย่างที่ไม่เหมาะสม ทำให้เกิดการเลือกปฏิบัติ ตัวอย่างเช่น AI คัดเลือกผู้สมัครงานเพศชายมากกว่าเพศหญิง อันมีสาเหตุจากข้อมูลส่วนใหญ่ที่ใช้สอนโมเดลนั้นเป็นข้อมูลของเพศชาย จึงทำให้ AI ให้น้ำหนักในการเลือกผู้สมัครงานเพศชายมากกว่าผู้สมัครงานเพศหญิง เป็นต้น โดยการป้องกันการเกิดอคติที่มาจากข้อมูล (Data Bias) สามารถทำได้โดยการเลือกข้อมูลที่ใช้ในกระบวนการสอนโมเดล ด้วยความระมัดระวังและคำนึงถึงมาตรการในการคัดเลือกข้อมูล ตัวอย่างข้อมูลที่ใช้ และการทดสอบโมเดล เพื่อค้นหาความผิดพลาด และดำเนินการปรับปรุงแก้ไขต่อไป

(2) อคติที่มาจากการออกแบบและสร้างโมเดล (Bias Introduced by Engineering Decisions) ซึ่งอาจเกิดได้จากความผิดพลาดในการตัดสินใจของบุคคลที่เกี่ยวข้องในการออกแบบและสร้างโมเดล ทั้งโดยเจตนา และไม่เจตนา รวมถึงกระบวนการออกแบบ และสร้างโมเดลที่ไม่มีมาตรการควบคุมที่เพียงพอ ซึ่งสามารถป้องกันอคตินี้ได้โดยกำหนดให้บุคลากรที่มีส่วนร่วมในกระบวนการสร้างโมเดล มีความหลากหลาย และมีกระบวนการสอบทาน (Review) และทดสอบ (Test) ในขั้นตอนการสร้างโมเดลที่ดี รวมถึง (ในกรณีที่จำเป็น) จัดให้มีผู้เชี่ยวชาญด้านการประยุกต์ใช้ AI ให้คำปรึกษา เพื่อให้สามารถมองเห็นความเสี่ยงจากอคติที่อาจเกิดขึ้น และกำหนดแนวทางการลดความเสี่ยงดังกล่าว

2.2.2 ความสอดคล้องกับกฎหมายและหลักจริยธรรม (Ethics)

การใช้ AI ได้อย่างเหมาะสมและมีประสิทธิภาพนั้น จำเป็นต้องมีความสอดคล้องกับกฎหมาย และหลักจริยธรรมที่เกี่ยวข้อง โดยผู้ประกอบธุรกิจควรดำเนินการ ดังนี้

(1) รวบรวมกฎหมายและหลักจริยธรรม ซึ่งรวมถึงหลักปฏิบัติ ค่านิยม พันธกิจ และนโยบายขององค์กรที่เกี่ยวข้องกับการใช้งาน AI

(2) พิจารณาความเหมาะสมของการใช้งาน AI โดยคำนึงถึงความสอดคล้องกับกฎหมาย และหลักจริยธรรมตามข้อ (1)

(3) ใช้กฎหมายและหลักจริยธรรมตามข้อ (1) เป็นพื้นฐานในการออกแบบ พัฒนา และใช้งาน AI ขององค์กร

2.2.3 ความรับผิดชอบ (Accountability)

หลักการในด้านความรับผิดชอบนี้มุ่งที่การกำหนดความรับผิดชอบ (Accountability) และความเป็นเจ้าของ (Ownership) สำหรับกิจกรรมที่มีการใช้งาน AI อย่างชัดเจน เพื่อให้การนำ AI มาใช้งานเป็นไปอย่างมีประสิทธิภาพ

หลักการในด้านความรับผิดชอบแบ่งออกเป็น 2 ด้านที่สำคัญ ดังนี้

(1) ความรับผิดชอบภายในองค์กร (Internal Accountability)

ผู้ประกอบการที่มีการนำเทคโนโลยี AI มาใช้งาน ควรขยายขอบเขตหน้าที่และความรับผิดชอบของผู้บริหารระดับสูง (Senior Management) ให้ครอบคลุม ตั้งแต่การอนุมัตินโยบายการใช้งานไปจนถึงการติดตามผลและความเสี่ยงที่อาจเกิดขึ้นจากการใช้งานระบบ AI การมอบหมายอำนาจหน้าที่ในการตัดสินใจเกี่ยวกับ AI ควรคำนึงถึงระดับของผลกระทบที่เกิดจากการตัดสินใจและความสอดคล้องกับกรอบการกำกับดูแลของผู้ประกอบธุรกิจ (Internal Governance Framework) ตัวอย่างเช่น ผู้ประกอบธุรกิจที่กำหนดให้ Chief Financial Officer (“CFO”) เป็นผู้มีอำนาจในการอนุมัติเรื่องการซื้อขายหน่วยลงทุน เมื่อนำ AI มาใช้ส่งคำสั่งซื้อขายหน่วยลงทุนแบบอัตโนมัติ ผู้ประกอบธุรกิจควรกำหนดให้ CFO หรือผู้บริหารระดับสูงกว่า มีส่วนร่วมในการอนุมัติโมเดลของ AI และข้อมูลที่ใช้งานโดย AI เพื่อให้สอดคล้องกับกรอบการกำกับดูแลขององค์กร อย่างไรก็ตาม หาก CFO มีความรู้ ความเชี่ยวชาญ และประสบการณ์ด้าน AI ที่ยังไม่เพียงพอ ผู้ประกอบธุรกิจสามารถแต่งตั้งที่ปรึกษาด้านการประยุกต์ใช้ AI หรือแต่งตั้งคณะกรรมการที่มี CFO เป็นองค์ประกอบเพื่อรับผิดชอบในเรื่องดังกล่าวตามความเหมาะสม เป็นต้น

(2) ความรับผิดชอบภายนอกองค์กร (External Accountability)

ผู้ประกอบการควรกำหนดช่องทางการติดต่อเพื่อให้ผู้ใช้งาน ลูกค้า เจ้าของข้อมูลส่วนบุคคล หรือผู้ที่อาจได้รับผลกระทบจากการทำงานของ AI สามารถติดต่อสอบถามข้อมูล ร้องเรียน และแจ้งประเด็นปัญหาหรือความผิดพลาดเกี่ยวกับการใช้งาน AI ได้อย่างสะดวก ซึ่งจะช่วยให้ผู้ประกอบการรับทราบประเด็นปัญหาและความผิดพลาดที่เกิดขึ้น และดำเนินการแก้ไขได้อย่างทันท่วงที

นอกจากนี้ คำถาม ข้อร้องเรียน ประเด็นปัญหา และความผิดพลาดต่าง ๆ ที่ได้รับแจ้ง ควรถูกรวบรวมและใช้ประโยชน์สำหรับการตรวจสอบและทบทวนในภายหลัง และวางแผนป้องกันปัญหาที่อาจเกิดขึ้นในอนาคต

2.2.4 ความโปร่งใส (Transparency)

ความโปร่งใสเป็นหลักการที่สำคัญประการหนึ่งที่ผู้ประกอบการควรคำนึงถึงในการนำ AI มาใช้งาน เนื่องจาก AI เป็นเทคโนโลยีที่ซับซ้อนและยากในการเข้าใจกลไกการทำงาน แม้ว่าจะมีพัฒนาการด้านความแม่นยำที่เพิ่มมากขึ้นกว่าในอดีต แต่ข้อผิดพลาดจากการทำงานของ AI ยังคงมีโอกาสเกิดขึ้นได้ในระดับที่มากขึ้นแตกต่างกันไปตามโมเดลและข้อมูลที่ใช้งาน ดังนั้น การใช้งาน AI จึงควรคำนึงถึงเรื่องความโปร่งใส โดยต้องสามารถอธิบายได้ว่า AI ถูกนำมาใช้ในกระบวนการ (Process) หรือระบบงานใด และหลักการทำงานหรือการตัดสินใจของ AI เป็นอย่างไร (Explainability) นอกจากนี้ ความโปร่งใส ยังแสดงให้เห็นได้จากความสามารถในการตรวจสอบข้อมูลย้อนหลังสำหรับกิจกรรมที่เกิดขึ้น (Traceability) ว่าเป็นไปตามหลักการหรือความตั้งใจที่แจ้งไว้ต่อลูกค้าหรือไม่

ระดับของความโปร่งใสที่เหมาะสมในการเปิดเผยข้อมูลเกี่ยวกับการใช้งาน AI นั้น ขึ้นอยู่กับบริบท และลักษณะของการใช้งาน AI ผู้ประกอบธุรกิจควรพิจารณาการเปิดเผยข้อมูลที่สะท้อนถึงขอบเขตและข้อจำกัด ของการใช้งาน รวมถึงการเปลี่ยนแปลงเงื่อนไขโมเดล ตลอดจนความเสี่ยงและความปลอดภัยของ กิจกรรมที่มีการใช้งาน AI

อย่างไรก็ดี ผู้ประกอบธุรกิจควรตรวจสอบให้มั่นใจว่าข้อความประชาสัมพันธ์ หรือเนื้อหา ของรายงานต่อผู้ลงทุน หรือเอกสารเกี่ยวกับการตลาดที่อ้างถึงการใช้งาน AI ขององค์กรนั้น มีความถูกต้อง มีหลักฐานรองรับ และไม่ทำให้เข้าใจผิดเกี่ยวกับขอบเขต ความสามารถ หรือประสิทธิภาพที่แท้จริงของระบบ (AI-washing) โดยอย่างน้อยควรพิจารณาถึง (1) การระบุขอบเขตการใช้งานตามข้อเท็จจริง (2) การเปิดเผย ข้อจำกัดและความเสี่ยงควบคู่กับประโยชน์ (3) มีกระบวนการทบทวนข้อความที่อ้างถึง AI ก่อนเผยแพร่ และ (4) ไม่อ้างผลการทดสอบหรือประสิทธิภาพโดยไม่มีวิธีการวัดที่สามารถตรวจสอบย้อนหลังได้

นอกจากนี้ ผู้ประกอบธุรกิจควรจัดเก็บข้อมูลที่เกี่ยวข้องกับกระบวนการสร้างและทดสอบโมเดล¹ หรือบันทึกเหตุการณ์ (Log) อย่างเพียงพอและเหมาะสมกับความเสี่ยงของการใช้งาน AI เพื่อรองรับ การตรวจสอบในภายหลังเมื่อพบประเด็นปัญหา หรือเมื่อได้รับการร้องขอจากลูกค้าหรือผู้ตรวจสอบ (Auditor) ตามความสมเหตุสมผล

จากหลักการความเป็นธรรม จริยธรรม ความรับผิดชอบ และความโปร่งใส เชิงปฏิบัติ (FEAT in Practice) ข้างต้น เพื่อให้สามารถนำไปใช้ได้จริงในระดับปฏิบัติ ผู้ประกอบธุรกิจควรพิจารณาหลักความเป็นธรรม (Fairness) หลักจริยธรรม (Ethics) หลักความรับผิดชอบ (Accountability) และหลักความโปร่งใส (Transparency) ตามข้อ 2.2.1-2.2.4 แบบบูรณาการ โดยมีใช้พิจารณาแยกขาดจากกัน กล่าวคือ การออกแบบหรือการใช้งาน AI ที่มีความถูกต้องทางเทคนิคเพียงอย่างเดียว อาจยังไม่เพียงพอ หากผลลัพธ์มีความเอนเอียง (bias) ขาดคำอธิบาย ที่เหมาะสม หรือทำให้ผู้ใช้บริการไม่สามารถทราบขอบเขตและข้อจำกัดของระบบได้อย่างชัดเจน

ดังนั้น ผู้ประกอบธุรกิจจึงควรกำหนดตัวอย่างเชิงปฏิบัติของแต่ละหลักการให้สอดคล้องกับบริบทของ องค์กร เช่น ความเป็นธรรม (Fairness) อาจสะท้อนผ่านการทดสอบ bias และการทบทวนผลลัพธ์ข้ามกลุ่มลูกค้า ด้านจริยธรรม (Ethics) อาจสะท้อนผ่านการกำหนดขอบเขตการใช้งานที่ไม่ก่อให้เกิดอันตรายหรือการใช้ข้อมูล เกินความจำเป็น ความรับผิดชอบ (Accountability) อาจสะท้อนผ่านการกำหนดผู้มีอำนาจตัดสินใจ และการเก็บ หลักฐานอย่างเพียงพอ ส่วนความโปร่งใส (Transparency) อาจสะท้อนผ่านการเปิดเผยข้อมูลที่เหมาะสมแก่ ผู้ใช้บริการและความสามารถในการตรวจสอบย้อนหลัง

2.2.5 การกำกับดูแลโดยมนุษย์ตามระดับความเสี่ยง (Human Oversight)

ผู้ประกอบธุรกิจควรคำนึงถึงรูปแบบการกำหนดระดับการกำกับดูแลโดยมนุษย์ (Human Oversight Configuration) ให้เหมาะสมกับความเสี่ยงและผลกระทบของ use case โดยอาจอยู่ในรูปแบบ

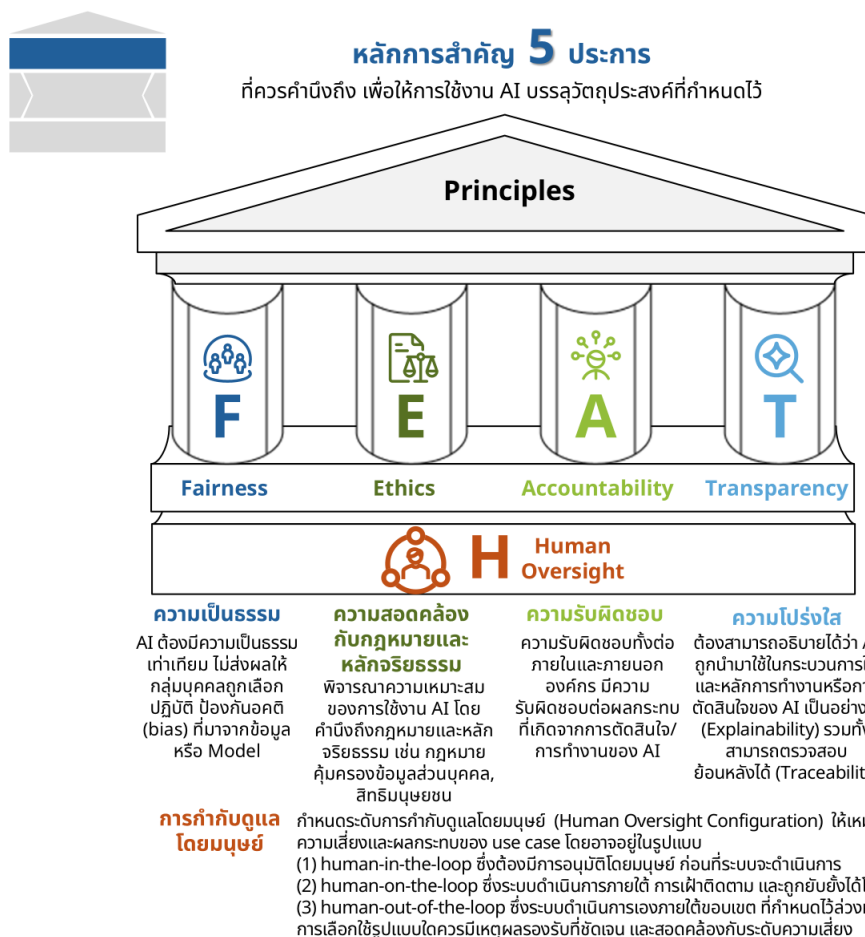
(1) human-in-the-loop ซึ่งต้องมีการอนุมัติโดยมนุษย์ ก่อนที่ระบบจะดำเนินการ หรือ

¹ ตัวอย่างข้อมูลที่ควรจัดเก็บ เช่น แหล่งที่มาของข้อมูล (Data Provenance) การประมวลผลหรือเปลี่ยนแปลงใด ๆ ที่เกิดขึ้นกับข้อมูลในกระบวนการ จัดเตรียมข้อมูล (Data Lineage) การออกแบบและการทำงานของอัลกอริทึม และชุดข้อมูลพร้อมกับผลลัพธ์ในกระบวนการสอน (Training) และทดสอบ โมเดล (Test) เป็นต้น

(2) human-on-the-loop ซึ่งระบบสามารถดำเนินการได้ภายใต้การเฝ้าติดตามและสามารถถูกยับยั้งได้โดยมนุษย์ หรือ

(3) human-out-of-the-loop ซึ่งระบบดำเนินการเองภายใต้ขอบเขตที่กำหนดไว้ล่วงหน้า

ทั้งนี้ การเลือกใช้รูปแบบใดควรมีเหตุผลรองรับที่ชัดเจนและสอดคล้องกับระดับความเสี่ยงของกิจกรรมนั้น สำหรับ use case ที่เกี่ยวข้องกับผู้ลงทุน การสื่อสารกับลูกค้า การเข้าถึงข้อมูลอ่อนไหว การซื้อขาย หรือการส่งคำสั่งแบบอัตโนมัติ หรือการกระทำใดที่อาจก่อให้เกิดผลกระทบในวงกว้าง ผู้ประกอบธุรกิจควรให้ความสำคัญกับการมีมนุษย์อยู่ในวงจรการตัดสินใจหรืออย่างน้อยต้องสามารถยับยั้งหรือย้อนกลับการกระทำของระบบได้ภายในระยะเวลาที่เหมาะสมกับความเสี่ยง



ภาพที่ 5 หลักการที่ควรคำนึงถึง (Principles)

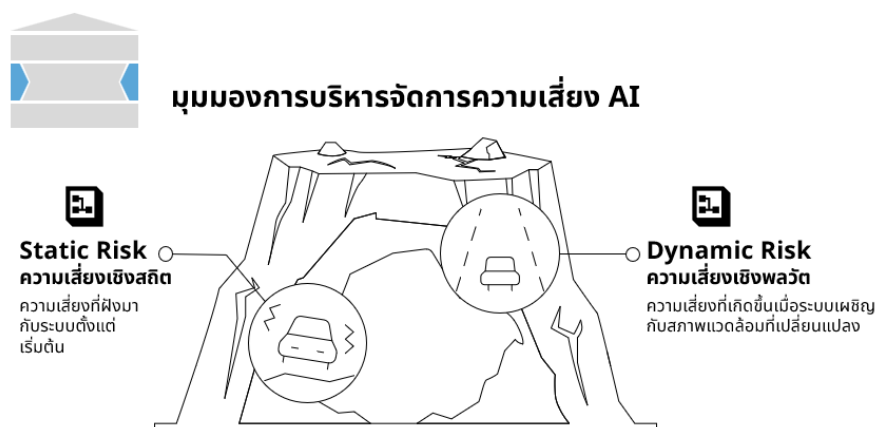
2.3 การบริหารจัดการความเสี่ยง (Risk Management)

AI เป็นเทคโนโลยีที่กำลังมีบทบาทสำคัญในภาคตลาดทุน เนื่องจากมีศักยภาพสูงในการวิเคราะห์ข้อมูล และช่วยในการตัดสินใจซึ่งสามารถช่วยลดค่าใช้จ่ายและต้นทุนในการดำเนินธุรกิจ พร้อมทั้งเพิ่มประสิทธิภาพในการทำงานให้ดียิ่งขึ้น อย่างไรก็ตาม การใช้งาน AI ในการดำเนินธุรกิจย่อมนำมาซึ่งความเสี่ยงใหม่ ๆ ที่หากขาดการบริหารจัดการอย่างเหมาะสม อาจนำมาซึ่งความเสียหายด้านการเงิน ภาพลักษณ์องค์กร และเกิดการละเมิดกฎระเบียบที่เกี่ยวข้องได้ โดยเฉพาะอย่างยิ่ง การใช้งาน AI ที่อาจกระทบต่อความมั่นคงของระบบการเงินและเสถียรภาพของระบบเศรษฐกิจโดยรวม ผู้ประกอบธุรกิจควรพิจารณาความเสี่ยงและผลกระทบอย่างรอบคอบ ก่อนนำ AI ไปพัฒนาหรือประยุกต์ใช้กับระบบบริการ และทบทวน ติดตามความเสี่ยงและปรับปรุงอย่างต่อเนื่อง ให้เป็นไปตามกฎหมายและหลักเกณฑ์ที่เกี่ยวข้องทั้งในปัจจุบันและอนาคต

กรอบการกำกับดูแลฉบับนี้ใช้แนวคิดการจำแนกความเสี่ยงของ AI โดยจำแนกตามลักษณะพฤติกรรมของระบบและปฏิสัมพันธ์ที่มีต่อสภาพแวดล้อม เพื่อให้เข้าใจว่าความเสี่ยงนั้นมีที่มาอย่างไรและให้ครอบคลุมตลอดวงจรชีวิต และเน้นย้ำว่าการกำกับดูแลการใช้งาน AI ไม่ใช่การประเมินความเสี่ยงตอนเริ่มต้นเพียงครั้งเดียว แต่เป็นกระบวนการต่อเนื่องที่ต้องอาศัยการตัดสินใจอย่างมีวิจารณญาณของมนุษย์ ควบคู่ไปกับการติดตามผลตลอดอายุการใช้งานของระบบตามวงจรชีวิตของ AI โดยสามารถแบ่งเป็น 2 มิติหลัก

(1) ความเสี่ยงเชิงสถิต (Static Risk): เปรียบเสมือน "ตัวรถและเครื่องยนต์" คือความเสี่ยงเชิงโครงสร้างที่ฝังมากระบบตั้งแต่วันแรก ทั้งจากข้อมูลที่เลือกใช้และตรรกะของแบบจำลอง ซึ่งเราสามารถตรวจสอบความพร้อมและควบคุมคุณภาพได้ทันทีในขั้นตอนการออกแบบและทดสอบก่อนออกวิ่งจริง (Mitchell et al., 2019; NIST, 2023)

(2) ความเสี่ยงเชิงพลวัต (Dynamic Risk): เปรียบเสมือน "พฤติกรรมการขับขึ้นบนถนนจริง" คือ ความเสี่ยงที่อุบัติขึ้นเมื่อระบบต้องเผชิญกับสภาพตลาดที่ผันผวนและปฏิสัมพันธ์กับผู้ใช้ ซึ่งพฤติกรรมเหล่านี้ อาจเปลี่ยนแปลงหรือทวีความรุนแรงขึ้นได้ตลอดเวลาตามสถานการณ์ที่เกิดขึ้นจริง จึงต้องอาศัยการเฝ้าระวังและปรับแต่งอย่างต่อเนื่องแม้จะผ่านการทดสอบมาแล้วก็ตาม (Gama et al., 2014; NIST, 2023) หัวใจสำคัญของการบริหารจัดการ Dynamic AI Risk คือ การกำหนดให้มีมนุษย์อยู่ในวงจรการตัดสินใจ (Human-in-the-loop) เพื่อปรับทิศทางหรือยับยั้งการทำงานของระบบให้สอดคล้องกับความเสี่ยงที่ยอมรับได้และบริบททางกฎหมายที่เปลี่ยนแปลงไป



ภาพที่ 6 มุมมองการบริหารจัดการความเสี่ยง AI (Risk Lens)

ความเสี่ยงจากการประยุกต์ใช้ AI มีความหลากหลายและแตกต่างกันไปตามบริบทของการประยุกต์ใช้ โดยหากพิจารณาในบริบทของตลาดทุน มีความเสี่ยงที่เกิดจากการใช้งานระบบ AI ดังนี้

(1) ด้านการสร้างผลกระทบต่อกลไกตลาด (Market Integrity & Abuse Risks)

องค์กรที่นำระบบ AI หรือ GenAI มาใช้ในกระบวนการซื้อขายหลักทรัพย์หรือสื่อสารกับลูกค้า อาจเผชิญความเสี่ยงที่ระบบสร้างรูปแบบคำสั่งซื้อขายซึ่งคล้ายพฤติกรรมต้องห้าม เช่น spoofing/layering โดยไม่ได้ตั้งใจ หรือสร้างเนื้อหาสื่อสาร/โฆษณาที่ทำให้ผู้ลงทุนเข้าใจผิด เช่น การกล่าวอ้างเกินจริงเกี่ยวกับประสิทธิภาพของ AI (“AI-washing”) หากไม่มีการกำกับดูแลและทดสอบโมเดลอย่างเพียงพอ อาจนำไปสู่การละเมิดกฎหมายหรือข้อบังคับด้าน Market Abuse ซึ่งส่งผลเสียต่อชื่อเสียงองค์กร ความเชื่อมั่นของผู้ลงทุน และอาจมีผลทางกฎหมาย ตัวอย่างมาตรการป้องกัน ได้แก่ การตรวจสอบและวิเคราะห์ผลลัพธ์อย่างระมัดระวัง เช่น สังเกตเอกสารที่มีข้อความจำนวนมากปรากฏขึ้นในคราวเดียว หรือการใช้ถ้อยคำที่ไม่ถูกต้อง เกินจริง หรือไม่สอดคล้องกับบริบท เป็นต้น

(2) ด้านแบบจำลองและข้อมูล (Model & Data Risks)

ระบบ AI ที่ใช้ข้อมูลฝึกหรือโมเดลในการตัดสินใจ อาจเกิดความผิดพลาดจากหลายปัจจัย ได้แก่ bias จากข้อมูลฝึก, drift เมื่อสภาพตลาดเปลี่ยน, hallucination ที่ LLM สร้างข้อมูลเท็จ หรือ prompt injection จากคำสั่งแฝง ตลอดจนข้อจำกัดด้านคุณภาพ แหล่งที่มา สิทธิการใช้ ความครบถ้วน ความเป็นปัจจุบัน และความเหมาะสมของข้อมูลกับวัตถุประสงค์ของ use case ผลกระทบคือ คำแนะนำหรือผลลัพธ์ที่ผิดพลาดไม่เป็นธรรม หรือขาดความน่าเชื่อถือ ส่งผลเสียต่อผู้ลงทุนและชื่อเสียงองค์กร รวมถึงความเสี่ยงด้านกฎหมายและระเบียบปฏิบัติหรือข้อกำหนด อาจจำเป็นต้องมีแนวทางควบคุม เช่น การแลกเปลี่ยนข้อมูล บันทึก พัฒนา ตรวจสอบ ประเมิน และดำเนินการตามขั้นตอนของวงจรชีวิต AI เพื่อให้สามารถตรวจสอบย้อนกลับได้ทุกกรณี

(3) ด้านความมั่นคงปลอดภัยทางไซเบอร์และความยืดหยุ่นในการดำเนินงาน (Cybersecurity & Operational Resilience Risks)

ความเสี่ยงในมุมมองนี้กล่าวถึงภัยความเสี่ยงใหม่ เช่น supply-chain compromise ช่องโหว่ใน library/API การโจมตีแบบ data/model poisoning การดำเนินการของเอเจนต์อัตโนมัติโดยไม่มีการกำกับดูแลของมนุษย์ รวมถึงกรณีระบบหยุดทำงานเมื่อ AI provider ล่ม ทั้งหมดนี้ส่งผลกระทบต่อธุรกิจ การให้บริการลูกค้า ความปลอดภัยของข้อมูล และเพิ่มโอกาสเกิด systemic risk หากหลายระบบมีการหยุดชะงักของธุรกิจและบริการลูกค้าพร้อมกัน จึงจำเป็นต้องสร้างความคงทนของระบบ AI ให้พร้อมรับมือเหตุขัดข้องและการโจมตีลด downtime และป้องกันความเสียหายต่อธุรกิจ

(4) ด้านบุคคลที่ 3 และการกระจุกตัวของผู้ให้บริการ (Third-party & Concentration Risks)

เมื่อองค์กรพึ่งพาผู้ให้บริการโมเดลพื้นฐานหรือโครงสร้างพื้นฐานภายนอกมากเกินไป เช่น Foundation Models, Cloud, API หรือ Vendor รายใหญ่ อาจเกิด single point of failure หากผู้ให้บริการหยุดชะงัก เป็นต้น รวมถึงปัญหาการปฏิบัติตามกฎหมาย การละเมิดสัญญาใช้ข้อมูล หรือข้อพิพาทด้าน sovereignty ของข้อมูล ผลกระทบตามมาคือบริการหยุดชะงัก ส่งผลกระทบต่อชื่อเสียงองค์กร และเพิ่มความเสี่ยงเชิงระบบ จึงควรต้องมีการทำ Third-party due diligence จัดทำ exit plan และ monitoring vendor อย่างต่อเนื่อง

(5) ด้านจรรยาบรรณ ผลประโยชน์สูงสุดของลูกค้า และการเก็บหลักฐาน (Conduct, Client Best Interest & Record-Keeping Risks)

แม้จะใช้ AI สนับสนุนงานการให้คำแนะนำหรือการจัดพอร์ตลงทุน องค์กรต้องรักษาหลัก best interest ของลูกค้า ต้องมีวิธีสื่อสารอย่างโปร่งใสและอธิบายได้ มี audit trail ที่ตรวจสอบย้อนหลังได้ หากขาด oversight อาจเกิดคำแนะนำไม่เหมาะสม ไม่สามารถพิสูจน์ความถูกต้องได้ ส่งผลเสียต่อผู้ลงทุนและชื่อเสียงองค์กร จึงอาจต้องมีการกำหนด Minimum Baseline Controls ให้ทุกทีมต้องทำร่วมกัน เช่น รายงานบันทึกเหตุผลในการตัดสินใจ เป็นต้น

(6) ด้านข้อมูลส่วนบุคคล ทรัพย์สินทางปัญญา และกฎหมายระหว่างประเทศ (Privacy Data & IP Risks)

ข้อมูลลูกค้าหรือข้อมูลภายในอาจรั่วไหลจากการใช้งาน AI การใช้ข้อมูลฝึกหรือป้อนให้ LLM โดยไม่มีสิทธิหรือฐานกฎหมายที่ถูกต้อง ภาระตามกฎหมายเมื่อมีการโอนข้อมูลข้ามประเทศ ความเสี่ยงด้านลิขสิทธิ์ และทรัพย์สินทางปัญญาของเนื้อหาที่ AI สร้างขึ้น อาจทำให้เกิดการละเมิดกฎหมายคุ้มครองข้อมูลส่วนบุคคล (PDPA, GDPR) ค่าปรับและความเสียหายทางชื่อเสียง ข้อพิพาทด้านลิขสิทธิ์และความรับผิดชอบทางกฎหมาย จึงจำเป็นต้องมีมาตรการป้องกันการรั่วไหลของข้อมูล คุ้มครองสิทธิส่วนบุคคล และลดความเสี่ยงด้านกฎหมายหรือลิขสิทธิ์

อย่างไรก็ดี นอกจากการจำแนกความเสี่ยงหลัก 6 ประเภทข้างต้น ผู้ประกอบธุรกิจควรพิจารณา มุมมองประกอบการประเมินความเสี่ยง (cross-risk perspectives) ที่อาจเกิดขึ้นได้ในหลาย use case เช่น การประเมินความเสี่ยงแบบต่อเนื่องและเมื่อเกิดเหตุการณ์สำคัญ ความเสี่ยงเชิงระบบและผลกระทบข้ามตลาด การพึ่งพาผู้ให้บริการรายสำคัญ การเปิดเผยข้อมูลและการป้องกัน AI-washing รวมถึงข้อกำหนดด้านกฎหมายหรือมาตรฐานต่างประเทศ ทั้งนี้ มุมมองดังกล่าวมิได้เป็นประเภทความเสี่ยงใหม่แยกจาก 6 ประเภทหลัก แต่เป็นปัจจัยประกอบการพิจารณาความเข้มข้นของมาตรการควบคุม การติดตามผล และการทบทวน use case ตามระดับความเสี่ยง ดังนี้

(1) การทบทวนความเสี่ยงอย่างต่อเนื่องและเมื่อเกิดเหตุการณ์ที่มีนัยสำคัญ (Continuous and Event-driven Risk Assessment) การประเมินความเสี่ยงจาก AI ไม่ควรเป็นกิจกรรมครั้งเดียวก่อนนำระบบขึ้นใช้งาน แต่ควรครอบคลุมตั้งแต่การริเริ่มแนวคิด การทดสอบ การใช้งานจริง การปรับปรุงโมเดล การเปลี่ยนแปลงแหล่งข้อมูลหรือผู้ให้บริการ ไปจนถึงการเลิกใช้งานระบบ โดยควรมีทั้งการทบทวนตามรอบระยะเวลาและเมื่อเกิดเหตุการณ์หรือการเปลี่ยนแปลงที่มีนัยสำคัญ

(2) การพึ่งพาบุคคลที่ 3 และการกระจุกตัวของผู้ให้บริการ (Third-party Dependency and Concentration) ผู้ประกอบธุรกิจควรประเมินการพึ่งพาผู้ให้บริการภายนอก เช่น foundation model, cloud infrastructure, external API, data broker หรือบริการเฉพาะทางอื่น โดยเฉพาะกรณีที่ผู้ประกอบธุรกิจหลายรายพึ่งพาผู้ให้บริการรายเดียวกัน ซึ่งอาจก่อให้เกิด single point of failure หรือขยายผลกระทบจากระดับองค์กรไปสู่ระดับตลาดได้

(3) ผลกระทบเชิงระบบและต่อผู้ร่วมตลาด (Systemic and Cross-Market Impact) ผู้ประกอบธุรกิจควรพิจารณาว่า use case หรือ การพึ่งพาผู้ให้บริการบางประเภทอาจก่อให้เกิดผลกระทบต่อผู้ร่วมตลาดหลายราย สภาพคล่อง ความผันผวน ความต่อเนื่องของบริการ หรือความเชื่อมั่นในตลาดทุนหรือไม่ โดยเฉพาะกรณีที่มีการใช้โมเดล ข้อมูล หรือโครงสร้างพื้นฐานร่วมกันในวงกว้าง

(4) การเปิดเผยข้อมูล คำกล่าวอ้างเกี่ยวกับ AI และผลกระทบต่อลูกค้า (Disclosure, Conduct and AI-washing Risks) สำหรับ use case ที่เกี่ยวข้องกับลูกค้า ผู้ลงทุน การตลาด หรือการเปิดเผยข้อมูลต่อสาธารณะ ผู้ประกอบธุรกิจควรประเมินความเสี่ยงจากการสื่อสารที่อาจเกินจริง ไม่ถูกต้อง หรือทำให้เข้าใจผิดเกี่ยวกับการใช้งานหรือความสามารถของ AI รวมถึงควรมีหลักฐานรองรับคำกล่าวอ้างและกระบวนการทบทวนข้อความที่เกี่ยวข้อง

(5) กฎหมาย มาตรฐานและข้อกำหนดต่างประเทศที่เกี่ยวข้องกับการใช้งาน AI (Cross-border and Extraterritorial Compliance Risks) ในกรณีที่ use case เกี่ยวข้องกับลูกค้า ผู้มีส่วนได้เสีย หรือผู้ให้บริการในต่างประเทศ ผู้ประกอบธุรกิจควรติดตามข้อกำหนดที่อาจมีผลต่อการใช้งาน AI การโอนข้อมูลข้ามประเทศ การเปิดเผยต่อผู้ใช้ การจัดเก็บหลักฐาน สัญญากับผู้ให้บริการ และการกำกับดูแลตามมาตรฐานหรือกฎหมายต่างประเทศที่เกี่ยวข้อง



ความเสี่ยงจากการประยุกต์ใช้ AI ในบริบทของภาคตลาดทุน และมุมมองประกอบการประเมินความเสี่ยง (cross-risk perspectives)

Market Integrity & Market Abuse Risks

ระบบ AI สร้างรูปแบบคำสั่งซื้อขายซึ่งคล้ายพฤติกรรมต้องห้าม เช่น Spoofing/ Layering หรือสร้างการเชื้ชวน/โฆษณาที่ทำให้ผู้ลงทุนเข้าใจผิด (AI-washing) หรือการที่ Model สำหรับการเทรดอัตโนมัติเกิดการเกยพ้องกัน ทำให้ตลาดเสียสมดุล (Herding Behavior)

Model & Data Risks

Bias จาก Training Data, AI Model Drift, Hallucination, การสร้างข้อมูลเท็จ, ถูกโจมตีด้วย Prompt Injection ส่งผลนำไปสู่ผลลัพธ์ที่ผิดพลาด ไม่เป็นธรรม ขาดความน่าเชื่อถือ

Cybersecurity & Operational Resilience Risks

เช่น Supply-chain Compromise, Library/ API Vulnerability, Data/Model Poisoning ซึ่งส่งผลต่อธุรกิจได้ จึงจำเป็นต้องสร้างความทนทานของระบบ AI และลด Downtime



Third-party & Concentration Risks

ความเสี่ยง Single Point of Failure เมื่อพึ่งพาผู้ให้บริการ AI Model / Infrastructure Provider มากเกินไป รวมถึงปัญหาด้านกฎหมาย การละเมิดสัญญาใช้ข้อมูล หรือ Data Sovereignty

Conduct, Client Best Interest & Record-Keeping Risks

AI ให้คำแนะนำที่ไม่เหมาะสม ผู้ประกอบธุรกิจไม่สามารถพิสูจน์ความถูกต้องได้ ส่งผลเสียต่อผู้ลงทุนและชื่อเสียงองค์กร เช่น การขาด Audit Trail, การไม่สื่อสารถึงที่มาของข้อมูลต่อลูกค้า

Privacy Data, IP, Regulatory Risks

- ข้อมูลลูกค้าหรือข้อมูลภายในรั่วไหลจากการใช้งาน AI
- การใช้ข้อมูลฝึกหรือป้อนให้ LLM โดยไม่มีสิทธิ์หรือฐานกฎหมายที่ถูกต้อง
- การละเมิดกฎหมายเมื่อมีการโอนข้อมูลข้ามประเทศ
- ความเสี่ยงด้านลิขสิทธิ์และทรัพย์สินทางปัญญาของเนื้อหาที่ AI สร้าง

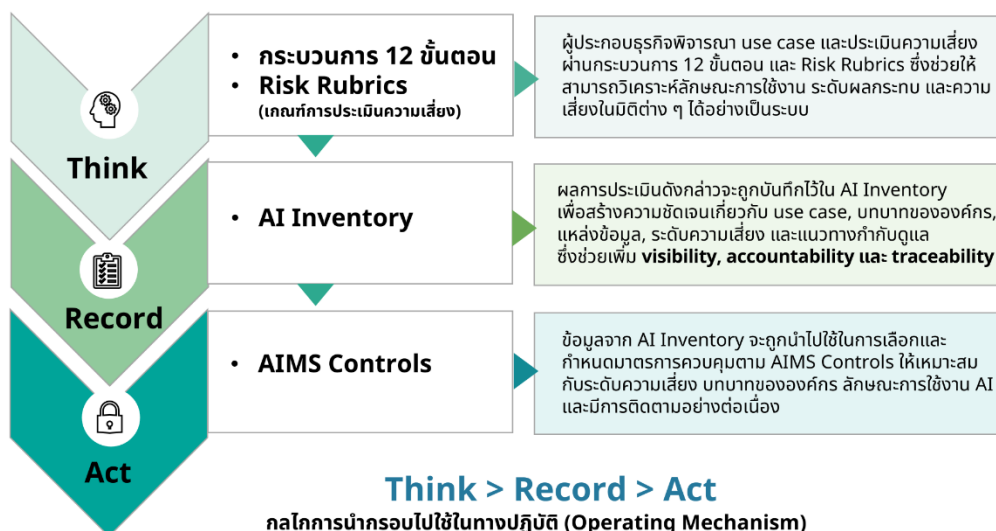
ปัจจัยประกอบการพิจารณาความเข้มของมาตรการควบคุม การติดตามผล และการทบทวน use case ตามระดับความเสี่ยง

ภาพที่ 7 ความเสี่ยงจากการประยุกต์ใช้ AI ในบริบทของภาคตลาดทุน

ภายหลังจากที่ผู้ประกอบการได้ประเมินความเสี่ยงจากการใช้งาน AI โดยพิจารณาทั้งในมิติของ ความเสี่ยงเชิงคงที่ (static risks) และเชิงพลวัต (dynamic risks) รวมถึงความเสี่ยงในบริบทของภาคตลาดทุน ขั้นตอนถัดไปคือการนำผลการประเมินดังกล่าวไปใช้กำหนดแนวทางกำกับดูแล ควบคุม และติดตามอย่างเป็น ระบบ ซึ่งกรอบการกำกับดูแลฉบับนี้ได้ประยุกต์กระบวนการดังกล่าวให้สามารถดำเนินการได้จริง ผ่านกลไก การนำกรอบการกำกับดูแลไปใช้งานจริง (Think > Record > Act) แนวทางดังกล่าวช่วยให้เกิดความเชื่อมโยง ระหว่างการประเมินความเสี่ยงกับการกำหนดมาตรการควบคุมอย่างต่อเนื่อง และสะท้อนหลัก Proportionality กล่าวคือ AI use case ที่มีความเสี่ยงสูงจะได้รับการประเมินและควบคุมอย่างเข้มข้นมากขึ้น ในขณะที่ AI use case ที่มีความเสี่ยงต่ำสามารถใช้แนวทางแบบยืดหยุ่นได้ตามความเหมาะสม

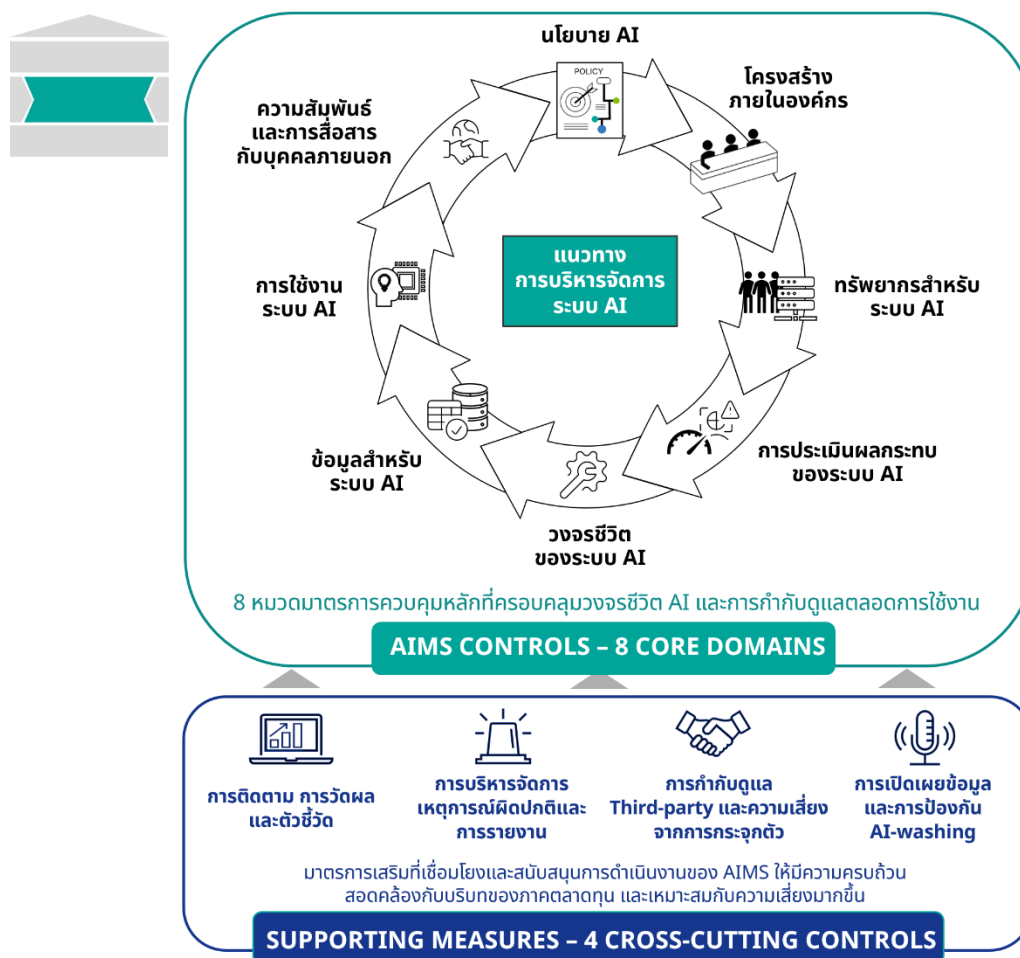
ตารางที่ 1 กลไกการนำกรอบการกำกับดูแลไปใช้งานจริง (Operating Mechanism)

Operating Mechanism	เครื่องมือ	ผลลัพธ์ที่คาดหวัง
Think	กระบวนการ 12 ขั้นตอน + Risk Rubrics (เกณฑ์การประเมินความเสี่ยง)	ผู้ประกอบการพิจารณา use case และประเมินความเสี่ยง ผ่านกระบวนการ 12 ขั้นตอน และ Risk Rubrics ซึ่งช่วยให้สามารถวิเคราะห์ลักษณะการใช้งาน ระดับผลกระทบ และความเสี่ยงในมิติต่าง ๆ ได้อย่างเป็นระบบ
Record	AI Inventory	ผลการประเมินดังกล่าวจะถูกบันทึกไว้ใน AI Inventory เพื่อสร้างความชัดเจนเกี่ยวกับ use case, บทบาทขององค์กร, แหล่งข้อมูล, ระดับความเสี่ยง และแนวทางกำกับดูแล ซึ่งช่วยเพิ่ม visibility, accountability และ traceability
Act	AIMS Controls	ข้อมูลจาก AI Inventory จะถูกนำไปใช้ในการเลือกและกำหนด มาตรการควบคุมตาม AIMS Controls ให้เหมาะสมกับระดับ ความเสี่ยง บทบาทขององค์กร และลักษณะการใช้งาน AI



ภาพที่ 8 กลไกการนำกรอบไปใช้ในทางปฏิบัติ (Operating Mechanism)

2.4 แนวทางการบริหารจัดการระบบ AI (AIMS Controls)



ภาพที่ 9 AIMS Controls

2.4.1 นโยบาย AI องค์กรต้องมีนโยบาย AI ที่เป็นลายลักษณ์อักษร ชัดเจน สอดคล้องกับนโยบายองค์กร และมาตรฐาน และควรทบทวนปรับปรุงอย่างสม่ำเสมอ เพื่อกำหนดทิศทางการใช้ AI อย่างรับผิดชอบและต้องมีความปลอดภัย ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) นโยบายต้องกำหนดให้ทุกระบบ AI เคารพความเป็นธรรม ไม่เลือกปฏิบัติ และคำนึงถึงจริยธรรมตลอดวงจรชีวิต
- (2) ระบุโครงสร้างการรับผิดชอบเชิงนโยบายและช่องทางตรวจสอบได้ชัดเจน
- (3) ฝังมาตรการรักษาความมั่นคงและปลอดภัยตั้งแต่ต้นน้ำ

2.4.2 โครงสร้างภายในองค์กร โดยกำหนดบทบาท ความรับผิดชอบ และช่องทางการรายงานปัญหาที่เกิดจากการใช้งาน AI อย่างชัดเจน เพื่อสร้างความโปร่งใสและการกำกับดูแลภายใน ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) แยกบทบาทระหว่างผู้พัฒนา ผู้อนุมัติ และผู้ตรวจสอบอิสระ
- (2) กำหนดบทบาท Human-in-the-loop/on-the-loop ในกระบวนการสำคัญ
- (3) มีช่องทางรายงานเหตุการณ์หรือปัญหาอย่างรวดเร็ว

2.4.3 ทรัพยากรสำหรับระบบ AI องค์กรต้องจัดหาและบันทึกทรัพยากรทรัพยากรที่เหมาะสม เช่น ข้อมูล เครื่องมือ ฮาร์ดแวร์ ซอฟต์แวร์ และบุคลากรที่มีทักษะ เพื่อให้ระบบ AI ทำงานได้อย่างมีประสิทธิภาพและปลอดภัย ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) จัดทำรายการทรัพยากรทั้งหมด พร้อมมาตรฐานสิทธิ์เข้าถึงขั้นต่ำ
- (2) กำหนดทักษะขั้นต่ำของบุคลากรและจัดอบรมต่อเนื่อง
- (3) บังคับใช้มาตรฐานความปลอดภัยในทุกขั้นตอน

2.4.4 การประเมินผลกระทบของระบบ AI ควรมีขั้นตอนในการประเมินผลกระทบของ AI ต่อบุคคล กลุ่มบุคคล และสังคมอย่างเป็นระบบ พร้อมบันทึกผลลัพธ์การใช้งานประกอบการตัดสินใจ ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) จัดทำแบบฟอร์มประเมินทั้งด้านเทคนิค จริยธรรม ผลกระทบต่อบุคคล กลุ่มบุคคล สังคม รวมถึง misuse/harm hypothesis ทุกครั้งก่อนใช้งานจริง
- (2) กำหนดเกณฑ์อนุมัติก่อนออกใช้งาน พร้อมบันทึก (memo) ผลการประเมินความเสี่ยงและความเสี่ยงคงเหลือ (residual risk) ทุกครั้งก่อนนำออกใช้งานหรือทดลองใช้งาน (Deploy/Pilot) จริง

2.4.5 วงจรชีวิตของระบบ AI ต้องควบคุมระบบ AI ตลอดวงจรชีวิต ตั้งแต่ออกแบบ พัฒนา ใช้งาน จนเลิกใช้ โดยกำหนดข้อกำหนด ทำการตรวจสอบ บันทึกและเก็บ log เหตุการณ์ ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) กำหนดขั้นตอนในมาตรฐานกระบวนการพัฒนา AI ที่ครอบคลุมทุกระยะ ตั้งแต่ ออกแบบ (Design) ฝึกสอนโมเดล (Train) ตรวจสอบความถูกต้อง (Validate) เริ่มใช้งาน (Release) ติดตามผล (Monitor) และยุติการใช้งาน (Retire)
- (2) บันทึกเวอร์ชันโมเดลและเกณฑ์การเริ่มใช้งาน และวางแผนการย้อนกลับ (Rollback)
- (3) มีแผนการติดตามหลังการใช้งานจริง (Post-market Monitoring Plan)

2.4.6 ข้อมูลสำหรับระบบ AI ต้องจัดการข้อมูลที่ใช้ในการพัฒนาและปรับปรุง AI ให้มีคุณภาพ ตรวจสอบแหล่งที่มา ปรับปรุงให้เหมาะสม และลดอคติ (bias) ที่อาจเกิดขึ้น ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) ตรวจสอบแหล่งที่มา ความถูกต้อง ความครบถ้วน ความเป็นส่วนตัว และการให้ความยินยอม ด้านทรัพย์สินทางปัญญา (IP Consent) สำหรับข้อมูลทุกชุดก่อนนำไปใช้งานจริง
- (2) ทำการปกปิดข้อมูล (Data Masking) การทำให้ไม่สามารถระบุตัวบุคคลได้ (Anonymization/De-identification) สำหรับข้อมูลส่วนบุคคล (PII) หรือข้อมูลลับในทุกขั้นตอนของการใช้งาน การฝึกสอน แบบจำลอง และการจัดการข้อมูลนำเข้า/ผลลัพธ์ และควรใช้ข้อมูลเท่าที่จำเป็น
- (3) เชื่อมโยงข้อมูลที่ใช้กับระบบ AI กับกลไกธรรมาภิบาลข้อมูลที่ต้องคำนึงอยู่ เช่น ทะเบียนข้อมูล หรือบัญชีข้อมูล เมทาดาตา แผนภาพการไหลของข้อมูล การกำหนดเจ้าของข้อมูล ระดับชั้นข้อมูล และบันทึกกิจกรรมการประมวลผลข้อมูล ตามความเหมาะสม เพื่อให้สามารถตรวจสอบแหล่งที่มา การใช้ และข้อจำกัดของข้อมูลได้

(4) กำหนดแนวทางการเก็บรักษา การเข้าถึง การใช้ การแบ่งปัน การลบ หรือการทำลายข้อมูลที่เกี่ยวข้องกับ AI use case ให้สอดคล้องกับวัตถุประสงค์ ระดับชั้นข้อมูล กฎหมาย นโยบายข้อมูล และระดับความเสี่ยงของ AI use case

2.4.7 การใช้งานระบบ AI ต้องกำหนดกระบวนการและเงื่อนไขการใช้อย่างชัดเจน ตรวจสอบให้มั่นใจว่าใช้งานตามวัตถุประสงค์ที่ตั้งใจ ไม่เบี่ยงเบนไปในทางที่อาจสร้างความเสี่ยง ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

(1) กำหนดขอบเขต เงื่อนไข วิธีการใช้งานให้ชัดเจน มีขั้นตอนปฏิบัติงาน (Standard Operating Procedure: SOP) สำหรับการแจ้งเหตุผิดปกติ การใช้งานผิดวัตถุประสงค์ และช่องทางการแจ้งเหตุโดยเฉพาะสำหรับกรณี GenAI หรือการใช้งานที่อ่อนไหว

(2) ออกแบบระบบป้องกัน (guardrails) เช่น safety filter การบังคับใช้นโยบาย และกำหนดให้มีบุคคลควบคุมการตัดสินใจ (Human-in-the-loop) สำหรับกรณีที่มีความเสี่ยงสูง

(3) ทดสอบระบบด้วยทีมจำลองการโจมตี (Red Team) และการทดสอบเชิงปรับแต่ง (Adversarial Testing) ก่อนและหลังการนำระบบ AI ออกใช้งานเพื่อตรวจสอบหาช่องโหว่

นอกจากนี้ ผู้ประกอบธุรกิจอาจคำนึงถึงการบริหารจัดการเหตุการณ์ผิดปกติที่เกิดจากการใช้งาน AI (AI Incident Management) มีวัตถุประสงค์เพื่อเสริมสร้างความพร้อมปรับตัว (Resilience) ให้องค์กรสามารถรับมือกับเหตุการณ์ผิดปกติของระบบ AI โดยไม่กระทบต่อความต่อเนื่องของบริการและความเชื่อมั่นของผู้ลงทุน โดยสามารถประเมินและคัดเลือกผสมแนวทางเดิมด้าน ERM และ Cybersecurity เข้ากับบริบทของ AI อย่างเป็นระบบ มีวัตถุประสงค์ให้ครอบคลุมอย่างเป็นขั้นตอน ตั้งแต่การกำหนดตัวชี้วัดเพื่อเฝ้าระวังและตรวจจับความผิดปกติ การวิเคราะห์สาเหตุและประเมินผลกระทบต่อนักลงทุนเพื่อจัดลำดับความเร่งด่วน การจำกัดวงความเสียหายผ่านการหยุดหรือสลับระบบ การแก้ไขต้นเหตุและฟื้นฟูระบบอย่างระมัดระวัง จนถึงการถอดบทเรียนและปรับปรุงมาตรการควบคุมอย่างต่อเนื่อง ทั้งหมดนี้เป็นตัวอย่างกระบวนการเชิงวงจรที่ช่วยยกระดับการกำกับดูแล AI ให้มีประสิทธิภาพและยั่งยืนในระยะยาว

2.4.8 ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก ควรจัดการความรับผิดชอบและความเสี่ยงกับผู้ขาย พันธมิตร และลูกค้าอย่างชัดเจน ตั้งเกณฑ์คัดเลือกและติดตามผู้ให้บริการ AI รวมถึงจัดเตรียมและสื่อสารข้อมูลที่จำเป็นแก่ผู้ใช้ ลูกค้า หรือคู่ค้าอย่างโปร่งใส เช่น ประโยชน์ ความเสี่ยง วิธีใช้ หรือเหตุการณ์ผิดปกติ เพื่อสร้างความเข้าใจและความเชื่อมั่น คำนึงถึงหลัก Shared Responsibility โดยองค์กรยังคงต้องรับผิดชอบต่อผลลัพธ์ (Accountability) ควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

(1) กำหนดเกณฑ์การคัดเลือกผู้ให้บริการ ตรวจสอบความน่าเชื่อถือ ทบทวนข้อตกลงระดับบริการ (SLA) และตัวชี้วัด (KPI) (ถ้ามี) อย่างสม่ำเสมอ

(2) เปิดเผยข้อมูลสำคัญแก่ผู้ให้บริการ/ลูกค้า เช่น วิธีใช้ ข้อจำกัด และการระบุผลลัพธ์อย่างชัดเจน การแสดงสัญลักษณ์หรือข้อความกำกับ (label) ที่ผู้ใช้งานเข้าใจได้ง่ายว่าเนื้อหาถูกสร้างโดย AI หรือแจ้งให้เจ้าของข้อมูลทราบหากมีการนำข้อมูลอ่อนไหวไปฝึกสอนหรือใช้ในการประมวลผลด้วยโมเดล

(3) มีช่องทางรับข้อเสนอแนะ ร้องเรียน หรือแจ้งเหตุผิดปกติจากลูกค้า คู่ค้า หรือผู้ให้บริการ และมีคู่มือการรับมือเหตุการณ์ที่เกี่ยวข้องกับความเสียหายจากบุคคลภายนอก

นอกจากนี้ ยังมีมาตรการเสริมเพิ่มเติมเพื่อให้สามารถบริหารจัดการระบบ AI ได้อย่างมีประสิทธิภาพ และเสริมความพร้อมด้านการติดตามและการวัดผล การรายงานเหตุการณ์ การกำกับดูแลบุคคลที่ 3 หรือผู้ให้บริการภายนอก และความเสี่ยงจากการกระจุกตัว รวมถึงการเปิดเผยข้อมูลและการป้องกันการสื่อสารเกี่ยวกับ AI ที่อาจทำให้เข้าใจผิด ทั้งนี้ การนำมาตรการดังกล่าวไปใช้ควรพิจารณาตามลักษณะ ความซับซ้อน และระดับความเสี่ยงของแต่ละ use case โดยมีวัตถุประสงค์เพื่อเป็นแนวทางเสริมจาก AIMS Controls ที่ช่วยให้การเลือกและดำเนินมาตรการในชั้น Act มีความครบถ้วน สอดคล้องกับบริบทของภาคตลาดทุน และเหมาะสมกับความเสี่ยงมากขึ้น ซึ่งจะกล่าวถึงในหัวข้อ 2.4.9 – 2.4.12 ดังนี้

2.4.9 การติดตาม การวัดผล และตัวชี้วัดสำหรับระบบ AI ผู้ประกอบธุรกิจควรจัดให้มีการติดตามผลการทำงานของระบบ AI อย่างต่อเนื่องโดยอาศัยตัวชี้วัดที่เหมาะสมกับลักษณะการใช้งาน ระดับความเสี่ยง และผลกระทบที่อาจเกิดขึ้น ทั้งในเชิงประสิทธิภาพ เชิงเทคนิค ความถูกต้องของผลลัพธ์ ความสม่ำเสมอ ความเป็นธรรม จำนวนเหตุการณ์ผิดปกติ การพึ่งพาบุคคลที่ 3 และระดับการแทรกแซงโดยมนุษย์ เพื่อให้สามารถระบุความเสี่ยงที่เพิ่มขึ้น ความผิดปกติ หรือสัญญาณเตือนล่วงหน้าได้ทันทั่วทั้ง โดยควรมีมาตรการควบคุม หรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) กำหนดตัวชี้วัดขั้นต่ำสำหรับแต่ละ use case เช่น accuracy, error rate, drift, override rate, incident rate, response time หรือ KPI อื่นที่เหมาะสมกับบริบท
- (2) สำหรับ use case ที่มีผลต่อผู้ลงทุน การซื้อขาย หรือความต่อเนื่องของบริการ ควรกำหนด threshold สำหรับการยกระดับ การหยุดใช้งานชั่วคราว หรือการทบทวนระบบทันที
- (3) จัดให้มี dashboard, log, รายงานสรุป หรือเอกสารเทียบเท่าที่สามารถตรวจสอบย้อนหลัง และนำเสนอต่อผู้บริหาร หรือหน่วยงานกำกับดูแลได้ตามความเหมาะสม

2.4.10 การบริหารจัดการเหตุการณ์ผิดปกติ และการรายงานต่อสำนักงาน ก.ล.ต. ผู้ประกอบธุรกิจควรจัดให้มีขั้นตอนการจำแนกเหตุการณ์ผิดปกติที่เกี่ยวข้องกับ AI และเกณฑ์การพิจารณาว่า เหตุการณ์ใดมีนัยสำคัญเพียงพอที่จะต้องรายงานต่อสำนักงาน ก.ล.ต. โดยคำนึงถึงผลกระทบต่อผู้ลงทุน ความถูกต้องของข้อมูล ความต่อเนื่องของบริการ การละเมิดข้อมูลส่วนบุคคล ความมั่นคงปลอดภัยไซเบอร์ และความเชื่อมั่นต่อการดำเนินธุรกิจหรือตลาดทุน โดยควรมีมาตรการควบคุม หรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) กำหนดนิยามของ AI incident, near miss และ significant AI incident พร้อมเกณฑ์การยกระดับภายในองค์กร
- (2) จัดเตรียมแบบฟอร์มรายงานเหตุการณ์ที่ครอบคลุมข้อเท็จจริงเบื้องต้น สาเหตุที่คาดหมาย ผลกระทบต่อผู้ลงทุนหรือระบบ มาตรการจำกัดวงความเสียหาย และแผนฟื้นฟู เพื่อให้สามารถรายงานต่อหน่วยงานกำกับดูแลได้โดยเร็วและเป็นมาตรฐาน
- (3) จัดให้มีการสรุปบทเรียนภายหลังเหตุการณ์ (post-incident review) และทบทวนมาตรการควบคุมเดิมอย่างน้อยทุกครั้งที่เกิดเหตุการณ์มีนัยสำคัญ

2.4.11 การกำกับดูแลบุคคลที่ 3 และความเสี่ยงจากการกระจุกตัว ผู้ประกอบธุรกิจควรมีกรอบการพิจารณา คัดเลือก ทำสัญญา ติดตาม และประเมินผู้ให้บริการ AI หรือผู้ให้บริการที่เกี่ยวข้องอย่างต่อเนื่อง โดยไม่พิจารณาเพียงคุณสมบัติทางเทคนิค หรือราคา แต่ต้องครอบคลุมถึงความโปร่งใสของผู้ให้บริการ ความมั่นคงปลอดภัย การคุ้มครองข้อมูลส่วนบุคคล สิทธิในทรัพย์สินทางปัญญา ความสามารถในการรองรับการตรวจสอบ และความพร้อมในการแจ้งเหตุการณ์หรือการเปลี่ยนแปลงที่มีนัยสำคัญให้แก่องค์กร โดยควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

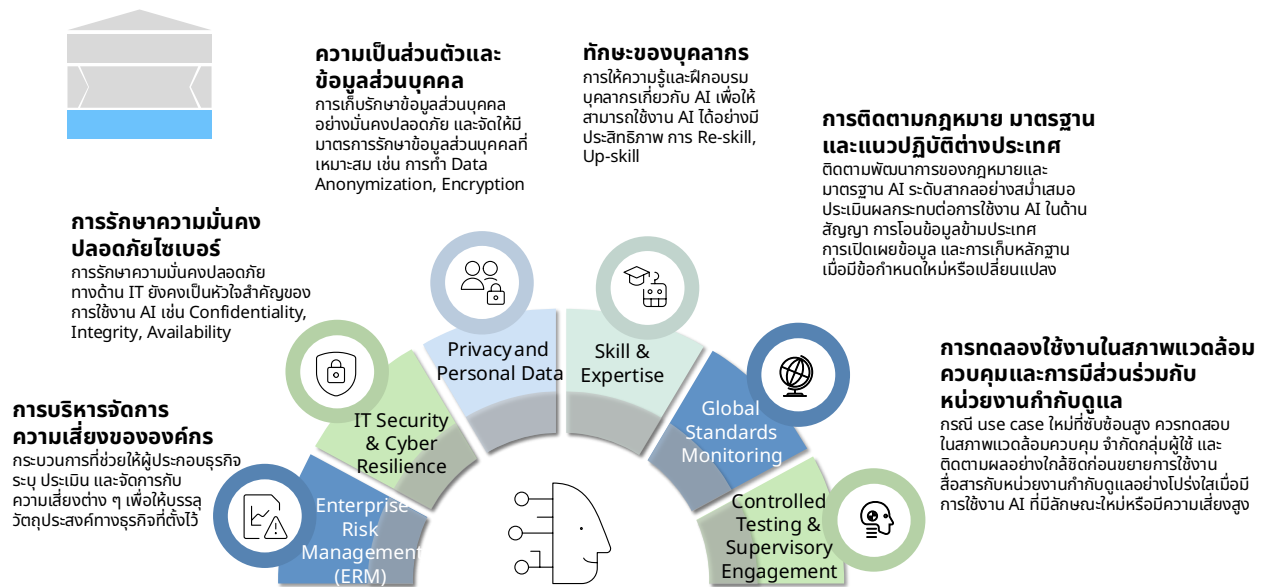
- (1) ดำเนินการ due diligence ของผู้ให้บริการที่ครอบคลุมอย่างน้อยด้าน model/data lineage, cyber security, subcontracting, incident notification, audit rights, portability และ exit plan
- (2) ประเมินความเสี่ยงจากการกระจุกตัว (concentration risk) ทั้งในระดับ use case และระดับองค์กร เช่น การใช้ผู้ให้บริการคลาวด์หรือ foundation model รายเดียวในหลายกระบวนการสำคัญ
- (3) จัดให้มีแผนสำรอง การย้ายระบบ หรือแนวทางลดการพึ่งพารายเดียว สำหรับ use case ที่มีความสำคัญสูง หรือกระทบต่อผู้ใช้บริการจำนวนมาก

2.4.12 การเปิดเผยข้อมูลต่อผู้ใช้บริการและการป้องกันการใช้งาน AI อย่างเกินจริง หรือทำให้เข้าใจผิด ผู้ประกอบธุรกิจควรเปิดเผยข้อมูลเกี่ยวกับการใช้งาน AI ต่อผู้ใช้บริการ ลูกค้า หรือผู้มีส่วนได้เสียอย่างเหมาะสมกับบริบทและความเสี่ยง โดยเฉพาะเมื่อ AI มีบทบาทในการให้คำแนะนำ การให้บริการอัตโนมัติ การประมวลผลข้อมูลส่วนบุคคล หรือการตัดสินใจที่มีผลกระทบต่อสิทธิ หรือผลประโยชน์ของผู้ใช้บริการ ทั้งนี้ การเปิดเผยควรมีความถูกต้อง ไม่เกินจริง ไม่ชวนให้เข้าใจผิด และไม่ปกปิดข้อจำกัดที่มีนัยสำคัญของระบบ โดยควรมีมาตรการควบคุมหรือเอกสารหลักฐานขั้นต่ำ ดังนี้

- (1) กำหนดหลักเกณฑ์ในการทบทวนข้อความประชาสัมพันธ์ ข้อความทางการตลาด รายงาน นักลงทุน หรือข้อความบนระบบดิจิทัลที่อ้างถึง AI ก่อนเผยแพร่ทุกครั้ง
- (2) เปิดเผยอย่างเหมาะสมว่าในบริการหรือกระบวนการใดมีการใช้ AI และบทบาทของ AI คืออะไร ข้อมูลประเภทใดที่ระบบอาจใช้ และผู้ใช้บริการมีช่องทางขอคำอธิบาย หรือขอให้บุคคลเข้ามาทบทวน ได้หรือไม่
- (3) จัดให้มีการเก็บหลักฐานรองรับการกล่าวอ้างเกี่ยวกับความสามารถ หรือประสิทธิภาพของระบบ AI เพื่อป้องกันการสื่อสารเกินจริง หรือ AI-washing

ดังนั้น ในการดำเนินการตามกรอบการกำกับดูแลการใช้งาน AI ให้เกิดผลได้จริง ผู้ประกอบธุรกิจจำเป็นต้องเข้าใจถึงวัตถุประสงค์และขอบเขต เหตุผลและความจำเป็น รวมถึงผลกระทบของการนำ AI มาใช้งาน และดำเนินการรวบรวมและบันทึกการใช้งานให้สามารถประเมินถึงความเสี่ยงและบริหารจัดการความเสี่ยง และเลือกใช้มาตรการควบคุมการใช้งาน AI ได้อย่างเหมาะสมและเป็นขั้นตอน ซึ่งหากองค์กรสามารถแบ่งแยกและระบุกลุ่มประเภทของการพัฒนาระบบงาน AI ได้อย่างชัดเจน ตัวอย่างเช่น ระบบ AI ที่องค์กรพัฒนาเอง (developer) หรือระบบ AI ที่องค์กรนำมาใช้ (user) หรือเป็นระบบงาน AI แบบผสมผสานที่องค์กรใช้ software AI แต่มีการ customized เพื่อให้เหมาะสมกับบริบทหรือหน่วยงานภายในองค์กร (customizer) เป็นต้น ก็จะสามารถกำหนดมาตรการควบคุมได้อย่างเหมาะสมกับระดับความเสี่ยง (Proportionality) (รายละเอียดในภาคผนวกหมายเลข 5)

2.5 ปัจจัยส่งเสริมให้บรรลุวัตถุประสงค์ (Enablers)



ภาพที่ 10 ปัจจัยส่งเสริมให้บรรลุวัตถุประสงค์ (Enablers)

การประยุกต์ใช้ AI ให้บรรลุวัตถุประสงค์ที่กำหนดไว้ ต้องคำนึงถึงปัจจัยส่งเสริมและสนับสนุนการนำ AI มาใช้ในองค์กรให้บรรลุวัตถุประสงค์ ดังนี้

2.5.1 การบริหารจัดการความเสี่ยงขององค์กร

การบริหารจัดการความเสี่ยงขององค์กร (ERM) เป็นกระบวนการที่ช่วยให้ผู้ประกอบธุรกิจระบุ ประเมิน และจัดการกับความเสี่ยงต่าง ๆ เพื่อให้บรรลุวัตถุประสงค์ทางธุรกิจที่ตั้งไว้ และสามารถช่วยเพิ่มโอกาสในการบรรลุวัตถุประสงค์ของการใช้งาน AI ได้ ดังนี้

(1) สร้างความเข้าใจความเสี่ยงที่เกี่ยวข้องกับ AI: ERM ช่วยให้ผู้ประกอบธุรกิจเข้าใจความเสี่ยงที่เกี่ยวข้องกับ AI ได้อย่างครอบคลุม โดยพิจารณาปัจจัยต่าง ๆ เช่น ประเภทของความเสี่ยง ระดับความรุนแรงของความเสี่ยง และผลกระทบต่อองค์กร ความเข้าใจความเสี่ยงเหล่านี้ จะช่วยให้องค์กรสามารถวางแผนและรับมือกับความเสี่ยงได้อย่างมีประสิทธิภาพ

(2) กำหนดมาตรการควบคุมความเสี่ยง: ERM ช่วยให้ผู้ประกอบธุรกิจกำหนดมาตรการควบคุมความเสี่ยงที่เหมาะสมกับความเสี่ยงที่ระบุไว้ มาตรการควบคุมความเสี่ยงเหล่านี้อาจรวมถึงการพัฒนานโยบายและแนวปฏิบัติ การฝึกอบรมพนักงาน และการลงทุนในเทคโนโลยี ความปลอดภัย

(3) ติดตามความเสี่ยง: ERM ช่วยให้ผู้ประกอบธุรกิจติดตามความเสี่ยงอย่างต่อเนื่อง เพื่อให้มั่นใจว่าความเสี่ยงเหล่านี้ยังคงอยู่ในระดับที่ยอมรับได้ การประเมินความเสี่ยงอย่างสม่ำเสมอ จะช่วยให้องค์กรสามารถระบุความเสี่ยงใหม่ ๆ และปรับปรุงมาตรการควบคุมความเสี่ยงที่มีอยู่ได้อย่างมีประสิทธิภาพ

2.5.2 การรักษาความมั่นคงปลอดภัยไซเบอร์

การรักษาความมั่นคงปลอดภัยทางเทคโนโลยีสารสนเทศยังคงเป็นหัวใจสำคัญของการใช้งาน AI เช่นเดียวกับการใช้งานเทคโนโลยีสารสนเทศอื่น ๆ โดยการออกแบบ พัฒนา และใช้งาน AI ควรคำนึงถึงความมั่นคงปลอดภัยทางเทคโนโลยีสารสนเทศ เพื่อป้องกันการถูกโจมตีโดยผู้ไม่ประสงค์ดี และการใช้งาน AI ที่ไม่เหมาะสม ซึ่งอาจส่งผลกระทบต่อลูกค้าและบุคคลที่เกี่ยวข้อง

ผู้ประกอบการควรจัดให้มีนโยบายและมาตรการควบคุมที่เหมาะสมกับความเสี่ยง โดยคำนึงถึงองค์ประกอบหลักของการรักษาความมั่นคงปลอดภัยทางเทคโนโลยีสารสนเทศ 3 ประการ ดังนี้

(1) การรักษาความลับ (Confidentiality): รักษาความลับของข้อมูลในทุกขั้นตอน ตั้งแต่การรวบรวมข้อมูล วิเคราะห์ข้อมูล และประมวลผลข้อมูล โดยให้สิทธิ์เข้าถึงข้อมูล เฉพาะแก่ผู้ที่มีสิทธิ์ และป้องกันการเข้าถึงข้อมูลโดยผู้ที่ไม่ได้สิทธิ์หรือผู้ที่ไม่เกี่ยวข้อง

(2) การรักษาความถูกต้องครบถ้วน (Integrity): มีมาตรการเพื่อป้องกันการแก้ไขหรือเปลี่ยนแปลงข้อมูลโดยไม่ได้รับอนุญาต เพื่อให้ข้อมูลมีความถูกต้องครบถ้วนและเชื่อถือได้อยู่เสมอ

(3) การรักษาสภาพความพร้อมใช้งาน (Availability): บริการของ AI ต้องสามารถเข้าถึงหรือใช้งานได้ เมื่อมีความจำเป็นต้องใช้ พร้อมทั้งสามารถใช้งานได้อย่างต่อเนื่อง (Continuity) โดยเฉพาะการนำ AI มาใช้งานกับกระบวนการทางธุรกิจซึ่งมีความสำคัญสูง

ทั้งนี้ เพื่อให้มั่นใจว่าการใช้งาน AI เป็นไปตามนโยบายด้าน IT Security ขององค์กร ผู้ประกอบการควรจัดให้มีการตรวจสอบด้าน IT (IT Audit) ที่ครอบคลุมระบบ AI โดยจัดให้มีการปรับปรุงข้อตรวจพบต่าง ๆ ภายในระยะเวลาที่เหมาะสมกับความเสี่ยง

2.5.3 ธรรมาภิบาลข้อมูลและความเป็นส่วนตัว

การกำกับดูแลข้อมูลสำหรับ AI ควรครอบคลุมอย่างน้อย 5 มิติ ได้แก่ คุณภาพข้อมูล การปกป้องข้อมูล ความเสี่ยงจากข้อมูลสูญหายหรือไม่พร้อมใช้ การปฏิบัติตามข้อกำหนดด้านข้อมูล และความเสี่ยงจากการเปิดเผยหรือรั่วไหลของข้อมูล โดยควรพิจารณาตลอดเส้นทางการใช้ข้อมูลของ AI use case ตั้งแต่ข้อมูลนำเข้าคำสั่งหรือข้อความที่ป้อนเข้าสู่ระบบ ฐานความรู้ที่ระบบดึงมาใช้ ผลลัพธ์ที่ระบบสร้างขึ้น บันทึกการใช้งาน และการเชื่อมต่อกับระบบหรือผู้ให้บริการภายนอกตามความเกี่ยวข้องของแต่ละ use case โดยผู้ประกอบการสามารถใช้กลไก Data Governance ที่มีอยู่เป็นฐานในการกำกับดูแลข้อมูลที่เกี่ยวข้อง เช่น การกำหนดโครงสร้างผู้รับผิดชอบข้อมูล นโยบายข้อมูล ระดับชั้นของข้อมูล บัญชีรายการข้อมูล เมทาเดตา อภิธานศัพท์ทางธุรกิจ การควบคุมคุณภาพข้อมูล และกระบวนการบริหารข้อมูลตามวงจรชีวิตข้อมูล มาใช้ประกอบการประเมิน AI use case เพื่อให้ทราบว่าข้อมูลที่นำเข้า ใช้ประมวลผล สร้างผลลัพธ์ บันทึกเป็น log หรือถูกแลกเปลี่ยนกับระบบหรือผู้ให้บริการภายนอก มีผู้รับผิดชอบ ข้อจำกัด และมาตรการควบคุมที่เหมาะสมหรือไม่ นอกจากนี้ ยังควรพิจารณาเพิ่มเติมว่าข้อมูลที่ใช้เหมาะสมกับวัตถุประสงค์ของ AI use case หรือไม่ ใช้ตามสิทธิและข้อจำกัดที่เกี่ยวข้องหรือไม่ มีคุณภาพเพียงพอหรือไม่ มีความเสี่ยงจาก bias, drift, data exposure หรือ data loss หรือไม่ และมีมาตรการควบคุมที่เหมาะสมกับระดับความอ่อนไหวของข้อมูลหรือไม่

หาก AI มีการใช้งานข้อมูลส่วนบุคคลในการประมวลผล ผู้ประกอบธุรกิจควรมีการเก็บรักษาข้อมูลส่วนบุคคลอย่างมั่นคงปลอดภัย และจัดให้มีมาตรการรักษาข้อมูลส่วนบุคคลที่เหมาะสม เช่น มีการเข้ารหัสลับข้อมูล (Encryption) และมีการทำข้อมูลส่วนบุคคลให้เป็นข้อมูลนิรนาม (Data Anonymization) เป็นต้น เพื่อลดความเสี่ยงจากกรณีข้อมูลรั่วไหล และควรปกป้องข้อมูลส่วนบุคคลตลอดวงจรชีวิต AI โดยคำนึงถึงกฎหมายที่เกี่ยวข้อง เช่น PDPA และกฎหมายลิขสิทธิ์ และต้องใช้เท่าที่จำเป็นตามวัตถุประสงค์และทำลายหรือลบข้อมูลเมื่อเลิกใช้ มีมาตรการควบคุมเหมาะสมและความเป็นส่วนตัว ให้ความสำคัญกับทุกขั้นตอนตั้งแต่รวบรวม ใช้งาน เก็บรักษา ไปจนถึงทำลาย ต้องโปร่งใสและปลอดภัยตามมาตรฐาน และหากเกิดเหตุการณ์ข้อมูลรั่วไหล ต้องแจ้งเจ้าของข้อมูลและหน่วยงานกำกับทันทีพร้อมแผนเยียวยา

ในกรณีที่ผู้ประกอบธุรกิจมีการใช้ข้อมูลเพื่อประเมินความพร้อมด้าน AI ภายในองค์กรหรือเข้าร่วมการเปรียบเทียบกับกลุ่มอุตสาหกรรม ควรกำหนดมาตรการคุ้มครองข้อมูลที่เกี่ยวข้องให้เหมาะสม เช่น การจำกัดสิทธิการเข้าถึงข้อมูลตามบทบาท การเปิดเผยผลประโยชน์ในระดับที่ไม่ทำให้ระบุตัวบุคคล ฝ่ายงาน หรือองค์กรได้เกินจำเป็น และการบันทึกการเข้าถึงข้อมูล เพื่อให้การใช้ข้อมูลเพื่อการประเมินหรือ benchmarking สอดคล้องกับวัตถุประสงค์ กฎหมายคุ้มครองข้อมูลส่วนบุคคล นโยบายธรรมาภิบาลข้อมูล และมาตรฐานความมั่นคงปลอดภัยที่เกี่ยวข้อง

2.5.4 ทักษะของบุคลากร

การนำเทคโนโลยีมาใช้ในการประกอบธุรกิจ รวมถึงการใช้งาน AI ที่ไม่เหมาะสม ไม่ถูกต้อง อาจสร้างความเสียหายต่อการดำเนินธุรกิจได้ ดังนั้น ผู้ประกอบธุรกิจต้องดำเนินการ เพื่อให้มั่นใจได้ว่าบุคลากรทุกระดับตั้งแต่ผู้ปฏิบัติงาน ไปจนถึงคณะกรรมการของผู้ประกอบธุรกิจ มีความเข้าใจเพียงพอเกี่ยวกับการใช้งาน AI ผู้ประกอบธุรกิจควรจัดให้มีบุคลากรที่รับผิดชอบในด้านการพัฒนา ทดสอบ นำไปใช้งาน และควบคุมดูแลการทำงานของ AI ที่เพียงพอและเหมาะสมกับรูปแบบการใช้งาน AI ขององค์กร การขาดทักษะ ความเชี่ยวชาญ และประสบการณ์ด้าน AI ของบุคลากรอาจส่งผลให้เกิดปัญหาในการบำรุงรักษา (Maintenance) และการพัฒนาประสิทธิภาพของโมเดลของ AI ซึ่งจะทำให้ผู้ประกอบธุรกิจต้องพึ่งพาศูนย์กลางภายนอก (Third Party) มากเกินไป (Over-Reliance) ดังนั้น ผู้บริหารระดับสูงและบุคลากรที่รับผิดชอบควรได้รับการเพิ่มพูนทักษะ ความรู้ หรือการฝึกอบรมเพื่อให้เข้าใจถึงความท้าทาย ปัญหา และความเสี่ยงที่เกี่ยวข้องกับเทคโนโลยี AI และสามารถกำหนดมาตรการควบคุมความเสี่ยงอย่างเหมาะสมก่อนที่จะตัดสินใจนำ AI มาใช้ในเชิงธุรกิจ อย่างไรก็ตาม การพัฒนาบุคลากรด้าน AI ควรมุ่งเน้นมากกว่าการสร้างความรู้หรือการอบรมทั่วไป โดยควรพิจารณาความสามารถในหลายระดับ ได้แก่ (1) ความสามารถขององค์กรในการนำ AI ไปใช้ให้เกิดผลในวงกว้าง (2) ความสามารถตามบทบาทของแต่ละฝ่ายงาน (3) ความรู้ความเข้าใจเกี่ยวกับ AI และทักษะเชิงปฏิบัติที่สามารถนำไปใช้กับงานจริงได้ ทั้งนี้ องค์กรควรประเมินช่องว่าง (Gap Analysis) ด้านความรู้และทักษะอย่างเป็นระบบ เพื่อให้สามารถออกแบบการพัฒนา การกำกับดูแล และการนำ AI ไปใช้ในแต่ละบทบาทได้เหมาะสมกับความเสี่ยงและลักษณะการใช้งาน

2.5.5 การติดตามกฎหมาย มาตรฐาน และแนวปฏิบัติที่เกี่ยวข้องจากต่างประเทศ

เมื่อผู้ประกอบการธุรกิจมีการดำเนินธุรกิจข้ามพรมแดน ใช้บริการผู้ให้บริการต่างประเทศ หรือให้บริการแก่ผู้ใช้บริการที่อาศัยอยู่ภายใต้กฎหมาย หรือข้อกำหนดของต่างประเทศ ผู้ประกอบการธุรกิจควรจัดให้มีกลไกติดตามพัฒนาการของกฎหมาย มาตรฐาน และแนวปฏิบัติที่เกี่ยวข้องกับ AI อย่างสม่ำเสมอ เพื่อประเมินผลกระทบต่อการใช้งาน AI ในมิติด้านสัญญา การโอนข้อมูลข้ามประเทศ การเปิดเผยข้อมูล การคุ้มครองผู้ใช้บริการ และการเก็บหลักฐานการติดตามดังกล่าวควรบูรณาการกับกระบวนการกำกับดูแลและบริหารความเสี่ยงขององค์กร และควรมีการทบทวน use case ที่เกี่ยวข้อง เมื่อมีข้อกำหนดใหม่ หรือมีการเปลี่ยนแปลงที่อาจกระทบต่อความสอดคล้องตามกฎหมายของธุรกิจ

อย่างไรก็ดี ผู้ประกอบการธุรกิจอาจนำมาตรฐานสากลและมาตรฐานที่เกี่ยวข้องมาปรับใช้ในการเสริมสร้างธรรมาภิบาล AI ผ่านกลไกการกำกับดูแลตนเอง (self-regulation) อาทิ การจัดทำแนวปฏิบัติที่ดี (code of practice) ร่วมกันภายในกลุ่มอุตสาหกรรมหรือสมาคมวิชาชีพ หรือการขอรับการรับรองการปฏิบัติตามมาตรฐานจากหน่วยงานที่น่าเชื่อถือทั้งภายในประเทศและต่างประเทศ เป็นต้น

2.5.6 การทดลองใช้งานในสภาพแวดล้อมควบคุม และการสร้างการมีส่วนร่วมกับหน่วยงานกำกับดูแล

สำหรับ use case ใหม่ที่มีความซับซ้อนสูงหรือยังมีความไม่แน่นอนในเชิงกฎหมายและความเสี่ยง ผู้ประกอบการธุรกิจควรพิจารณาการทดสอบในสภาพแวดล้อมควบคุม การจำกัดกลุ่มผู้ใช้ การกำหนดวงเงินหรือขอบเขตการใช้งาน และการติดตามผลอย่างใกล้ชิด ก่อนขยายการใช้งานจริง โดยเฉพาะ use case ที่เกี่ยวข้องกับ Generative AI, Agentic AI, การซื้อขายอัตโนมัติ หรือการสื่อสารกับลูกค้า

นอกจากนี้ ผู้ประกอบการธุรกิจควรมีการสื่อสารกับหน่วยงานกำกับดูแล อย่างโปร่งใส เมื่อมีการใช้งาน AI ที่มีลักษณะใหม่ มีความซับซ้อนหรือมีความเสี่ยงสูง เพื่อให้สามารถรับฟังข้อสังเกต เตรียมความพร้อมด้านข้อมูลและเอกสาร และลดความไม่แน่นอนจากการตีความในภายหลัง

ภาคผนวก

1. จุดมุ่งหมายของการนำกรอบการกำกับดูแลไปใช้และการใช้งานภาคผนวก
[Recommended Start] [For All Organizations]
2. กระบวนการ 12 ขั้นตอนสำหรับการบริหารจัดการและประเมินผลกระทบ AI (AI Impact Assessment Flow)
[Recommended Start] [For All Organizations]
3. ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยง (Risk Rubrics and Scoring for AI Use Cases)
[For Risk Assessment] [For High-risk Use Cases]
4. AI Inventory
[For All Organizations] [Quick Win] [Visibility Tool]
5. AIMS Controls
[For Control Owners] [AIMS Reference]
6. ตัวอย่างการประยุกต์ใช้เครื่องมืออย่างง่าย Lite Example - GenAI Summary Tool
[Quick Win] [For Low-risk Use Cases] [Start Here if unsure]
7. ตัวอย่างการประยุกต์ใช้เครื่องมือกรณี Use Case ความเสี่ยงสูง - Automated Trading Support
[For Emerging Use Cases] [For High-risk/Agentic Use Cases] [Full Assessment Illustration]
8. ความเสี่ยงสำคัญจากการใช้งาน AI ในภาคตลาดทุน
[Reference Tool] [Domain Lens]
9. คำถามที่พบบ่อย (FAQs)
[Implementation Guidance] [Read When in Doubt]
10. ตัวอย่างตัวชี้วัดและข้อมูลสำหรับการกำกับติดตามการใช้งาน AI
[Monitor] [Ongoing Oversight] [Risk-based]
11. ตัวอย่างแบบรายงานเหตุการณ์ที่เกี่ยวข้องกับ AI
[Report] [Incident Readiness] [Event-based]
12. ตัวอย่างแบบประเมินผู้ให้บริการ AI และความเสี่ยงจากการกระจุกตัว
[Third-party Risk] [Vendor Assessment] [Concentration Risk]
13. ตัวอย่างแนวทางการเปิดเผยการใช้ AI และการป้องกัน AI-washing
[Disclosure] [Transparency] [Anti-AI-washing]
14. ตัวอย่างคำถามสำหรับการกำกับติดตาม การทบทวนภายใน และการเตรียมความพร้อมสำหรับ supervisory dialogue
[Review] [Internal Review] [Supervisory Dialogue]

ภาคผนวก 1 จุดมุ่งหมายของการนำกรอบการกำกับดูแลไปใช้และการใช้งานภาคผนวก

กรอบการกำกับดูแลฉบับนี้จัดทำขึ้นในฐานะกรอบแนวทางแบบผสมผสาน (Hybrid Principle-based and Risk-based Approach) ที่ทำหน้าที่เป็นสะพานเชื่อมระหว่างหลักการและแนวทางสากลกับการนำไปปฏิบัติที่เหมาะสมกับบริบทของภาคตลาดทุนไทย โดยมุ่งสนับสนุนการสร้างคุณค่าจากการใช้งาน AI อย่างสมดุลผ่านผลลัพธ์เชิงคุณค่าและเชิงกลยุทธ์ Trust, Innovation, Efficiency และ Resilience (TIER) เพื่อให้การส่งเสริมนวัตกรรมและการเพิ่มประสิทธิภาพดำเนินไปควบคู่กับการบริหารความเสี่ยงและการรักษาความมั่นคงปลอดภัยของระบบ คู่มือฉบับปรับปรุงนี้มีได้ออกแบบมาเพื่อบังคับใช้ในลักษณะ one-size-fits-all หากแต่เป็นเครื่องมือสนับสนุนการตัดสินใจสำหรับคณะกรรมการและผู้บริหาร โดยอาศัยโครงสร้างแบบเลเยอร์ที่เชื่อมโยงหลักการด้านธรรมาภิบาลและจริยธรรม เช่น หลักการ Fairness, Ethics, Accountability และ Transparency (FEAT) เข้ากับการพิจารณาความเสี่ยงที่ครอบคลุมทั้งในมิติคงที่ (static risk) และมิติที่เปลี่ยนแปลงไปตามบริบท (dynamic risk) ดังที่ได้กล่าวไว้แล้วข้างต้น ทั้งนี้ คู่มือยังให้ความสำคัญกับการพิจารณาอย่างรอบคอบ (due care) ผ่านการจัดให้มีกระบวนการประเมินและบันทึกการตัดสินใจที่เหมาะสมภายในกรอบการบริหารจัดการ AI และระบบสนับสนุนที่เกี่ยวข้อง โดยเปิดโอกาสให้องค์กรสามารถเลือกและปรับใช้สัดส่วนของมาตรการควบคุมให้สอดคล้องกับลักษณะการใช้งาน AI ความพร้อมขององค์กร และระดับความเสี่ยงของแต่ละ use case เพื่อลดภาระที่ไม่จำเป็น และไม่ทำให้กรอบการกำกับดูแลเป็นอุปสรรคต่อการพัฒนาและการใช้ประโยชน์จากเทคโนโลยี

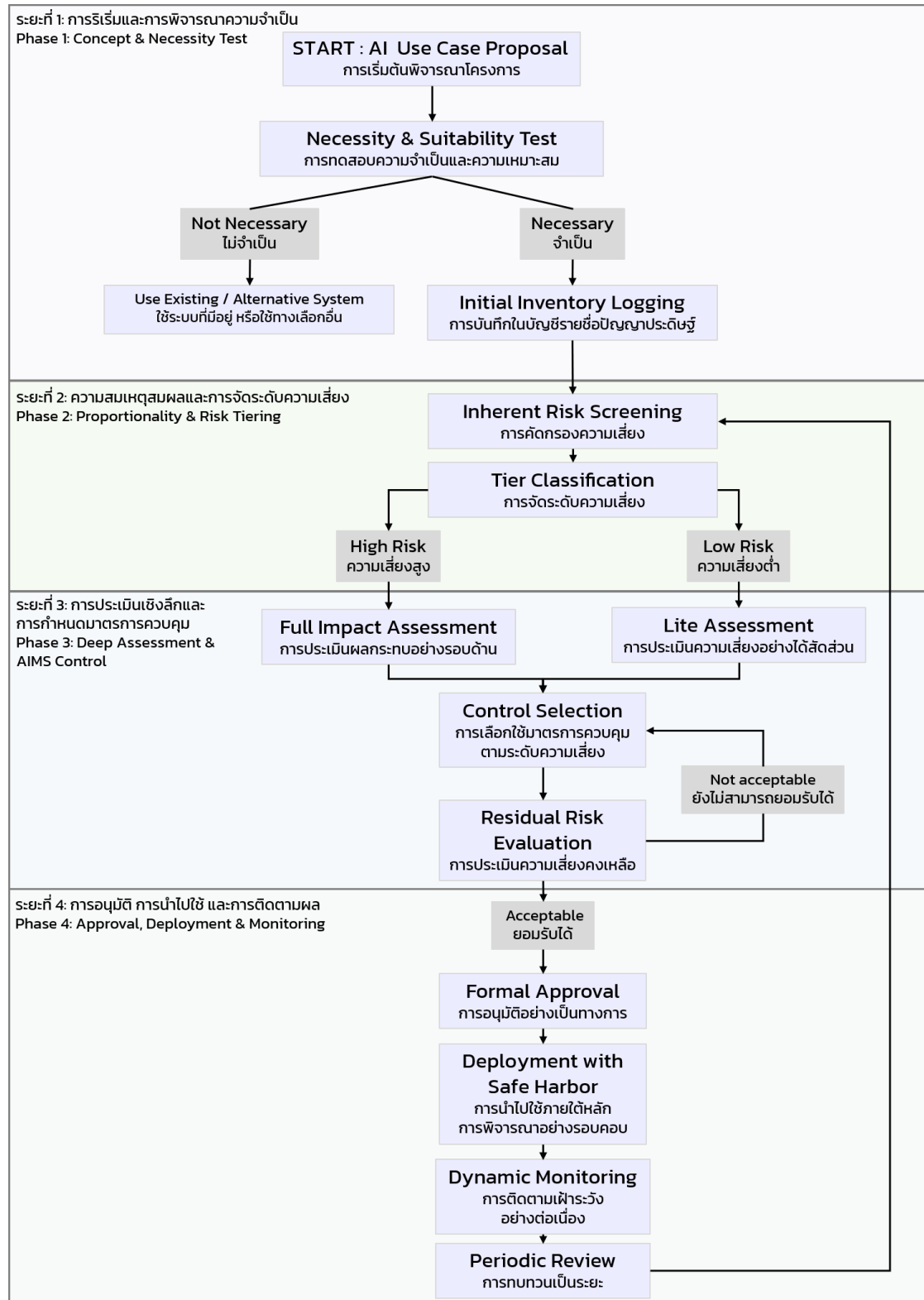
องค์กรไม่จำเป็นต้องมี AI governance ครบถ้วนสมบูรณ์ตั้งแต่เริ่มต้น แต่ควรมีระดับความพร้อมขั้นต่ำ (minimum readiness) ที่เหมาะสมกับลักษณะการใช้งาน AI และผลกระทบที่อาจเกิดขึ้น โดยอย่างน้อยควรสามารถ

- ระบุได้ว่ามีการใช้งาน AI ในส่วนใดขององค์กร
- เข้าใจวัตถุประสงค์และผลกระทบของ use case ที่มีนัยสำคัญ
- กำหนดผู้รับผิดชอบในการใช้งานและการกำกับดูแลอย่างชัดเจน
- มีมาตรการพื้นฐานด้านข้อมูล ความมั่นคงปลอดภัย และการกำกับดูแลโดยมนุษย์
- สามารถทบทวนการใช้งานเมื่อบริบท เทคโนโลยี หรือความเสี่ยงเปลี่ยนแปลง

การเริ่มใช้งานกรอบการกำกับดูแลนี้สามารถดำเนินการได้อย่างเป็นลำดับ โดยเริ่มจากการทำความเข้าใจขั้นตอนการพิจารณาและประเมิน ผ่านกระบวนการพิจารณา 12 ขั้นตอน ซึ่งทำหน้าที่เป็นกลไกการคิด (Think) และการตัดสินใจเชิงกำกับดูแล เพื่อให้ทราบว่า AI ถูกนำมาใช้ทำอะไร อยู่ในบริบทใด และมีผลกระทบต่อใครบ้าง เมื่อสามารถคิดและกำหนดกรอบการใช้งานได้แล้ว องค์กรจึงบันทึกผลการพิจารณาลงใน AI Inventory (Record) เพื่อสร้างความชัดเจนด้านความรับผิดชอบและการตรวจสอบย้อนหลัง โดยในขั้นตอนนี้สามารถใช้กรอบการประเมินความเสี่ยงจาก Risk Rubrics (เกณฑ์การประเมินความเสี่ยง) เพื่อช่วยจัดระดับความเสี่ยงของแต่ละ use case ได้อย่างเป็นระบบ จากนั้น เมื่อทราบบทบาทการใช้งานขององค์กร ลักษณะการใช้งาน และระดับความเสี่ยงแล้ว องค์กรจึงพิจารณาเลือกและดำเนินการตามมาตรการควบคุม (Act) ที่สอดคล้องกับ AI Management System (AIMS) ได้อย่างเหมาะสมกับบริบทและความเสี่ยง ทั้งนี้ ระดับความเข้มของมาตรการไม่จำเป็นต้องเท่ากันในทุกกรณี แต่ควรปรับตามหลัก Proportionality เพื่อให้สามารถเริ่มต้นการกำกับดูแล AI ได้อย่างเหมาะสม ใช้งานได้จริง และพัฒนาให้มีความสมบูรณ์ยิ่งขึ้นตามระดับความพร้อมขององค์กร

ภาคผนวก 2 กระบวนการ 12 ขั้นตอนสำหรับการประเมินผลกระทบ AI (AI Impact Assessment Flow)

กระบวนการนี้ออกแบบมาเพื่อให้องค์กรใช้เป็นตัวอย่างในการประเมินและตรวจสอบระบบ AI ที่มีผลกระทบอย่างมีนัยสำคัญต่อการดำเนินงานภายใน ต่อผู้ลงทุน หรือต่อเสถียรภาพของตลาดทุน โดยแบ่งออกเป็น 4 ระยะหลัก 12 ขั้นตอน



ภาพที่ 11 กระบวนการประเมิน 12 ขั้นตอน

Operational Decision Flow: จากการประเมินสู่การใช้งาน

Operational Decision Flow และแนวคิด PIA ใน Framework นี้ได้รับการพัฒนาโดยอาศัยการศึกษาแนวปฏิบัติจากเอกสารการประเมินผลกระทบของระบบ AI ขั้นสูงในระดับสากล กระบวนการนี้แบ่งเป็น 4 ระยะหลัก โดยเน้นที่การคัดกรอง (Screening) เพื่อให้ระบบที่ความเสี่ยงต่ำสามารถผ่านกระบวนการได้รวดเร็ว (Efficiency) และระบบที่ความเสี่ยงสูงได้รับการควบคุมอย่างเข้มงวด (Trust & Resilience) เพื่อให้มั่นใจว่าการพิจารณาครอบคลุมมิติความเสี่ยงที่สำคัญ ทั้งด้านความเป็นส่วนตัว ความได้สัดส่วน อิทธิพลต่อการตัดสินใจ และผลกระทบเชิงระบบ ควรพิจารณาใช้งานทุกระยะและทุกขั้นตอน แต่จัดสรรความสำคัญตามระดับความเสี่ยง

ทั้งนี้ กระบวนการทั้ง 12 ขั้นตอนมีวัตถุประสงค์เพื่อใช้เป็นแนวทางในการพิจารณาและทำความเข้าใจการใช้งาน AI ในเชิงกระบวนการตัดสินใจและการกำกับดูแล ดังนั้น การกำหนดและดำเนินมาตรการควบคุมในทางปฏิบัติให้เป็นไปตามกรอบ AI Management System (AIMS) ที่กำหนด ซึ่งจะกล่าวถึงในลำดับถัดไป

ตารางที่ 2 ตัวอย่างรายละเอียดกระบวนการ 12 ขั้นตอนสำหรับการบริหารจัดการและประเมินผลกระทบ AI

ขั้นตอน	รายละเอียดวิธีการ
ระยะที่ 1: การริเริ่มและการพิจารณาความจำเป็น (Phase 1: Concept & Necessity Test)	
ขั้นตอนที่ 1 การริเริ่มโครงการ AI (Initiation)	ผู้ประกอบธุรกิจควรกำหนดให้มีหน่วยงานหรือสายงานที่รับผิดชอบในการเสนอและดูแล AI use case อย่างชัดเจน รวมถึงระบุวัตถุประสงค์ของการใช้งาน AI และบริบทของการนำไปใช้ เพื่อให้เกิดความชัดเจนด้านความรับผิดชอบ (accountability) ตั้งแต่ระยะเริ่มต้น
ขั้นตอนที่ 2 การทดสอบความจำเป็น และความเหมาะสม (Necessity & Suitability Test)	ก่อนตัดสินใจใช้ AI ผู้ประกอบธุรกิจควรพิจารณาอย่างรอบคอบว่า การใช้งานดังกล่าวมีความจำเป็น หรือมีทางเลือกอื่นที่สามารถบรรลุวัตถุประสงค์เดียวกันได้ โดยมีความเสี่ยงหรือผลกระทบต่อข้อมูลส่วนบุคคลและต่อผู้ลงทุนในระดับที่น้อยกว่าหรือไม่ รวมถึงอาจต้องพิจารณาเพิ่มเติมว่า use case มีลักษณะเกี่ยวข้องกับข้อมูลสำคัญ ข้อมูลส่วนบุคคล ข้อมูลอ่อนไหว ข้อมูลลับ ฐานความรู้/RAG API แหล่งข้อมูลภายนอก ผู้ให้บริการภายนอก หรือความสามารถในการเสนอหรือดำเนินการหลายขั้นตอนแทนมนุษย์หรือไม่ ทั้งนี้ การตัดสินใจควรคำนึงถึงประโยชน์ที่ได้รับและความเสี่ยงที่อาจเกิดขึ้น กรณีที่ไม่มีความจำเป็นในการใช้ AI ผู้ประกอบธุรกิจควรพิจารณาใช้ระบบหรือกระบวนการทางเลือกอื่นแทน ทั้งนี้ หากเห็นว่ามีมีความจำเป็น ควรบันทึกเหตุผลประกอบการตัดสินใจดังกล่าวไว้อย่างเหมาะสม
ขั้นตอนที่ 3 การบันทึกในบัญชีรายชื่อ ปัญหาประติษฐ์ (Initial Inventory Logging)	เมื่อมีการตัดสินใจดำเนินการต่อ ผู้ประกอบธุรกิจควรบันทึก AI use case ลงในบัญชีรายชื่อปัญหาประติษฐ์ (“AI Inventory”) ตั้งแต่ระยะเริ่มต้น เพื่อส่งเสริมความโปร่งใสและการกำกับดูแลอย่างเป็นระบบตลอดวงจรการใช้งาน หากพบประเด็นด้านข้อมูล ระบบ ผู้ให้บริการ หรือความเป็นอิสระในการดำเนินการ ให้บันทึกใน AI Inventory เพื่อใช้ประกอบการจัดระดับความเสี่ยงและเลือก AIMS Controls ที่เกี่ยวข้อง
ระยะที่ 2: ความสมเหตุสมผลและการจัดระดับความเสี่ยง (Phase 2: Proportionality & Risk Tiering)	
ขั้นตอนที่ 4 การคัดกรองความเสี่ยง (Inherent Risk Screening)	ผู้ประกอบธุรกิจควรประเมินความเสี่ยงเบื้องต้นของ AI use case โดยพิจารณาจากลักษณะการใช้งาน วัตถุประสงค์ของการใช้งานว่าอยู่ในกลุ่มไหน และผลกระทบที่อาจเกิดขึ้น เช่น ผลกระทบต่อผู้ลงทุน ความซับซ้อนของอัลกอริทึม การใช้ข้อมูลส่วนบุคคล หรือระดับอิทธิพลของ AI ต่อการตัดสินใจ ทั้งนี้ อาจใช้เกณฑ์การประเมินความเสี่ยง (Risk Rubrics) ช่วยในการประเมินได้ ซึ่งได้ถูกกล่าวไว้ในภาคผนวกหมายเลข 3

ขั้นตอน	รายละเอียดวิธีการ
ขั้นตอนที่ 5 การจัดระดับความเสี่ยง (Tier Classification)	จากผลการคัดกรอง ผู้ประกอบธุรกิจควรจัดระดับความเสี่ยงของ AI use case เพื่อกำหนดความเข้มข้นของการประเมินและมาตรการควบคุมในขั้นถัดไป โดยทั่วไปอาจพิจารณาแบ่งเป็นกรณีที่มีความเสี่ยงต่ำ ปานกลาง และกรณีที่มีความเสี่ยงสูง ทั้งนี้ เพื่อให้การกำกับดูแลเป็นไปตามหลัก Proportionality
ระยะที่ 3: การประเมินเชิงลึกและการกำหนดมาตรการควบคุม (Phase 3: Deep Assessment & AIMS Control)	
ขั้นตอนที่ 6 การประเมินผลกระทบ อย่างรอบด้าน (Full Impact Assessment)	สำหรับ AI use case ที่มีความเสี่ยงสูง ผู้ประกอบธุรกิจควรดำเนินการประเมินเชิงลึก โดยพิจารณาผลกระทบในหลายมิติ อาทิ ความเป็นส่วนตัวของข้อมูล ความเอนเอียงและความเป็นธรรมของโมเดล ความสามารถในการอธิบายผลลัพธ์ ผลกระทบต่อผู้ลงทุน ความเป็นธรรมของตลาด และความเสี่ยงจากการพึ่งพาบุคคลภายนอก อย่างไรก็ตาม หากเป็น AI use case ที่มีความเสี่ยงต่ำ อาจพิจารณาในประเด็นสำคัญเฉพาะที่เกี่ยวข้องได้ เช่น กรณีใช้งาน Gen AI ครอบคลุมทั่วไป เพื่อสนับสนุนงานภายใน อาจใช้ checklist รายข้อ ประเมินได้โดยจัดทำตัวอย่างไว้ในภาคผนวกหมายเลข 6 และ 7
ขั้นตอนที่ 7 การเลือกใช้มาตรการ ควบคุมตามระดับ ความเสี่ยง (Control Selection)	จากผลการประเมิน ผู้ประกอบธุรกิจควรกำหนดมาตรการควบคุมที่เหมาะสมกับระดับความเสี่ยงของการใช้งาน AI เช่น มาตรการด้านความมั่นคงปลอดภัยทางไซเบอร์ การกำหนดให้มนุษย์มีส่วนร่วมในการตัดสินใจ (human oversight) การทดสอบความทนทานของระบบ หรือการตรวจสอบจากภายนอกในกรณีที่มีความเสี่ยงสูง (เมื่อทราบขอบเขตและระดับความเสี่ยงแล้ว ก็สามารถเลือกตัวอย่างมาตรการควบคุม ได้) ซึ่งได้ถูกกล่าวไว้ในภาคผนวกหมายเลข 5
ขั้นตอนที่ 8 การประเมินความเสี่ยง คงเหลือ (Residual Risk Evaluation)	ผู้ประกอบธุรกิจควรประเมินว่าความเสี่ยงที่ยังคงเหลืออยู่หลังจากนำมาตรการควบคุมมาใช้แล้วอยู่ในระดับที่องค์กรสามารถยอมรับได้หรือไม่ หากไม่สามารถยอมรับได้ ควรพิจารณาปรับปรุงมาตรการควบคุมเพิ่มเติมหรือทบทวนการใช้งาน AI ดังกล่าวใหม่
ระยะที่ 4: การอนุมัติ การนำไปใช้ และการติดตามผล (Phase 4: Approval, Deployment & Monitoring)	
ขั้นตอนที่ 9 การอนุมัติอย่างเป็นทางการ (Formal Sign-off)	ก่อนการนำ AI ไปใช้งาน ผู้ประกอบธุรกิจควรจัดให้มีการอนุมัติจากผู้มีอำนาจที่เหมาะสมตามระดับความเสี่ยงของ AI use case เช่น ผู้บริหารสายงาน หรือคณะกรรมการที่ได้รับมอบหมายด้านการกำกับดูแล AI ในกรณีที่มีความเสี่ยงสูง
ขั้นตอนที่ 10 การนำไปใช้ภายใต้หลัก การพิจารณาอย่าง รอบคอบ (Deployment with Safe Harbor)	ผู้ประกอบธุรกิจควรจัดเก็บหลักฐานการพิจารณา การตัดสินใจ และเหตุผลประกอบการกำหนดมาตรการควบคุมไว้อย่างเป็นระบบ เช่น แหล่งที่มาและสิทธิการใช้ข้อมูล ผลการทดสอบก่อนใช้งาน ขอบเขตการดำเนินการของระบบหรือ agent เงื่อนไขการอนุมัติโดยมนุษย์ แผนสำรองหรือการหยุดใช้งานชั่วคราว และหลักฐานการอนุมัติจากผู้มีอำนาจตามความเหมาะสม เพื่อแสดงให้เห็นถึงการดำเนินการด้วยความระมัดระวังอันควร (due care)
ขั้นตอนที่ 11 การติดตามเฝ้าระวังอย่าง ต่อเนื่อง (Dynamic Monitoring)	ภายหลังการนำไปใช้ ผู้ประกอบธุรกิจควรมีการติดตามและเฝ้าระวังการทำงานของ AI อย่างต่อเนื่อง เพื่อให้สามารถตรวจจับความผิดปกติ การเปลี่ยนแปลงของพฤติกรรมโมเดล (model drift) หรือเหตุการณ์ที่อาจก่อให้เกิดผลกระทบต่อผู้ลงทุนหรือการดำเนินงานได้อย่างทันท่วงที
ขั้นตอนที่ 12 การทบทวนเป็นระยะ (Periodic Review)	ผู้ประกอบธุรกิจควรทบทวนการใช้งาน AI เป็นระยะตามรอบเวลาที่กำหนด หรือเมื่อมีการเปลี่ยนแปลงที่มีนัยสำคัญ เพื่อให้การกำกับดูแล AI เป็นกระบวนการที่ต่อเนื่องและสามารถปรับให้สอดคล้องกับบริบทและความเสี่ยงที่เปลี่ยนแปลงไป

ภาคผนวก 3 ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยง (Risk Rubrics and Scoring for AI Use Case)

องค์กรสามารถกำหนดชุดเกณฑ์การประเมินความเสี่ยง AI ตามบริบทและความเสี่ยงที่มีนัยสำคัญต่อธุรกิจ โดยใช้เกณฑ์ดังกล่าวอย่างสม่ำเสมอสำหรับทุกกรณีการใช้งานประเภทเดียวกัน เพื่อให้สามารถเปรียบเทียบและจัดลำดับผลการประเมินความเสี่ยงได้อย่างเป็นธรรมชาติ ทั้งนี้ อาจมีการใช้เกณฑ์เพิ่มเติมในบางกรณีเพื่อการวิเคราะห์เชิงลึก เช่น หากมีผลกระทบต่อผู้ลงทุน ลูกค้า ตลาด ระบบสำคัญ การปฏิบัติตามกฎเกณฑ์ มีความซับซ้อนหรือความเป็นอิสระในการทำงานเกินกว่าการใช้งานทั่วไป หรือมีการพึ่งพาข้อมูล ผู้ให้บริการ โมเดล หรือระบบภายนอกที่มีนัยสำคัญ ผู้ประกอบธุรกิจควรใช้ Risk Rubrics ครบทั้ง 10 มิติ เพื่อให้การประเมินครอบคลุมความเสี่ยงที่เกี่ยวข้องอย่างเพียงพอ

(ตัวอย่าง) เกณฑ์การให้คะแนนความเสี่ยง



ภาพที่ 12 ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยง

ตารางที่ 3 ตัวอย่างเกณฑ์การให้คะแนนความเสี่ยงในการประเมินความเสี่ยงของ Use Case

เกณฑ์การให้คะแนนความเสี่ยง (1-5: คะแนนสูง เสี่ยงมาก)

เกณฑ์	คำอธิบาย (สรุป)	ระดับ 1	ระดับ 2	ระดับ 3	ระดับ 4	ระดับ 5
Market Impact	โอกาส/ขนาดผลกระทบต่อราคา/สภาพคล่อง/quote stability	ไม่มีผลต่อราคา/สภาพคล่อง ใช้เพื่อ internal analytics	ผลกระทบเฉพาะสินทรัพย์เดี่ยว ปริมาณจำกัด	อาจกระทบราคา/quote ในบางช่วงเวลา	มีนัยต่อสภาพคล่องหรือ price discovery	เสี่ยงบิดเบือนตลาดหรือกระทบ market integrity และกระทบต่อสินทรัพย์ที่มีมูลค่าเกิน X ลบ. ตามที่องค์กรกำหนด

เกณฑ์	คำอธิบาย (สรุป)	ระดับ 1	ระดับ 2	ระดับ 3	ระดับ 4	ระดับ 5
Investor Harm	ผลกระทบต่อ suitability/ mis-selling/ การตัดสินใจ/ลูกค้า	ไม่มี interaction กับลูกค้า	ให้ข้อมูลทั่วไป ไม่มีคำแนะนำ	สนับสนุนการตัดสินใจบางส่วน	มีผลต่อคำแนะนำหรือ suitability	เสี่ยง mis-selling หรือสร้างความเสียหายเป็นวงกว้าง และกระทบต่อลูกค้ามากกว่า X% ของฐานลูกค้าทั้งหมด
Systemic/ Contagion/ Market-wide dependency	เสี่ยงลามข้ามระบบ (correlations, crowding, dependency)	ใช้เฉพาะภายใน ไม่มี correlation	ใช้ในวงจำกัด	ใช้ร่วมกับ market signals	เสี่ยง herding/ crowding	เสี่ยงลามข้ามระบบหรือหลายสถาบัน
Model Risk	ความซับซ้อน/ความเสถียร/robustness	Rule-based/ deterministic	ML explainable สามารถอธิบายที่มาของผลลัพธ์	Complex ML ต้อง monitor	Black-box/ foundation model	Agentic/self-adaptive สูง
Data Risk	คุณภาพ/lineage/ provenance/ classification/ sensitive data/ bias/drift	Curated/ audited data	Internal data มี lineage	External data บางส่วน	Multi-source/ limited provenance	Dynamic/ real-time/ bias & drift สูง
Cyber/ Adversarial	โจมตีต่อโมเดล/ ข้อมูล/prompt injection	Isolated/ no external access	Basic security controls	API exposure	GenAI/ prompt-based	High-value target/ adversarial risk สูง
Third Party	พึ่งพา LLM/Cloud/ API รายใหญ่	In-house/ diversified	Vendor เดียวแต่เปลี่ยนได้	Cloud/LLM รายหลัก	Single dominant provider	Lock-in/ outage ส่งผลรุนแรง
Explainability Gap	ความสามารถ อธิบาย/ตรวจสอบ	Fully explainable	Explainable with effort	Partial explainability	Limited transparency	ไม่สามารถ อธิบายได้อย่าง มีนัยสำคัญ
Human-in-the-Loop	ระดับ oversight/ segregation of duties	Manual decision/ override และมี กระบวนการ ตรวจสอบซ้ำ (Maker-Checker) ที่ชัดเจน	Human approval required	Human-on-the-loop	Exception-based review	Fully autonomous
Regulatory Exposure	เข้าเกณฑ์ EU AI Act High-Risk/ข้อกำหนด เฉพาะ	ไม่เข้าเกณฑ์ กำกับ	Soft guidance	อยู่ใน scope regulator	EU AI Act High Risk	Prohibited/ enforcement focus

ภาคผนวก 4 AI Inventory

องค์ประกอบสำคัญของบัญชีรายชื่อปัญญาประดิษฐ์ (AI Inventory)

บัญชีรายชื่อปัญญาประดิษฐ์ควรถูกออกแบบให้เป็น ทะเบียนเชิงการตัดสินใจ (decision-anchored register) ที่สะท้อน

- บริบทการใช้งานจริง
- ระดับ “อิทธิพลต่อการตัดสินใจ” ของ AI
- ความรับผิดชอบเชิงบทบาท (ระบบ AI ที่องค์กรพัฒนาเอง (developer) หรือ ระบบ AI ที่องค์กรนำมาใช้ (user) หรือเป็นระบบงาน AI แบบผสมผสานที่องค์กรใช้ software AI แต่มีการ customized เพื่อให้เหมาะสมกับบริบทหรือเนื้องานภายในองค์กร (customizer))
- ผลลัพธ์ของการประเมินความเสี่ยงตามแนวคิด AI Impact Assessment

เพื่อให้สามารถมองเห็นภาพรวมของการใช้งาน AI ในองค์กร ควบคู่กับการบันทึกรายละเอียดเชิงกำกับดูแล กรอบการกำกับดูแลฉบับนี้เสนอให้จัดทำบัญชีรายชื่อการใช้งาน AI ในสองระดับ ได้แก่ (1) ตารางสรุปภาพรวมของกรณีการใช้งานทั้งหมด (Summary table) และ (2) ตารางรายละเอียดสำหรับแต่ละกรณีการใช้งาน (AI Inventory) โดยใช้เป็นเครื่องมือในการบันทึกเหตุผลและมาตรการกำกับดูแลอย่างเป็นระบบ

(1) ตารางสรุปภาพรวมของกรณีการใช้งานทั้งหมด (Summary Table)

วัตถุประสงค์ของตารางนี้ เพื่อแสดงภาพรวมของการใช้งาน AI ทั้งหมดในองค์กร โดยใช้วัตถุประสงค์การใช้งานในตลาดทุน และระดับอิทธิพลต่อการตัดสินใจเป็นตัวแปรหลักในการประเมินความเสี่ยงเชิงกำกับดูแลในระดับผู้บริหาร

ตารางที่ 4 ตัวอย่าง AI Use Case Summary (Register)

Use Case ID	ชื่อ Use Case	วัตถุประสงค์การใช้งานในตลาดทุน	Client-Facing/ Investor Impact	Influence Level	Overall Risk Tier	Business Owner	สถานะ
AI-UC-01		☐ Investment & Trading ☐ Advice & Services ☐ Risk & Compliance ☐ Internal Support	☐ Yes ☐ No	☐ Advisory ☐ Decision-Influencing ☐ Autonomous	☐ Low ☐ Medium ☐ High		☐ Pilot ☐ Active ☐ Retired
AI-UC-02							
AI-UC-03							
AI-UC-04							
AI-UC-05							

หมายเหตุ

วัตถุประสงค์การใช้งานในตลาดทุน (4 Objectives) - ระบุบทบาทหลักของ AI ใน ecosystem ของตลาดทุน ซึ่งสะท้อนระดับความสำคัญและลักษณะความเสี่ยงโดยรวมของการใช้งาน ได้แก่

- Investment & Trading - การซื้อขาย/พอร์ต/สภาพคล่อง
- Advice & Services - คำแนะนำ/การสื่อสารกับผู้ลงทุน
- Risk & Compliance - การกำกับดูแล/ตรวจจับความเสี่ยง
- Internal Support - งานสนับสนุนภายในองค์กร

Client-Facing/Investor Impact (Yes/No) - ระบุว่าผลลัพธ์ของ AI ถูกนำไปใช้ตัดสินใจที่ส่งผลกระทบต่อสิทธิ ผลประโยชน์ หรือการเข้าถึงบริการของลูกค้าหรือนักลงทุนโดยตรงหรือไม่

Influence Level - ระบุระดับอิทธิพลของ AI ต่อการตัดสินใจ

- Advisory - ให้ข้อมูล/ข้อเสนอ
- Decision-Influencing - จัดอันดับ/ขึ้นนำการเลือก
- Autonomous - ตัดสินใจและดำเนินการเอง

Overall Risk Tier - ระบุระดับความเสี่ยงรวมที่ได้จากการประเมินตามเกณฑ์การประเมินความเสี่ยง

- Low (ต่ำ)
- Medium (ปานกลาง)
- High (สูง)

Business Owner - ระบุฝ่ายงานผู้รับผิดชอบ use case

สถานะ - ระบุสถานะของ use case ว่าอยู่ในกระบวนการไหน

- Pilot (ต้นแบบ)
- Active (ใช้งานอยู่)
- Retired (เลิกใช้)

(2) ตารางรายละเอียดสำหรับแต่ละกรณีการใช้งาน (AI Inventory)

Use Case ID: AI-UC-__

ชื่อ Use Case:

หน่วยงานเจ้าของ (Business Owner):

วันที่จัดทำ/ทบทวน:

ตารางที่ 5 ตัวอย่าง Detailed AI Inventory (Use Case Detail)

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
1) Context & Strategic Importance - AI นี้ถูกใช้เพื่อบรรลุวัตถุประสงค์ใดบ้าง - ผลลัพธ์ของ AI ส่งผลกระทบต่อใครโดยตรงบ้าง (ผู้ลงทุน/ตลาด/องค์กร) - การใช้งานนี้มีนัยเชิง fiduciary หรือ conduct risk หรือไม่	- Purpose/วัตถุประสงค์ - Scope/Use Context - วัตถุประสงค์การใช้งานในตลาดทุน: ☐ Investment & Trading ☐ Advice & Services ☐ Risk & Compliance ☐ Internal Support - Client-Facing/Investor Impact: ☐ Yes ☐ No	- AIMS Controls - - นโยบาย AI - โครงสร้างภายในองค์กร - การประเมินผลกระทบของระบบ AI
2) Role & Development Responsibility - องค์กรมีบทบาทเพียงผู้ใช้ หรือ เข้าไปกำหนดพฤติกรรมระบบ AI - ระดับการพัฒนาส่งเสริมให้เกิดความรับผิดชอบ developer/customizer เพียงใด	- Organizational AI Role: ☐ ระบบ AI ที่องค์กรพัฒนาเอง (developer) ☐ ระบบ AI ที่องค์กรนำมาใช้ (user) ☐ ระบบงาน AI แบบผสมผสานที่องค์กรใช้ software AI แต่มีการ customized เพื่อให้เหมาะสมกับบริบทหรือหน่วยงานภายในองค์กร (customizer) - Development Characteristics: ☐ None ☐ Prompt Eng./Orchestration ☐ RAG ☐ Rule/Logic ☐ Fine-tuning ☐ Custom-trained	- AIMS Controls - - โครงสร้างภายในองค์กร - ทรัพยากรสำหรับระบบ AI - วงจรชีวิตของระบบ AI

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	• Deployment Model: ☐ In-house ☐ SaaS ☐ Hybrid • Supporting Roles: ☐ Risk ☐ Compliance ☐ IT ☐ DPO	
3) Data, System & Influence Characteristics • ข้อมูลที่ใช้มีความอ่อนไหวหรือสามารถระบุตัวบุคคลได้หรือไม่ • AI มีอิทธิพลต่อการตัดสินใจมากกว่ามนุษย์มากน้อยเพียงใด • มีการพึ่งพาข้อมูลหรือโมเดลจากภายนอกหรือไม่	• AI Technology/Model Type • Input Data Types (Internal/External/Market/Client) • PII/Sensitive Data: ☐ Low ☐ Medium ☐ High • Data Source: ☐ Internal ☐ External • Influence Level: ☐ Advisory ☐ Decision-Influencing ☐ Autonomous <i>[กรณีเสี่ยงสูง หรือกระทบข้อมูลสำคัญ จำเป็นต้องประเมินในส่วนนี้เพิ่มเติม]</i> ☐ ใช้ข้อมูลสำคัญ/ข้อมูลส่วนบุคคล/ข้อมูลอ่อนไหว/ข้อมูลลับ – (ดูเพิ่มเติม 2.4.6 Data controls) ☐ ใช้ฐานความรู้ RAG/API/แหล่งข้อมูลภายนอก - (ดูเพิ่มเติม 2.4.6 Data controls + 2.4.8 Provider/dependency controls) ☐ พึ่งพา model/platform/cloud/data provider ภายนอก - (ดูเพิ่มเติม 2.4.8 Provider/dependency controls)	• AIMS Controls - 3. ทรัพยากรสำหรับระบบ AI 5. วงจรชีวิตของระบบ AI 6. ข้อมูลสำหรับระบบ AI • AI Security Guideline สกมช. – หมวด Data Protection, Model Security & Supply chain Risk

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	☐ มี tool access หรือเชื่อมต่อระบบปฏิบัติงาน - (ดูเพิ่มเติม 2.4.5 AI lifecycle + 2.4.7 AI system use) ☐ มี autonomous หรือ semi-autonomous action - (ดูเพิ่มเติม 2.4.4 Impact assessment + 2.4.7 AI system use + human oversight) ☐ มี write/execute permission หรือเสนอ action ที่มีผลต่อกระบวนการสำคัญ - (ดูเพิ่มเติม 2.4.7 AI system use + approval/action log controls) ☐ มีผลกระทบที่ย้อนกลับได้ยาก หรือจำเป็นต้องมี fallback/suspension/recovery mechanism - (ดูเพิ่มเติม monitoring/fallback/incident controls + ภาคผนวก 10-11) ☐ อื่น ๆ : โปรดระบุ- (พิจารณา AIMS Controls ที่เกี่ยวข้องตามลักษณะเฉพาะของ use case)	
4) Decision Path & Human Oversight • มนุษย์เข้ามาในขั้นตอนใดของการตัดสินใจ (ก่อน/หลังผลกระทบ) • มีอำนาจ override หรือหยุดการทำงานของ AI จริงหรือไม่ • Oversight ที่ออกแบบไว้เพียงพอหรือไม่เมื่อความเสี่ยงเพิ่ม	• Human-in-the-Loop Position: ☐ Pre ☐ Post ☐ Advisory • Human Authority: ☐ Review ☐ Override ☐ Approve only • Oversight Effectiveness: ☐ High ☐ Medium ☐ Low (อธิบายสั้น)	• AIMS Controls - 2. โครงสร้างภายในองค์กร 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI • AI Security Guideline สกมช. – หมวด Logging, Traceability & Accountability
5) Impact & Risk (เชื่อม 10 Rubrics) • ใครได้รับผลกระทบจากการใช้งาน AI บ้าง • ความเสี่ยงหลักเกิดจาก model, data, conduct หรือ systemic effect ไດ	• Overall Risk Tier: ☐ Low ☐ Medium ☐ High • Primary Impact Areas: ☐ Client/Investor ☐ Market Integrity ☐ Conduct/Fair Treatment	• AIMS Controls - 4. การประเมินผลกระทบของระบบ AI 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI • AI Security Guideline สกมช. – หมวด Threat Modeling & Adversarial Risk

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
• มาตรการที่ใช้จัดการความเสี่ยงเพียงพอหรือไม่ อย่างไร • ใครเป็นผู้รับผิดชอบและยอมรับ residual risks	☐ Operational Efficiency ☐ Operational Resilience/Business Continuity ☐ Reputational ☐ Regulatory/Legal Compliance ☐ Privacy/Confidentiality ☐ Financial/Revenue Impact ☐ Systemic/Contagion Risk ☐ Third-party/Outsourcing Dependency ☐ Other: _____ • Key Risk Drivers (ศึกษา Risk Rubrics เพิ่มเติม) ☐ R1 Market Impact/Market Integrity Risk ☐ R2 Investor Harm/Client Impact Risk ☐ R3 Systemic/Contagion Risk ☐ R4 Model Risk: Bias/Drift/Hallucination/Accuracy ☐ R5 Data Risk: Quality/Privacy/Confidentiality/IP ☐ R6 Cyber/Adversarial Risk ☐ R7 Third Party and Concentration Risk ☐ R8 Explainability/Transparency Gap ☐ R9 Human Oversight/Automation Bias Risk ☐ R10 Regulatory/Legal/Conduct Exposure • Primary Risk Owner..... • Residual Risk Acceptance Body (Committee) 	
6) Deployment, Monitoring & Resilience • หาก AI ทำงานผิดพลาด จะตรวจจับได้เร็วเพียงใด มีมาตรการบริหารจัดการเหตุการณ์ฉุกเฉินเพียงพอหรือไม่ • สามารถหยุดหรือย้อนกลับไปได้ไหม • โมเดลสำรอง (fallback) ได้หรือไม่	• Key Controls in Place (High-level) • Review Frequency: ☐ Monthly ☐ Quarterly ☐ Annually ☐ Event-driven • Monitoring Focus: ☐ Bias/Fairness ☐ Model Drift/Performance Degradation	• AIMS Controls - 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI 8. ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก • AI Security Guideline สกมช. - หมวด Incident Response, Resilience & Recovery

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
• ใครเป็นผู้ตัดสินใจเมื่อเกิด incident	☐ Hallucination/Inaccurate Output ☐ Prompt Injection/Misuse ☐ Data Leakage/Sensitive Data Exposure ☐ Complaints/User Feedback ☐ Suitability/Appropriateness of Advice ☐ Market Abuse/Abnormal Trading Pattern ☐ Cybersecurity/Adversarial Events ☐ Third-party/Vendor Performance ☐ Availability/Service Disruption ☐ Regulatory/Compliance Breach ☐ Other: _____ • Incident/Kill Switch: ☐ Yes ☐ No • Fallback Process..... • Third-party Provider (if any)	

หมายเหตุ

ตารางนี้มีวัตถุประสงค์เพื่อสนับสนุนการพิจารณาและบันทึกเหตุผลเชิงกำกับดูแล สำหรับแต่ละกรณีการใช้งาน AI โดยการกำหนดและดำเนินการควบคุมในทางปฏิบัติให้เป็นไปตามระบบการจัดการ AI (AIMS) ที่กำหนดไว้ในเนื้อหาหลักและภาคผนวกอื่น ทั้งนี้ ตาราง AI Inventory นี้ทำหน้าที่เป็นกลไกเชื่อมโยงการประเมินความเสี่ยงกับการเลือกใช้มาตรการควบคุมอย่างได้สัดส่วนโดยไม่ซ้ำซ้อนกับ AIMS ซึ่งองค์กรสามารถประยุกต์ใช้ได้ตามความเหมาะสม สอดคล้องกับความเสี่ยงและบริบทการใช้งาน

ภาคผนวก 5 AIMS Controls

นิยามบทบาทของ Developer, User และ Customizer

Developer (ผู้พัฒนาและควบคุมการออกแบบระบบ AI)

บุคคลหรือทีมภายในองค์กรที่รับผิดชอบการออกแบบ พัฒนา ทดสอบ ตรวจสอบ ปรับปรุง หรือควบคุมองค์ประกอบหลักของโมเดลหรือระบบ AI เพื่อวัตถุประสงค์เฉพาะขององค์กร รวมถึงการกำหนดสถาปัตยกรรมระบบ การจัดการข้อมูล การตรวจสอบความถูกต้องของโมเดล การกำหนดมาตรการควบคุม และการบริหารวงจรชีวิตของระบบ AI ให้มีความปลอดภัย เชื่อถือได้ และสอดคล้องกับข้อกำหนดขององค์กร เช่น การพัฒนาโมเดล ML ภายในองค์กร การสร้าง RAG Architecture เอง หรือ การพัฒนา Agentic AI Workflow เอง เป็นต้น

หน้าที่หลัก

- กำกับดูแลด้าน AI Lifecycle และ Data Governance
- ออกแบบ ทดสอบ และตรวจสอบความถูกต้องของโมเดลหรือระบบ AI
- จัดทำเอกสารกำกับดูแลและหลักฐานกำกับ เช่น Model Card, Data Sheet, Validation Report, Testing Evidence และ Release Gates
- กำกับดูแลความเสี่ยงทางเทคนิคจาก API, LLM, Cloud Service หรือ Foundation Model
- สนับสนุนการตรวจสอบย้อนหลัง ความโปร่งใส และการบริหารความเสี่ยงตลอดวงจรชีวิตระบบ AI

User (ผู้ใช้งานระบบ AI)

หน่วยงานหรือบุคคลที่นำระบบ AI ไปใช้จริงในกระบวนการทำงาน การให้บริการลูกค้า การวิเคราะห์ การลงทุน การวิจัย หรือการกำกับดูแล โดยมีหน้าที่ใช้งานระบบ AI ให้เป็นไปตามนโยบาย มาตรการควบคุม และข้อกำหนดขององค์กร โดยไม่ได้พัฒนา ปรับแต่ง หรือเชื่อมต่อกับระบบ AI ในเชิงลึก เช่น เจ้าหน้าที่ใช้ GenAI ช่วยสรุปเอกสาร นักวิเคราะห์ใช้ AI ช่วยวิเคราะห์ข้อมูล หรือฝ่ายกำกับดูแลใช้ AI ช่วยตรวจสอบข้อมูล เป็นต้น

หน้าที่หลัก

- ใช้งาน AI อย่างรับผิดชอบและสอดคล้องกับนโยบายองค์กร
- ตรวจสอบความเหมาะสมของผลลัพธ์ก่อนนำไปใช้
- เปิดเผยหรือสื่อสารการใช้ AI ต่อผู้มีส่วนได้เสียตามข้อกำหนด
- รายงานเหตุการณ์ ความผิดปกติ หรือผลกระทบที่พบจากการใช้งาน
- สนับสนุนการติดตามผลหลังใช้งานจริง
- มีส่วนร่วมในการประเมินผลกระทบและปรับปรุง controls

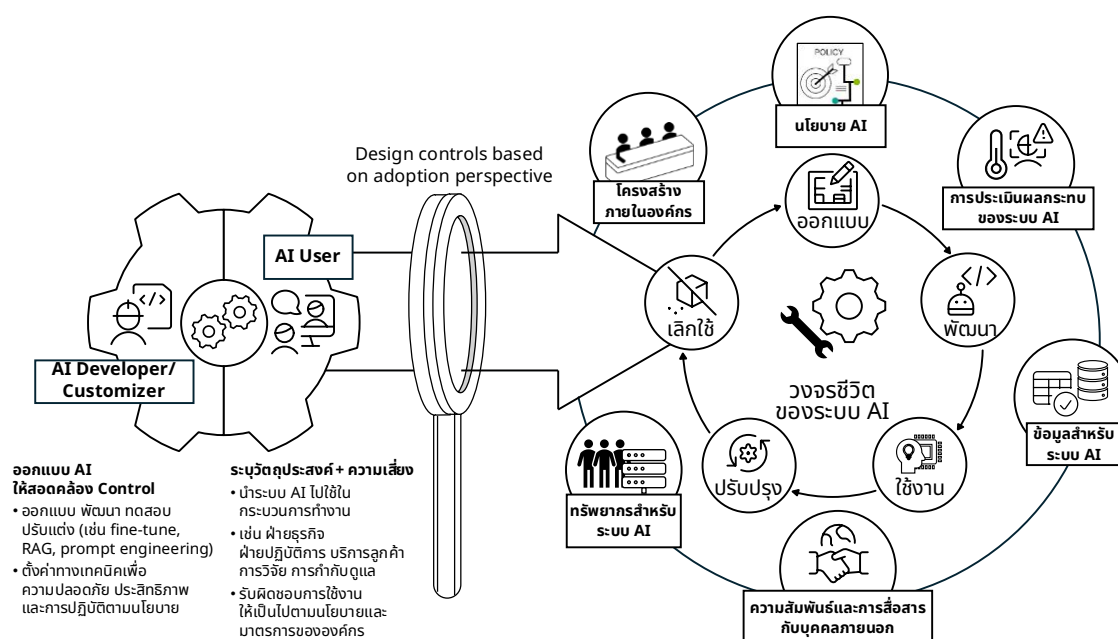
Customizer (ผู้ปรับแต่งและบูรณาการระบบ AI)

บุคคลหรือทีมภายในองค์กรที่นำโมเดล ระบบ หรือบริการ AI ที่มีอยู่แล้วมาปรับแต่ง เชื่อมต่อ กำหนดค่า หรือออกแบบ workflow เพื่อให้เหมาะสมกับกระบวนการทำงานขององค์กร โดยไม่ได้พัฒนาโมเดลหลักหรือแพลตฟอร์ม AI ตั้งแต่ต้น แต่มีอิทธิพลต่อพฤติกรรม ขอบเขตการทำงาน ความสามารถ หรือระดับความเป็นอิสระของระบบ AI เช่น Prompt Engineering, Prompt Orchestration, Agentic AI Orchestration, RAG Configuration หรือ API/Tool/Knowledge Base Integration เป็นต้น

หน้าที่หลัก

- ออกแบบ workflow ระหว่าง AI เครื่องมือ และระบบงานภายใน
- กำหนดขอบเขตการทำงาน หรือ Action Boundary
- กำหนดระดับความเป็นอิสระ หรือ Autonomy Level
- จัดการการเชื่อมต่อข้อมูล เครื่องมือ และระบบภายนอก
- กำหนด controls เช่น Human Oversight, Approval Checkpoint, Tool Restriction และ Runtime Monitoring
- จัดทำ AI Inventory และเอกสารที่สะท้อนลักษณะการใช้งานจริง
- สนับสนุนการกำกับดูแล third party และการประเมินผลกระทบจากการเชื่อมต่อระบบ

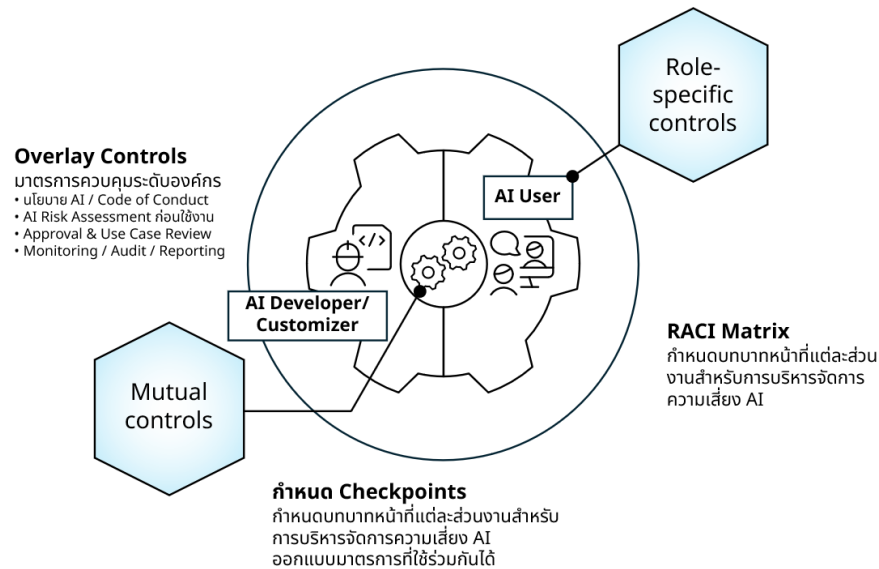
ทั้งนี้ การแบ่งบทบาทเพื่อเลือกใช้มาตรการควบคุมเป็น Developer, User และ Customizer มีวัตถุประสงค์เพื่อช่วยให้ผู้ประกอบธุรกิจระบุลักษณะการมีส่วนร่วมขององค์กรต่อระบบ AI ได้ชัดเจนขึ้น โดย Customizer เป็นบทบาทกึ่งกลางสำหรับกรณีที่องค์กรไม่ได้พัฒนาโมเดลหลักเอง แต่มีการปรับแต่ง เชื่อมต่อ กำหนดค่าหรือออกแบบ workflow ที่มีผลต่อพฤติกรรมของระบบ AI อย่างไรก็ตาม การเลือกมาตรการควบคุมควรพิจารณาตามขอบเขตความรับผิดชอบขององค์กรในแต่ละ use case และระดับความเสี่ยงตามหลัก Proportionality



ภาพที่ 13 การแบ่งบทบาทเพื่อเลือกใช้มาตรการควบคุม

AIMS Controls: Mutual and Role-specific Controls

"สำหรับมิติด้านความมั่นคงปลอดภัยไซเบอร์เชิงลึก สามารถอ้างอิงตามมาตรฐาน AI Security ของ สกมช."



ภาพที่ 14 Mutual and Role-specific Controls

ตารางที่ 6 ตัวอย่าง AIMS Controls: Mutual and Role-specific Controls

AIMS CONTROLS	DEVELOPER/CUSTOMIZER	USER
1. นโยบาย AI ☒ Mutual Controls ☐ Role-specific Controls	- จัดทำ AI Policy ที่ครอบคลุมหลักการจริยธรรม AI ที่มีการส่งเสริม secure by design ในทุกขั้นตอนของ AI Lifecycle เริ่มตั้งแต่ขั้นตอนออกแบบไปจนถึงการยกเลิกใช้งาน และสอดคล้องกับมาตรฐานและแนวทางปฏิบัติของหน่วยงานกำกับดูแลอื่น รวมถึงมีการทบทวนนโยบายตามรอบระยะเวลาที่กำหนด - เหตุผล: ให้นโยบายเป็น “รากฐาน” ของการกำกับ, ลด conduct risk และ misalignment ตั้งแต่ต้นน้ำ, สอดคล้องการจัดการความเสี่ยงต่อเนื่อง - Minimum Evidence: Policy เวอร์ชันล่าสุด + บันทึกอนุมัติ, แนวทาง/มาตรฐานที่ใช้อ้างอิง, Record การสื่อสาร/อบรม	- จัดทำ นโยบายการใช้งานที่ยอมรับได้ (Use Policy/AUP) ที่อธิบายขอบเขตและข้อห้ามการใช้งานให้มีความรับผิดชอบและมีจริยธรรม รวมไปถึงนโยบายการสื่อสารกับลูกค้าหรือผู้ใช้บริการ ไปจนถึงการใช้งานที่เกี่ยวข้องกับข้อมูลและความปลอดภัย และสอดคล้องกับมาตรฐานและแนวทางปฏิบัติของหน่วยงานกำกับดูแลอื่น รวมถึงมีการทบทวนนโยบายตามรอบระยะเวลาที่กำหนด - เหตุผล: จุดใช้งานมีการโปร่งใสต่อผู้ใช้/ลูกค้า, ลด mis selling/misleading - Minimum Evidence: Use/AUP, template การแจ้งลูกค้า, บันทึกอบรม frontline
2. โครงสร้างภายในองค์กร ☒ Mutual Controls ☐ Role-specific Controls	- ผังบทบาทหน้าที่ (RACI Matrix) เช่น เจ้าของโมเดล ผู้ดูแลข้อมูล ผู้ตรวจสอบ ผู้รับผิดชอบ ความปลอดภัย และผู้มีสิทธิอนุมัติ - ช่องทางแจ้งเหตุผิดปกติ และ เส้นทางส่งต่อหากปัญหายกระดับ (Escalation path)	กำหนดผู้รับผิดชอบทั้งในการใช้งานสำหรับธุรกิจและผู้รับผิดชอบด้านมาตรการควบคุม ผู้มีสิทธิตรวจสอบ-อนุมัติผลจาก AI และจัดทำแผนฝึกอบรมเพื่อให้มั่นใจว่าผู้รับผิดชอบดังกล่าวมี

AIMS CONTROLS	DEVELOPER/CUSTOMIZER	USER
	เหตุผล: บทบาทชัด ลดช่อง oversight, เน้น accountability ตั้งแต่ design ถึง monitor Evidence: RACI matrix, ToR คณะกรรมการ, Escalation SOP	ความรู้ความเข้าใจความเชี่ยวชาญในการใช้งาน AI อย่างรับผิดชอบ เหตุผล: ลด automation bias/over reliance, ยืนยันความพร้อมบุคลากรที่รับผิดชอบตัดสินใจ Evidence: คำบรรยายบทบาท, บันทึกอบรม/ทดสอบความรู้, รายชื่อ Reviewer
3. ทรัพยากรสำหรับระบบ AI ☐ Mutual Controls ☒ Role-specific Controls	ระบุและจัดทำรายชื่อทรัพยากรสำคัญ เช่น ข้อมูล เครื่องมือ ระบบและการประมวลผล ทีมงาน พร้อมตารางแสดงหน้าที่ตามทักษะของทีม เพื่อควบคุมคุณภาพและความปลอดภัยของระบบ AI เหตุผล: คุณภาพ/ความปลอดภัยขึ้นกับทรัพยากร, กำหนดให้ระบุ/บริหารทรัพยากรสำคัญ Evidence: Inventories, Configuration Management Database (CMDB) extracts (หลักฐานรายการทรัพยากรและ dependency ของระบบ), Hardened baseline/benchmark, แผนพัฒนาทักษะ	
4. การประเมินผลกระทบของระบบ AI ☒ Mutual Controls ☐ Role-specific Controls	- กำหนดกระบวนการเพื่อประเมินผลกระทบและดำเนินการประเมินผลกระทบและความเสี่ยงจาก AI รวมถึงกรณีถูกนำไปใช้ผิด เช่น โจมตีด้วย prompt injection หรือ data poisoning ซึ่งควรครอบคลุม foreseeable misuse, การจัดทำ checklist Go/No Go พร้อมเงื่อนไข และควรมีการบันทึกผลการประเมินและจัดเก็บไว้ตามระยะเวลาที่กำหนด - จัดทำตัวอย่างกรณี abuse case หรือการจำลองภัยคุกคาม เพื่อเตรียมความพร้อมประเมินผลกระทบและดำเนินการบริหารความเสี่ยงที่อาจเกิดขึ้นจากการใช้ AI use case เหตุผล: เน้นผลกระทบต่อบุคคล/สังคม, ประเมินทบทวนตลอดวงจรชีวิต (รวม misuse) Evidence: Impact report, Threat model, Risk register link, Sign off ก่อน release	กำหนดกระบวนการเพื่อประเมินผลกระทบและดำเนินการประเมินผลกระทบในบริบทต่าง ๆ (Contextual impact) ที่อาจเกิดขึ้นได้จริง (ลูกค้า/ตลาด/ปฏิบัติการ), การจัดทำ checklist Go/No Go พร้อมเงื่อนไข และควรมีการบันทึกผลการประเมินและจัดเก็บไว้ตามระยะเวลาที่กำหนด เหตุผล: ความเสียหายเกิดที่จุดใช้งาน, IOSCO เน้น investor protection/market integrity - ต้องประเมินในบริบทจริงก่อนใช้งานในวงกว้าง Evidence: Pilot report, Go live checklist, Residual risk memo

AIMS CONTROLS	DEVELOPER/CUSTOMIZER	USER
5. วงจรชีวิตของระบบ AI ☒ Mutual Controls ☐ Role-specific Controls	• การออกแบบและระบุขั้นตอนวงจรชีวิตของ AI และกำหนดความปลอดภัยในทุกขั้นตอน • กำหนดให้มีมาตรการตรวจสอบและทบทวนระบบ AI เกณฑ์การใช้งานและจัดทำเป็นเอกสาร • จัดทำทะเบียนเวอร์ชันโมเดล เกณฑ์และเงื่อนไขสำหรับการพิจารณาอนุมัติก่อนเผยแพร่ และแผนสำรองกรณีต้องถอนกลับมาใช้เวอร์ชันก่อนหน้า (Rollback plan) • กำหนดและจัดทำเอกสารเกี่ยวกับการดำเนินงานระบบ AI การควบคุมและติดตามตรวจสอบการทำงาน และมีการรายงานอย่างต่อเนื่อง รวมถึงมีการบันทึก logs การใช้งาน • กำหนดเทคโนโลยีที่ใช้ในระบบ AI และจัดทำเอกสารทางเทคนิค (Technical Documentation) อย่างเพียงพอ สำหรับผู้มีส่วนได้เสียแต่ละประเภทที่เกี่ยวข้อง เช่น ผู้ใช้ คู่ค้า หน่วยงานกำกับดูแล เป็นต้น เหตุผล: ครอบคลุมทุกช่วงวงจรชีวิต, กำหนด risk mgmt. เป็นวงจรต่อเนื่อง Evidence: SDLC artifacts, Registry entries, Release/rollback records	• Playbook คู่มือขั้นตอนใช้งาน วิธีติดตามผลวิธีรับมือเหตุการณ์ผิดปกติ และมอบหมายผู้รับผิดชอบในการติดตาม และเฝ้าระวังเหตุการณ์หลังใช้งาน (post market monitoring) • กำหนดและจัดทำเอกสารเกี่ยวกับการดำเนินงานระบบ AI การควบคุมและติดตาม ตรวจสอบการทำงาน และมีการรายงานอย่างต่อเนื่อง รวมถึงมีการบันทึก logs การใช้งาน เหตุผล: ผู้ใช้งานคือเจ้าของ phase operate/monitor, IOSCO ย้ำความสำคัญของการเฝ้าติดตามต่อเนื่องต่อความซื่อสัตย์ของตลาด/นักลงทุน Evidence: KPIs/KRIs (drift, FP/FN, abnormal impact), Incident logs, PIR (post incident review)
6. ข้อมูลสำหรับระบบ AI ☐ Mutual Controls ☒ Role-specific Controls	• การติดตามแหล่งที่มาของข้อมูล คุณภาพข้อมูล การเตรียมข้อมูลและวิธีการจัดเตรียมข้อมูล ตรวจสอบอคติ และเช็คเรื่องสิทธิส่วนบุคคล/ลิขสิทธิ์ รวมถึงกรณีปรับแต่งด้วยข้อมูล (fine tune/RAG) เชื่อมข้อมูลที่ใช้กับ AI กับกลไกธรรมาภิบาลข้อมูลที่มีอยู่ เช่น Data Catalog/Metadata/ROPA/Data Flow/Data Owner ตามความเหมาะสม • กำหนดแนวทางควบคุม input, prompt, file upload, RAG/API data, output, log และ action record ตามระดับชั้นและความเสี่ยง เหตุผล: เน้นมาตรฐานข้อมูล - หัวใจของ model risk, ลดการละเมิดสิทธิส่วนบุคคล/ลิขสิทธิ์ และลดความเสี่ยงข้อมูลรั่วไหล ใช้ข้อมูลเกินจำเป็น เก็บข้อมูลเกินขอบเขต Evidence: Data sheets, Source licences/ ToS, DQ dashboard, Bias/fairness report	

AIMS CONTROLS	DEVELOPER/CUSTOMIZER	USER
	Data catalog/metadata/ROPA, Prompt/output/log handling rule, RAG/API source list, access matrix, retention/deletion rule	
7. การใช้งานระบบ AI ☐ Mutual Controls ☒ Role-specific Controls	- ออกแบบและวางแผน ข้อจำกัดเพื่อความปลอดภัย (Guardrails/constraints) เพื่อให้การใช้งานระบบ AI เป็นไปตามวัตถุประสงค์ เช่น การคัดกรองเพื่อความปลอดภัย กำหนดขอบเขตใช้งาน (safety filters, policy enforcement, rate limits) ระบบตรวจสอบโดยมนุษย์ในจุดเสี่ยงสูง (HITL design) - เหตุผล: ลด harm/automation bias เป็นรากฐานให้ user ใช้ได้ปลอดภัย - Evidence: Safety config, Enforcement tests, Decision logs template	- จัดทำคู่มือหรือแนวทางปฏิบัติงานมาตรฐานที่ระบุวัตถุประสงค์ในการใช้งาน AI อย่างมีความรับผิดชอบ และระบุขั้นตอนอย่างละเอียดเป็นขั้นเป็นตอน เพื่อให้พนักงานทุกคนดำเนินงานในกระบวนการเดียวกันอย่างถูกต้อง สม่าเสมอ และลดข้อผิดพลาด (SOP) การจัดให้มี Training ผู้ใช้งาน, การรวบรวมและวิเคราะห์การใช้งาน (Usage analytics) เพื่อจับพฤติกรรมผิดปกติ เช่น พนักงานควรมีความรู้ ความเข้าใจเกี่ยวกับการใช้งานข้อมูลตาม Data governance (Data Classification) - กำหนดให้มีการตรวจสอบให้แน่ใจว่าการใช้งานระบบ AI เป็นไปตามวัตถุประสงค์และจัดทำเอกสารประกอบ - เหตุผล: จุดใช้งานคือแหล่งเกิดอันตราย, SOP และการฝึกช่วยลด human error/over reliance, analytics สนับสนุน post market monitoring - Evidence: SOP/Playbook, Training logs, Usage dashboards & alerts
8. ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก ☐ Mutual Controls ☒ Role-specific Controls	- จัดทำเอกสารโมเดล AI ที่มีการสรุปวัตถุประสงค์ความสามารถและข้อจำกัดของโมเดล ประเด็นเสี่ยง จุดล้มเหลวที่พบแล้วของโมเดล (purpose, capability, limitation, metrics, risks, known failure modes) ในรูปแบบที่ user สามารถนำไปใช้สื่อสารภายนอกได้ - เหตุผล: โปร่งใส/ตรวจสอบได้, ช่วยให้ user ประเมินบริบทการใช้งานและหน้าที่เปิดเผยต่อผู้เกี่ยวข้อง - Evidence: Model Card, Limitation statement, Change log/Release notes - กระบวนการตรวจสอบผู้ให้บริการภายนอก (Vendor due diligence) เช่น LLM/API/Cloud สิทธิ์ในการตรวจสอบ	- แจ้งเตือนลูกค้าหรือผู้ใช้งานว่า "กำลังโต้ตอบกับ AI" ให้คำแนะนำเรื่องข้อจำกัด-วิธีใช้ที่ปลอดภัย ป้ายกำกับเอาต์พุตตามกฎหมาย (เช่น "กำลังโต้ตอบกับ AI"), คำแนะนำการใช้อย่างปลอดภัย/ข้อจำกัด; ป้าย output/watermark - เหตุผล: ลด mis use/ความคาดหวังผิด, รองรับหลักการ transparency และข้อกำหนดบางเขตอำนาจ - Evidence: UX notice screenshots, Help center articles, Output labels/watermark - ติดตามคุณภาพบริการจาก vendor เช่น SLA/KPI แชร้การเกิด incident และมีการแจ้ง patch กำหนดมาตรฐานจำนวนครั้งของการเกิด

AIMS CONTROLS	DEVELOPER/CUSTOMIZER	USER
	เปลี่ยนแปลง หรือออกจากบริการ การทดสอบความพร้อมรับมือเหตุการณ์ และผลการตรวจสอบถูกส่งต่อให้ user เพื่อใช้กำหนดระดับ disclosure/contingency • ระบุ dependency จาก model, platform, cloud, API, data provider หรือเครื่องมือภายนอก และกำหนดแนวทางรองรับ change/outage/exit เหตุผล: ลด concentration/outsourcing risk ที่ IOSCO ระบุเป็นช่องโหว่ระบบ, การบริหารคู่ค้า/ลูกค้าอย่างเป็นระบบ Evidence: Vendor profile and risk report, SLA/DPA, Exit run book, Resilience test reports	เหตุการณ์ต่อลูกค้าให้ทราบ และต้อง feedback กลับไปยัง developer/customizer เพื่อปรับปรุงโมเดลหรือ control เหตุผล: ผู้ใช้กระทบโดยตรงเมื่อผู้ให้บริการเปลี่ยน/ล้ม, ต้องรักษามาตรฐานสื่อสารกับลูกค้าในเหตุการณ์ผิดปกติ Evidence: Vendor scorecard (รายละเอียดการประเมินผู้ให้บริการภายนอกหรือ third party), Monthly/quarterly review minutes, Customer comms templates

หมายเหตุ

Minimum Evidence เป็นตัวอย่างหลักฐานขั้นต่ำที่สามารถใช้แสดงการควบคุมได้ ผู้ประกอบธุรกิจอาจใช้หลักฐานอื่นที่มีสาระเทียบเคียงกันได้ตามลักษณะ use case บทบาทขององค์กร บริบทการใช้งานและระดับความเสี่ยง

ภาคผนวก 6 ตัวอย่างการประยุกต์ใช้เครื่องมืออย่างง่าย Lite Example — GenAI

Summary Tool

1. ตัวอย่างการประเมิน 4 ระยะ 12 ขั้นตอน

วัตถุประสงค์ของส่วนนี้

ตัวอย่างนี้แสดงให้เห็นว่า ในการพิจารณาตามกระบวนการ 12 ขั้นตอนสามารถทำแบบกระชับได้สำหรับกรณี use case ความเสี่ยงต่ำ โดยยังคงครอบคลุมทั้งเรื่องความจำเป็น บริบท ความเสี่ยง การควบคุม และการติดตามผล

ตารางที่ 7 ตัวอย่างการประเมินกระบวนการ 12 ขั้นตอนของ GenAI Summary Tool

ระยะ	คำถามที่ควรคำนึงถึง	ตัวอย่างคำตอบ: GenAI Summary Tool
Phase 1: Necessity & Context	AI นี้ใช้ทำอะไร จำเป็นหรือไม่ และควรลงทะเบียนเป็น use case หรือไม่	ใช้ GenAI ช่วยสรุปรายงานวิจัยจากต่างประเทศเพื่อช่วยให้นักวิเคราะห์อ่านเร็วขึ้น มีความจำเป็นเชิงประสิทธิภาพเพราะเอกสารมีจำนวนมาก จึงควรลงทะเบียนเป็น use case ใน AI Inventory
Phase 2: Risk & Proportionality	ใช้กับใคร ใช้ข้อมูลอะไร กระทบลูกค้า หรือผู้ลงทุนหรือไม่ และความเสี่ยงเบื้องต้นอยู่ระดับใด	ใช้ภายในฝ่ายวิเคราะห์ ไม่ส่งให้ลูกค้าโดยตรง ใช้รายงานวิจัยสาธารณะหรือเอกสารที่ไม่ใช่ข้อมูลลับ ไม่มี PII หรือ sensitive data จัดเป็น Low Risk
Phase 3: Control Thinking	ต้องมี controls ขั้นต่ำอะไรเพื่อจัดการความเสี่ยงหลัก และความเสี่ยงคงเหลือยอมรับได้หรือไม่	ความเสี่ยงหลักคือ hallucination และ data leakage จึงควรมีข้อห้ามป้อนข้อมูลลับ human review access control และ disclaimer ความเสี่ยงคงเหลือยอมรับได้ เพราะใช้เป็นข้อมูลประกอบและมีคนตรวจทาน
Phase 4: Deployment & Review	ใครอนุมัติ ใช้งานภายใต้เงื่อนไขใด ติดตามอย่างไร และทบทวนเมื่อใด	อนุมัติโดยหัวหน้าฝ่ายวิจัย/IT ตามอำนาจภายใน จัดทำ usage guideline ติดตาม feedback กรณีสรุปผิด และทบทวนปีละครั้งหรือเมื่อเปลี่ยนโมเดล/ผู้ให้บริการ

หมายเหตุ

use case นี้สามารถใช้การประเมินแบบกระชับได้ เนื่องจากเป็นการใช้งานภายใน มี AI เป็นเพียงเครื่องมือช่วยสนับสนุนการทำงาน ซึ่งใน use case นี้คือการสรุปเอกสารผลงานวิจัย ไม่มีผลต่อผู้ลงทุนโดยตรง ไม่มีการใช้ข้อมูลอ่อนไหว หรือข้อมูลที่ถูกจัดชั้นความลับ และมีมนุษย์ตรวจสอบผลลัพธ์ก่อนนำไปใช้ ทั้งนี้ ผลการพิจารณาควรถูกบันทึกลง AI Inventory ตาม template ปกติ เพื่อให้เกิด visibility และ traceability

2. ตัวอย่าง AI Inventory แบบ Lite-compatible

กรณีศึกษา: GenAI Research Summary Tool

Use Case ID: AI-UC-INT-01

ชื่อ Use Case: GenAI Research Summary Tool

หน่วยงานเจ้าของ (Business Owner): ฝ่ายวิเคราะห์หลักทรัพย์

วันที่จัดทำ/ทบทวน:

หมายเหตุ: ตารางนี้ใช้ template เดียวกับ AI Inventory ตามปกติ แต่กรอกในระดับความลึกที่เหมาะสมกับ use case ความเสี่ยงต่ำ โดยยกตัวอย่าง risk rubrics บางส่วนมาร่วมวิเคราะห์เป็นตัวอย่าง ทั้งนี้ AIMS Controls ในคู่มือถูกออกแบบให้ครอบคลุมทั้ง mutual controls และ role-specific controls สำหรับ Developer/Customizer/User ตลอดจนวงจรชีวิต AI

ตารางที่ 8 ตัวอย่างการจัดทำ AI Inventory ของ GenAI Summary Tool

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
1) Context & Strategic Importance • AI นี้ถูกใช้เพื่อบรรลุวัตถุประสงค์ใด บริบทไหน • ผลลัพธ์ของ AI ส่งผลกระทบต่อใครโดยตรงบ้าง (ผู้ลงทุน/ตลาด/องค์กร) • การใช้งานนี้มีนัยเชิง fiduciary หรือ conduct risk หรือไม่	• Purpose/วัตถุประสงค์สรุปบทวิเคราะห์จากแหล่งข้อมูลต่างประเทศเพื่อช่วยให้นักวิเคราะห์อ่านและคัดกรองข้อมูลได้เร็วขึ้น..... • Scope/Use Contextใช้ภายในฝ่ายวิเคราะห์เท่านั้น ไม่ส่งผลลัพธ์ให้ลูกค้าหรือผู้ลงทุนโดยตรง..... • วัตถุประสงค์การใช้งานในตลาดทุน: ☐ Investment & Trading ☐ Advice & Services ☐ Risk & Compliance ☒ Internal Support • Client-Facing/Investor Impact: ☐ Yes ☒ No	• AIMS Controls - 1. นโยบาย AI 2. โครงสร้างภายในองค์กร 4. การประเมินผลกระทบของระบบ AI
2) Role & Development Responsibility • องค์กรมีบทบาทเพียงผู้ใช้ หรือเข้าไปกำหนดพฤติกรรมระบบ AI • ระดับการพัฒนาส่งเสริมให้เกิดความรับผิดชอบ developer/customizer เพียงใด	• Organizational AI Role: ☐ ระบบ AI ที่องค์กรพัฒนาเอง (developer) ☒ ระบบ AI ที่องค์กรนำมาใช้ (user) ☐ ระบบงาน AI แบบผสมผสานที่องค์กรใช้ software AI แต่มีการ customized เพื่อให้เหมาะสมกับบริบทหรืองานภายในองค์กร (customizer) • Development Characteristics: ☐ None	• AIMS Controls - 2. โครงสร้างภายในองค์กร 3. ทรัพยากรสำหรับระบบ AI 5. วงจรชีวิตของระบบ AI

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	☒ Prompt Eng./Orchestration (กรณีมี prompt template หรือ workflow การสรุปที่องค์กรกำหนดเอง) หรือ ☒ None (หากใช้งาน SaaS/API โดยไม่กำหนด prompt หรือ workflow เฉพาะ) ☐ RAG ☐ Rule/Logic ☐ Fine-tuning ☐ Custom-trained • Deployment Model: ☐ In-house ☒ SaaS ☐ Hybrid • Supporting Roles: ☐ Risk ☒ Compliance ☒ IT ☐ DPO	
3) Data, System & Influence Characteristics • ข้อมูลที่ใช้มีความอ่อนไหวหรือสามารถระบุตัวบุคคลได้หรือไม่ • AI มีอิทธิพลต่อการตัดสินใจมากกว่ามนุษย์มากน้อยเพียงใด • มีการพึ่งพาข้อมูลหรือโมเดลจากภายนอกหรือไม่	• AI Technology/Model Type Generative AI/LLM..... • Input Data Types (Internal/External/Market/Client) Public research documents/เอกสารวิเคราะห์จากแหล่งข้อมูลภายนอกที่ไม่ใช่ข้อมูลลับระดับสูง..... • PII/Sensitive Data: ☒ Low ☐ Medium ☐ High • Data Source: ☐ Internal ☒ External • Influence Level: ☒ Advisory ☐ Decision-Influencing ☐ Autonomous	• AIMS Controls - 3. ทรัพยากรสำหรับระบบ AI 5. วงจรชีวิตของระบบ AI 6. ข้อมูลสำหรับระบบ AI • AI Security Guideline สกมช. – หมวด Data Protection, Model Security & Supply chain Risk

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
4) Decision Path & Human Oversight • มนุษย์เข้ามาในขั้นตอนใดของการตัดสินใจ (ก่อน/หลังผลกระทบ) • มีอำนาจ override หรือหยุดการทำงานของ AI จริงหรือไม่ • Oversight ที่ออกแบบไว้เพียงพอหรือไม่เมื่อความเสี่ยงเพิ่ม	• Human-in-the-Loop Position: ☒ Pre ☐ Post ☐ Advisory • Human Authority: ☒ Review ☐ Override ☐ Approve only • Oversight Effectiveness: ☒ High ☐ Medium ☐ Low (อธิบายสั้น) High เพราะ ใช้ AI เป็นเครื่องมือช่วยสรุป ไม่ใช่แทนการวิเคราะห์ของมนุษย์.....	• AIMS Controls - 2. โครงสร้างภายในองค์กร 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI • AI Security Guideline สกมช. – หมวด Logging, Traceability & Accountability
5) Impact & Risk (เชื่อม 10 Rubrics) • ความเสี่ยงหลักเกิดจาก model, data, conduct หรือ systemic effect ไດ • มาตรการที่ใช้จัดการความเสี่ยงเพียงพอหรือไม่ อย่างไร • ใครเป็นผู้รับผิดชอบและยอมรับ residual risk	• Overall Risk Tier: ☒ Low ☐ Medium ☐ High • Primary Impact Areas: ☒ Operational Efficiency ☒ Reputational ☒ Privacy/Confidentiality ☐ Client/Investor ☐ Market Integrity ☐ Systemic/Contagion Risk • Key Risk Drivers (สรุป 2–3 bullet) ☒ R4 Model Risk: Hallucination/Accuracy ☒ R5 Data Risk: Confidentiality ☒ R6 Cyber/Adversarial Risk: Prompt Injection ☒ R7 Third Party: SaaS/API provider ☒ R9 Human Oversight/Automation Bias Risk • Primary Risk Owner Head of Research.....	• AIMS Controls - 4. การประเมินผลกระทบของระบบ AI 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI • AI Security Guideline สกมช. – หมวด Threat Modeling & Adversarial Risk

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	Residual Risk Acceptance Body (Committee)ผู้บริหารฝ่ายวิจัยร่วมกับ IT หรือหน่วยงานที่รับผิดชอบตามโครงสร้างภายใน.....	
6) Deployment, Monitoring & Resilience - หาก AI ทำงานผิดพลาด จะตรวจจับได้เร็วเพียงใด มีมาตรการบริหารจัดการเหตุการณ์ฉุกเฉินเพียงพอหรือไม่ - สามารถหยุดหรือถอยกลับไปใช้โมเดลสำรองได้หรือไม่ - ใครเป็นผู้ตัดสินใจเมื่อเกิด incident	- Key Controls in Place (High-level) Usage guideline, data input restriction, safe prompting practice, human review, access control, disclaimer, feedback channel..... - Review Frequency: ☐ Monthly ☐ Quarterly ☒ Annually ☒ Event-driven - Monitoring Focus: ☒ Hallucination/Inaccurate Output ☒ Data Leakage/Sensitive Data Exposure ☒ Prompt Injection/Misuse ☒ Complaints/User Feedback ☒ Third-party/Vendor Performance - Incident/Kill Switch: ☒ Yes ☐ No ☐ N/A - Fallback Processกลับไปอ่านและสรุปจากต้นฉบับโดยมนุษย์แทน..... - Third-party Provider (if any) OpenAI/ผู้ให้บริการ API ที่องค์กรใช้จริง..... - Key Controls (high level): Usage guideline, data input restriction, safe prompting practice, human review, access control, disclaimer, feedback channel - Review Frequency: Annually หรือเมื่อมีการเปลี่ยนโมเดล ผู้ให้บริการ ขอบเขตการใช้งาน หรือ major version update - Monitoring Focus: Hallucination feedback/ summary quality/inappropriate input	- AIMS Controls - 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI 8. ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก - AI Security Guideline สกมช. – หมวด Incident Response, Resilience & Recovery

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	- Incident/Suspension Capability: Yes - สามารถระงับสิทธิ์หรือปิดการเข้าถึง API ชั่วคราวเมื่อพบเหตุผิดปกติ - Fallback Process: กลับไปอ่านและสรุปจากต้นฉบับโดยมนุษย์ - Third-party Provider: OpenAI/ผู้ให้บริการ API ที่องค์กรใช้จริง	

หมายเหตุ

ตัวอย่างนี้แสดงการกรอก AI Inventory สำหรับ use case ความเสี่ยงต่ำในระดับที่กระชับ โดยยังคง 6 องค์ประกอบของ Detailed AI Inventory ครบถ้วน ทั้งนี้ การกรอกแบบ Lite ไม่ได้ลดความสำคัญของการบันทึก แต่ลดระดับความลึกของคำอธิบายให้เหมาะสมกับระดับความเสี่ยงของ use case หากลักษณะการใช้งานเปลี่ยนไป เช่น มีผลต่อผู้ลงทุน ใช้ข้อมูลอ่อนไหว หรือมีอิทธิพลต่อการตัดสินใจสำคัญ องค์กรควรขยายไปใช้การประเมินเชิงลึกและ controls ที่เข้มข้นขึ้นตาม AIMS

3. ตัวอย่าง AIMS Controls (Lite version)

วัตถุประสงค์ของส่วนนี้

ตาราง AIMS Controls ในเอกสารแบ่ง controls เป็น Mutual Controls และ Role-specific Controls สำหรับ Developer, User และ Customizer โดยครอบคลุม 8 ขั้นตอนหลักของ AIMS ได้แก่ นโยบาย โครงสร้าง องค์กร ทรัพยากร การประเมินผลกระทบ วงจรชีวิต AI ข้อมูล การใช้งาน และความสัมพันธ์ภายนอก

สำหรับ GenAI Summary Tool องค์กรมีบทบาทหลักเป็น User และมีลักษณะการใช้งานผ่าน SaaS/API จึงควรเลือก controls ขึ้นต่ำจากหมวดที่เกี่ยวข้อง ดังนี้

ตารางที่ 9 ตัวอย่างการเลือกมาตรการควบคุม AIMS Controls ของ GenAI Summary Tool

หมวด AIMS Controls	วัตถุประสงค์การเลือกใช้ มาตรการควบคุม สำหรับ GenAI Summary Tool	ประเภท Control	เหตุผลที่เลือกใช้	Minimum Evidence
1. นโยบาย AI	กำหนดว่าเครื่องมือนี้ใช้เพื่อช่วยสรุป และวิเคราะห์เบื้องต้นเท่านั้น ไม่ใช่ แพนบทวิเคราะห์สุดท้ายหรือ คำแนะนำการลงทุน	Mutual/ User-specific	ลดความเสี่ยงจากการใช้ AI เกินวัตถุประสงค์ และทำให้ ผู้ใช้เข้าใจขอบเขตการใช้งาน	Usage Guideline/ Acceptable Use Policy แบบสั้น
2. โครงสร้าง ภายในองค์กร	ระบุ Business Owner และ ผู้รับผิดชอบการใช้งาน เช่น ฝ่าย วิเคราะห์หลักทรัพย์ ร่วมกับ IT/ Compliance ในการกำกับดูแล	Mutual	ทำให้เกิด accountability และรู้ว่าใครเป็นผู้รับผิดชอบ เมื่อเกิดปัญหา	Owner record/ RACI แบบย่อ/ escalation contact
3. ทรัพยากร สำหรับระบบ AI	จำกัดสิทธิ์การใช้งานระบบหรือ API เฉพาะทีมที่เกี่ยวข้อง และควบคุม API key/account access	Mutual	ลด shadow AI การใช้งาน ผิดวัตถุประสงค์ และการ เข้าถึงระบบโดยผู้ไม่เกี่ยวข้อง	Access list/ Permission record

หมวด AIMS Controls	วัตถุประสงค์การเลือกใช้ มาตรการควบคุม สำหรับ GenAI Summary Tool	ประเภท Control	เหตุผลที่เลือกใช้	Minimum Evidence
4. การประเมินผล กระทบของระบบ AI	ประเมินแบบย่อว่าเป็น use case ภายใน ไม่ client-facing ใช้ข้อมูล ไม่อ่อนไหว และ AI มีระดับอิทธิพล เป็น Advisory	Mutual	ยืนยันว่า use case เข้าข่าย ความเสี่ยงต่ำและสามารถใช้ controls ขั้นต่ำได้อย่าง สมเหตุสมผล	AI Inventory/ risk tier note/ short assessment record
5. วงจรชีวิตของ ระบบ AI	ทบทวน use case อย่างน้อยปีละครั้ง หรือเมื่อเปลี่ยนโมเดล ผู้ให้บริการ ขอบเขตการใช้งาน หรือมี major version update จากผู้ให้บริการ	Mutual	ทำให้ governance ไม่หยุดที่ วันเริ่มใช้งาน และรองรับการ เปลี่ยนแปลงของ SaaS/LLM	Review record/ updated AI Inventory
6. ข้อมูลสำหรับ ระบบ AI	ห้ามป้อน MNPI ข้อมูลลูกค้า ข้อมูล ส่วนบุคคล หรือข้อมูลลับระดับสูง และตรวจสอบเงื่อนไข data opt-out/DPA/enterprise API terms	Mutual + User-specific	ป้องกัน data leakage, privacy risk, confidentiality risk และ การนำข้อมูลไปใช้ train model โดยไม่เหมาะสม	Data Input Rule/ AUP/DPA หรือ terms review note
7. การใช้งาน ระบบ AI	กำหนด safe prompting practices, human review, disclaimer และ feedback channel เช่น นักวิเคราะห์ต้อง ตรวจสอบผลสรุปกับต้นฉบับก่อนใช้อ้างอิง	User-specific	ลด hallucination, automation bias, prompt misuse และทำให้ผู้ใช้ ตระหนักถึงข้อจำกัดของ AI output	SOP/Safe Prompting Guide/ Review Log/ Disclaimer template/ Feedback log
8. ความสัมพันธ์ และการสื่อสารกับ บุคคลภายนอก	สอบถามผู้ให้บริการ SaaS/API เช่น privacy terms, SLA, data usage terms, incident notice และ ช่องทางติดต่อเมื่อเกิดปัญหา	User-specific	ลด third-party risk และ ทราบเงื่อนไขสำคัญของ ผู้ให้บริการที่อาจกระทบ การใช้งาน	Vendor record/ Terms of Service review/SLA note

หมายเหตุ

มาตรการควบคุมขั้นต่ำในตัวอย่างนี้ถูกเลือกจากข้อมูลที่บันทึกใน AI Inventory และจากตาราง AIMS Controls ตามบทบาทขององค์กรในฐานะ User และระดับความเสี่ยงของ use case ทั้งนี้ มาตรการดังกล่าวไม่ใช่รายการ controls ทั้งหมดของ AIMS แต่เป็นตัวอย่างการเลือก controls ในระดับที่ได้สัดส่วนกับ use case ความเสี่ยงต่ำ หากข้อมูล ลักษณะการใช้งาน ระดับอิทธิพลของ AI หรือผลกระทบต่อผู้ลงทุนเปลี่ยนแปลง องค์กรควรทบทวน Inventory และเลือก controls ที่เข้มข้นขึ้นตามความเหมาะสม

4. ตัวอย่าง Trigger: เมื่อใดต้องขยับจาก Lite ไป Full Assessment

เพื่อไม่ให้ตัวอย่างแบบย่อกลายเป็นการดำเนินการขั้นต่ำโดยอัตโนมัติ ผู้ประกอบธุรกิจสามารถวิเคราะห์และกำหนด trigger สำคัญที่จะต้องยกระดับไปประเมินในเชิงลึกต่อไป

ตารางที่ 10 ตัวอย่าง trigger หรือปัจจัยสำคัญที่ต้องพิจารณายกระดับไปประเมินเชิงลึก

Trigger	เหตุผลที่ต้องขยับไปประเมินเชิงลึก
นำผลสรุปไปส่งลูกค้าหรือผู้ลงทุน	เปลี่ยนจาก internal support เป็น investor-facing
ใช้ข้อมูลลูกค้า ข้อมูลส่วนบุคคล MNPI หรือข้อมูลลับ	ความเสี่ยงด้าน privacy, confidentiality และ conduct เพิ่มขึ้น
ใช้ผลลัพธ์ของ AI เป็นส่วนหนึ่งของคำแนะนำการลงทุน	Influence level ขยับจาก advisory ไป decision-influencing
เชื่อมต่อกับฐานข้อมูลภายในหรือระบบอัตโนมัติ	ความเสี่ยงด้าน access control, data leakage และ system dependency เพิ่มขึ้น
พบข้อผิดพลาดซ้ำ หรือมี complaint	risk profile เปลี่ยนและ controls เดิมอาจไม่เพียงพอ
เปลี่ยนโมเดล ผู้ให้บริการ หรือขอบเขตการใช้งาน	ต้องประเมิน risk ใหม่ตาม lifecycle governance

หมายเหตุ

ตัวอย่างการประยุกต์ใช้เครื่องมือในส่วนของภาคผนวกนี้แสดงให้เห็นว่า use case ความเสี่ยงต่ำสามารถใช้กรอบการกำกับดูแลในระดับที่กระชับได้ โดยยังคงหลัก Think > Record > Act ครบถ้วน กล่าวคือ องค์กรคิดผ่านประเด็นสำคัญตาม 4 ะยะบันทึกข้อมูลลง AI Inventory ตาม template ปกติ และเลือกมาตรการควบคุมขั้นต่ำจาก AIMS Controls ตามบทบาทและระดับความเสี่ยงของตน ทั้งนี้ หากลักษณะการใช้งานหรือผลกระทบเปลี่ยนแปลง องค์กรควรขยับไปใช้การประเมินเชิงลึกและ controls ที่เข้มข้นขึ้นตาม AIMS

ภาคผนวก 7 ตัวอย่างการประยุกต์ใช้เครื่องมือกรณี Use Case เสี่ยงสูง - Automated Trading Support

นิยาม: ระบบ Agentic AI ที่ช่วยวิเคราะห์ข้อมูลตลาดและสัญญาณที่กำหนดไว้ล่วงหน้า เพื่อเสนอคำสั่งซื้อขายหรือดำเนินการตาม workflow ภายใต้ข้อจำกัดที่องค์กรกำหนด เช่น trading limit, approved instruments, pre-trade control, human-in/on-the-loop monitoring และ kill switch ทั้งนี้ ระบบไม่ควรถูกออกแบบให้ดำเนินการนอกขอบเขตที่ได้รับอนุมัติหรือเปลี่ยนแปลงกลยุทธ์โดยไม่มีการกำกับดูแล และต้องครอบคลุมการควบคุมในระดับก่อนที่ระบบจะดำเนินการแต่ละครั้ง (runtime governance) ด้วย โดยผู้ประกอบธุรกิจควรพิจารณาให้มีองค์ประกอบสำคัญ 4 ด้าน ได้แก่ (1) กลไกยืนยันตัวตนของระบบ AI (Agent Identity) ที่ผูกกับผู้รับผิดชอบและขอบเขตอำนาจดำเนินการอย่างชัดเจน (2) การบันทึกข้อมูลประกอบก่อนดำเนินการ (Pre-execution Record) ที่ครอบคลุมทั้ง action ที่จะทำ ขั้นตอนการวิเคราะห์ เครื่องมือและข้อมูลที่ใช้พร้อม context ที่เกี่ยวข้อง (3) การตัดสินใจก่อนดำเนินการ (Pre-execution Disposition) ที่ให้ผลลัพธ์ชัดเจน เช่น ดำเนินการอัตโนมัติ ฝ่าฝืนติดตาม ส่งให้มนุษย์อนุมัติ ปฏิเสธ ตามระดับความเสี่ยงของแต่ละ action (4) การบันทึกเหตุการณ์ที่ตรวจสอบย้อนหลังได้ (Runtime Audit Log) ในรูปแบบที่แก้ไขไม่ได้ (append-only) เพื่อรองรับ compliance, audit และ review ทั้งนี้ ในกรณี workflow ที่มีหลายขั้นตอน การอนุมัติหรือผลการตัดสินใจในขั้นตอนหนึ่งไม่ถือเป็น การอนุมัติขั้นตอนถัดไปโดยอัตโนมัติ ระบบต้องผ่านการตรวจสอบใหม่ในทุก action

สำหรับระบบ AI ที่มีลักษณะ Agentic AI องค์กรควรพิจารณาความสามารถในการดำเนินการ (action capability) ขอบเขตอำนาจในการดำเนินการ (action authority) ระดับความเป็นอิสระ (autonomy level) และสิทธิ์การเข้าถึงของ agent (agent identity and privileges) เพื่อให้มั่นใจว่าระบบยังคงอยู่ภายใต้ขอบเขตการกำกับดูแล การตรวจสอบย้อนหลัง และการควบคุมโดยมนุษย์ตามระดับความเสี่ยงและผลกระทบของ use case ซึ่งแบ่งลักษณะได้ดังนี้

ตารางที่ 11 ตัวอย่างมิติคุณลักษณะของระบบ Agentic AI ที่ต้องพิจารณาเพิ่มเติม

มิติคุณลักษณะ	รายละเอียดที่ต้องพิจารณา	ที่มาและความสำคัญ
1. Agentic Capability (AI นี้ “มีความสามารถอะไร” และ “สามารถลงมือทำอะไรได้บ้าง”)	☒ Tool use - ระบบสามารถเรียกใช้เครื่องมือ/API ภายนอก ☒ Market data retrieval - ดึงข้อมูลตลาด ☒ Multi-step signal evaluation - วิเคราะห์หลายขั้นตอน ☒ Internal trading workflow integration - เชื่อมกับ workflow ภายใน ☒ Order proposal - เสนอคำสั่งซื้อขาย ☒ Bounded execution - สามารถ execute ภายใต้ขอบเขตที่กำหนด ☒ Pre-execution declaration - ระบบต้องรวม package action + reasoning + tool call + data source ก่อน execute ทุกครั้ง	ใช้เพื่อประเมินว่า AI เป็นเพียง “ระบบช่วยวิเคราะห์” หรือเริ่มมี “ความสามารถในการลงมือทำ (agency)” ซึ่งส่งผลต่อระดับ oversight, accountability และ controls ที่ต้องใช้ โดยแนวคิดนี้สอดคล้องกับแนวทาง Infocomm Media Development Authority (IMDA) และ IBM ที่เน้นการพิจารณา tool access, planning capability และ execution authority ของ Agentic AI และ SAFR ที่กำหนดให้ agent ต้องส่ง Governance Envelope ก่อนดำเนินการทุกครั้ง

มิติคุณลักษณะ	รายละเอียดที่ต้องพิจารณา	ที่มาและความสำคัญ
2. Action-space (AI นี้ “ได้รับอนุญาตให้ทำอะไรได้จริง” ภายใต้ governance boundary)	☒ Retrieve approved market data - AI สามารถดึงข้อมูลสภาพตลาด ☒ Evaluate predefined signals - วิเคราะห์สัญญาณที่กำหนดล่วงหน้าได้ ☒ Generate order proposal - เสนอคำสั่งได้ ☒ Route order through pre-trade controls - ส่งคำสั่งผ่านจุดตรวจสอบและเกณฑ์ที่กำหนด ☒ Execute only within approved hard limits, if permitted - execute ได้เฉพาะภายใต้ hard limits และ pre-trade controls ที่กำหนดไว้ล่วงหน้า ☒ Machine-readable mandate - ขอบเขตอำนาจถูกกำหนดในรูปแบบที่ระบบตรวจสอบได้อัตโนมัติ (เช่น structured policy, signed mandate) ☒ Multi-step independence - การอนุมัติในขั้นตอนหนึ่งไม่ carry-over ไปขั้นถัดไป	เป็นหัวใจของ governance สำหรับ Agentic AI เพราะ “สิ่งที่ AI ทำได้” สำคัญกว่าความฉลาดของโมเดลเพียงอย่างเดียว แนวคิด action-space bounding นี้ สอดคล้องกับหลัก least privilege, bounded autonomy และ safe delegation ในแนวทาง Agentic Governance ระดับสากล เช่น capability-based security และ AP2 mandate ของ SAFR ที่ให้ mandate อยู่ในรูปแบบ machine-verifiable
3. Actions Not Permitted (ทำในสิ่งที่ห้ามหรือไม่)	☐ Trade unapproved instruments - ซื้อขายหลักทรัพย์ที่ไม่ได้รับอนุญาต ☐ Override trading limits - มีสิทธิ์ข้ามข้อจำกัดด้านการซื้อขายที่กำหนดไว้ ☐ Modify model or strategy by itself - แก้ไข strategy เองได้ ☐ Execute outside approved mandate - ทำรายการนอก mandate ☐ Self-extend delegated authority - ระบบขยาย mandate ตัวเองผ่านการอนุมานหรือ chain-of-tools ☐ Bypass pre-execution checkpoint - ข้ามจุดตรวจสอบก่อนส่งคำสั่ง	การกำหนด “ข้อห้าม” ช่วยสร้าง governance boundary ที่ชัดเจน และลด risk จาก autonomy creep หรือ privilege escalation โดยเฉพาะในระบบที่มีผลกระทบต่อการตลาดทุน สอดคล้องกับ SAFR ที่ระบุว่า "an agent cannot extend the scope of a mandate through its own reasoning or inference"
4. Autonomy Level and Pre-execution disposition (AI นี้ “มีอิสระในการตัดสินใจและลงมือทำมากเพียงใด”)	A. Autonomy Level (ระดับความอิสระของระบบโดยรวม) ☒ Human-on-the-loop with real-time monitoring - มีมนุษย์ติดตามแบบ real-time และสามารถระงับการทำงานได้ ☒ Bounded autonomous execution within hard limits - AI สามารถ execute ภายใต้กรอบที่กำหนดได้ ☐ Fully autonomous execution without monitoring - ตัดสินใจอัตโนมัติ ไม่มีการติดตาม	ใช้เพื่อประเมิน “คุณภาพของ human oversight” ไม่ใช่เพียงมีมนุษย์อยู่ในกระบวนการหรือไม่ โดยสอดคล้องกับหลัก Human Oversight ของ OECD, NIST AI RMF และแนวคิด human governance และ Disposition Engine 4 outcomes ของ SAFR ที่กำหนดผลลัพธ์การตัดสินใจก่อนดำเนินการอย่างชัดเจนและผูกพัน

มิติคุณลักษณะ	รายละเอียดที่ต้องพิจารณา	ที่มาและความสำคัญ
	B. Pre-execution Disposition (การตัดสินใจก่อนดำเนินการต่อ action) ☒ Auto-Execute - action ที่อยู่ในกรอบ hard limit และผ่าน pre-trade control ให้ดำเนินการอัตโนมัติ ☒ Escalate - action ที่เกิน threshold ที่กำหนด ต้องส่งให้มนุษย์อนุมัติก่อน ☒ Deny - action ที่ละเมิด policy/mandate/hard limit ต้องถูกปฏิเสธก่อนถึงตลาด ☐ Observe - action ที่มี novelty หรือ anomaly signal ให้ดำเนินการแต่ log เพื่อเฝ้าติดตาม (ทางเลือก)	
5. Agent Identity (AI นี้ “ใช้ตัวตนและสิทธิ์การเข้าถึงอย่างไร”)	☒ Dedicated agent identity - Agent มี account เฉพาะ ☒ Least-privilege access - สิทธิ์จำกัดตามหน้าที่ ☒ Non-transferable credentials - ข้อมูลยืนยันตัวตน (credentials) ของ AI agent ห้ามใช้ร่วมกัน, ห้ามส่งต่อ, และต้องผูกกับตัวตนเฉพาะ ☒ Tool-call and order-action logging ใช้ credential ที่ตรวจสอบได้ และมี log ทุก action ☒ Identity verified before every action - ระบบตรวจสอบตัวตนก่อน execute ทุกครั้ง	เป็นหลัก traceability และ accountability เพื่อให้สามารถตรวจสอบย้อนหลังได้ว่า “AI ตัวใด ทำอะไร ผ่านเครื่องมือใด และเมื่อใด” ซึ่งสอดคล้องกับแนวคิด auditability และ non-repudiation ใน AI governance และ cybersecurity
6. Reversibility of Actions (ความสามารถในการย้อนกลับ ยกเลิก หรือจำกัดผลกระทบ)	คำสั่งบางประเภทเมื่อ execute แล้วอาจส่งผลกระทบต่อตลาดและไม่สามารถย้อนกลับได้ทันที จึงต้องมีการกำหนดขีดจำกัดไว้ล่วงหน้า (hard limit), pre-trade control, kill switch และ surveillance นอกจากนี้ อาจต้องมี Tamper-evident action log - บันทึกทุก action ในรูปแบบ append-only ที่แก้ไขไม่ได้ เพื่อ reconstruct เหตุการณ์ได้เสมอ	เป็นมิติสำคัญของ Proportionality เพราะยิ่ง action “ย้อนกลับไม่ได้” มาก governance และ resilience controls ต้องเข้มข้น โดยเชื่อมกับ operational resilience และ market integrity
7. Risk Tier (ระดับความเสี่ยง)	High Risk เนื่องจาก use case นี้ มีผลกระทบต่อตลาดการซื้อขาย และ operational resilience	Risk Tier ใช้ routing ไปยัง AIMS controls, การยกระดับการกำกับดูแล และความถี่การตรวจสอบ ที่เหมาะสม เช่น real-time monitoring, board visibility หรือ enhanced incident response

หมายเหตุ

สำหรับ Agentic AI ความเสี่ยงไม่ได้อยู่ที่ output เพียงอย่างเดียว แต่อยู่ที่ action-space, autonomy, tool access, agent identity และความสามารถดำเนินการจริง ซึ่งเป็นมิติที่ระดับสากลให้ความสำคัญเพิ่มจาก AI/GenAI ทั่วไป

1. ตัวอย่าง การประเมิน 4 ระยะ 12 ขั้นตอน

ตารางที่ 12 ตัวอย่างการประเมินกระบวนการ 12 ขั้นตอนของ Automated Trading Support

ขั้นตอน	รายละเอียดวิธีการ
ระยะที่ 1: การริเริ่มและการพิจารณาความจำเป็น (Phase 1: Concept & Necessity Test)	
ขั้นตอนที่ 1 การริเริ่มโครงการ AI (Initiation)	ฝ่ายลงทุนหรือฝ่ายซื้อขายเสนอใช้ Agentic AI เพื่อช่วยวิเคราะห์ข้อมูลตลาด ประเมินสัญญาณตามเงื่อนไขที่กำหนด และสนับสนุน workflow การส่งคำสั่งซื้อขายภายใต้กรอบควบคุมที่องค์กรอนุมัติ
ขั้นตอนที่ 2 การทดสอบความจำเป็น และความเหมาะสม (Necessity & Suitability Test)	มีความจำเป็นด้าน speed, consistency, real-time response และ operational efficiency อย่างไรก็ดี เนื่องจาก use case มีผลต่อการซื้อขายและอาจกระทบ market integrity จึงต้องพิสูจน์ว่า AI เพิ่มประสิทธิภาพโดยไม่สร้างความเสี่ยงต่อราคา สภาพคล่อง ผู้ลงทุน หรือเสถียรภาพของตลาดเกินกว่าระดับที่ยอมรับได้
ขั้นตอนที่ 3 การบันทึกในบัญชีรายชื่อ ปัญหาประดิษฐ์ (Initial Inventory Logging)	ลงทะเบียนเป็น High-risk/Market-impact/Agentic AI use case ใน AI Inventory เนื่องจากเกี่ยวข้องกับการวิเคราะห์สัญญาณตลาด การเสนอหรือส่งคำสั่งซื้อขาย การเชื่อมต่อบริษัทสำคัญ และการมี action-space ที่อาจก่อผลกระทบต่อกลไกตลาด และพิจารณาเพิ่มเติมว่าระบบมีลักษณะ agent ที่สามารถริเริ่มการดำเนินการเองหรือไม่ หากใช่ ให้ระบุขอบเขตอำนาจ (delegated authority) เช่น ประเภทคำสั่ง มูลค่าสูงสุด ความถี่สูงสุด และเงื่อนไขที่ต้อง escalate ให้มนุษย์อนุมัติ
ระยะที่ 2: ความสมเหตุสมผลและการจัดระดับความเสี่ยง (Phase 2: Proportionality & Risk Tiering)	
ขั้นตอนที่ 4 การคัดกรองความเสี่ยง (Inherent Risk Screening)	ความเสี่ยงโดยธรรมชาติสูง เนื่องจากเกี่ยวข้องกับการ market impact, erroneous order, model/data risk, data latency, cyber/adversarial risk, third-party dependency, automation bias, explainability gap และ regulatory exposure
ขั้นตอนที่ 5 การจัดระดับความเสี่ยง (Tier Classification)	จัดเป็น High Risk เนื่องจากระบบอาจมีผลต่อคำสั่งซื้อขาย ราคา สภาพคล่อง market integrity และความเชื่อมั่นของตลาด จึงต้องใช้การประเมินแบบเต็มและ controls ที่เข้มข้นกว่า [เพิ่มเติมกรณี Agentic AI] ต้องประเมินแยกจากระบบทำได้เพียงเสนอคำสั่ง หรือสามารถส่งคำสั่งภายใต้การกำหนดขีดจำกัดไว้ล่วงหน้า (hard limits) ได้ รวมถึงต้องระบุว่า action ใดต้อง human approval, action ใด execute ได้อัตโนมัติ และ action ใดถูกห้ามโดยเด็ดขาด [เพิ่มเติมกรณี Agentic AI] การส่งคำสั่งซื้อขายอาจเป็น action ที่ย้อนกลับไม่ได้หรือสร้างผลกระทบต่อตลาดทันที จึงต้องมีมาตรการควบคุมที่รัดกุมและเพียงพอด้วย
ระยะที่ 3: การประเมินเชิงลึกและการกำหนดมาตรการควบคุม (Phase 3: Deep Assessment & AIMS Control)	
ขั้นตอนที่ 6 การประเมินผลกระทบ อย่างรอบด้าน (Full Impact Assessment)	ประเมินครบทุกมิติของ 10 Risk Rubrics ได้แก่ market impact, investor harm, systemic/contagion, model risk, data risk, cyber/adversarial risk, third party and concentration risk, explainability gap, human oversight risk และ regulatory/conduct exposure

ขั้นตอน	รายละเอียดวิธีการ
ขั้นตอนที่ 7 การเลือกใช้มาตรการ ควบคุมตามระดับ ความเสี่ยง (Control Selection)	เลือก controls ตามระดับความเสี่ยงและลักษณะการใช้งาน AI ภายใต้แนวคิด AIMS โดยสำหรับ use case ที่มี action capability หรือมีผลกระทบต่อการซื้อขาย ควรให้ความสำคัญกับการควบคุมที่เข้มงวด (hard controls) และการควบคุมแบบเรียลไทม์ (real-time controls) เช่น Agent Identity เพื่อใช้ระบุ AI ตัวใดเป็นผู้ดำเนินการ, Access Control ว่ามีการจำกัดสิทธิ์ตามหลัก least privilege, การทำ Model Validation เพื่อใช้ประเมินความถูกต้อง ความเหมาะสม และข้อจำกัดของโมเดล, Tool-call Logging บันทึกการเรียกใช้เครื่องมือ เพื่อใช้ตรวจสอบย้อนหลังได้ว่า AI ใช้ tool ไหน เมื่อใด และเพื่ออะไร ทั้งนี้ สามารถศึกษามาตรการควบคุมเพิ่มเติมได้จาก IMDA (2026), IBM (2026)
ขั้นตอนที่ 8 การประเมินความเสี่ยง คงเหลือ (Residual Risk Evaluation)	ความเสี่ยงคงเหลืออาจยังอยู่ในระดับ medium-high แม้มี controls แล้ว เนื่องจากตลาดเปลี่ยนแปลงเร็วและระบบมี action capability การยอมรับ residual risk ต้องทำโดย governance body ระดับสูง เช่น Risk Committee/Investment Committee/Senior Management ตามโครงสร้างองค์กร [เพิ่มเติมกรณี Agentic AI] ไม่ควรเป็น human-in-the-loop ในเชิงรูปแบบ แต่ต้องวัดประสิทธิภาพของ oversight เช่น override rate, escalation rate, approval latency, kill-switch activation ที่เกี่ยวข้องกับการรับมือเหตุการณ์เบื้องต้น เป็นต้น และ post-approval incident rate เพื่อป้องกันการอนุมัติโดยขาดการตรวจสอบอย่างเพียงพอ (rubber-stamping) หรือการเชื่อ AI มากเกินไป จนละเลยพิจารณาของมนุษย์ (automation bias) [เพิ่มเติมกรณี Agentic AI] ต้องระบุ agent identity แยกจาก human user/service account และเก็บ log ว่า agent ใด เรียก tool ใด เสนอ action ใด ดำเนินการในนามใคร และภายใต้ permission ใด เพื่อให้ audit ได้ [เพิ่มเติมกรณี Agentic AI] ควรกำหนดกลไก runtime check ก่อนระบบส่งคำสั่งจริง โดยระบุผลลัพธ์การตัดสินใจที่เป็นไปได้ (เช่น ดำเนินการอัตโนมัติ ฝ่าฝืนติดตาม ส่งให้มนุษย์อนุมัติหรือปฏิเสธ) และเงื่อนไข kill switch ที่ระงับการทำงานของระบบทันที
ระยะที่ 4: การอนุมัติ การนำไปใช้ และการติดตามผล (Phase 4: Approval, Deployment & Monitoring)	
ขั้นตอนที่ 9 การอนุมัติอย่างเป็นทางการ (Formal Sign-off)	ควรมี formal sign-off จาก Business Owner, Risk, Compliance, IT/Cybersecurity, Model Risk หรือคณะกรรมการที่เกี่ยวข้อง โดยพิจารณาจาก risk tier, action-space, autonomy, controls, residual risk และ fallback plan
ขั้นตอนที่ 10 การนำไปใช้ภายใต้หลักการ พิจารณาอย่างรอบคอบ (Deployment with Safe Harbor)	ก่อน deploy ควรมีหลักฐานการทดสอบความพร้อมก่อนใช้งาน เช่น AI Inventory, model/system documentation, test results, stress testing, pre-trade control evidence, agent identity record, access list, approved data source/instrument list, rollback/fallback plan และ incident playbook
ขั้นตอนที่ 11 การติดตามเฝ้าระวังอย่าง ต่อเนื่อง (Dynamic Monitoring)	ติดตามแบบ real-time หรือ near real-time ในมิติ order anomalies, limit breach, abnormal trading pattern, data latency, model drift, tool-call errors, override rate, kill-switch activation, post-trade exceptions และ market abuse indicators

ขั้นตอน	รายละเอียดวิธีการ
ขั้นตอนที่ 12 การทบทวนเป็นระยะ (Periodic Review)	ทบทวนอย่างน้อยรายไตรมาส หรือ event-driven เมื่อเปลี่ยน model, data source, trading scope, tool connector, agent permission, market condition, dependency หรือเกิด incident [เพิ่มเติมกรณี Agentic AI] ต้องทบทวนเพิ่มเติมเมื่อมีการเพิ่ม tool ใหม่ เปลี่ยน autonomy level เพิ่ม write/execute permission เพิ่ม memory capability หรือให้ทำงานร่วมกับ agent อื่น

1.1 ตัวอย่าง Risk Rubrics — ตัวอย่างการกำหนดระดับคะแนน

ตัวอย่างนี้เป็นการสาธิตวิธีการให้คะแนนเท่านั้น ผู้ประกอบธุรกิจสามารถกำหนดคะแนนตามบริบทจริง ปริมาณการซื้อขาย ผลกระทบต่อผู้ลงทุน ขอบเขตอำนาจของ agent และ controls ที่มีอยู่

ตารางที่ 13 ตัวอย่างการประเมินความเสี่ยงตาม Risk Rubrics ของ Automated Trading Support

Risk Rubric	คะแนน	เหตุผล
R1 Market Impact/Market Integrity	5	เกี่ยวข้องกับคำสั่งซื้อขายและอาจกระทบราคา/สภาพคล่อง
R2 Investor Harm/Client Impact	3–4	อาจกระทบผู้ลงทุนทางอ้อมผ่านราคา สภาพคล่อง หรือ execution quality
R3 Systemic/Contagion	4	หากใช้ในวงกว้างหรือหลายระบบมี behavior คล้ายกัน อาจเกิด pro-cyclicality
R4 Model Risk	4	มีความเสี่ยงจาก drift, erroneous signal, hallucination ใน agent reasoning หรือ model failure
R5 Data Risk	4	data latency, data quality, confidentiality ของ strategy/limits สำคัญมาก
R6 Cyber/Adversarial	5	มีความเสี่ยงจาก adversarial attack, API compromise, tool misuse, indirect prompt injection
R7 Third Party and Concentration Risk	4	อาจพึ่งพา market data, model provider, cloud/API หรือ trading infrastructure
R8 Explainability Gap	4	ต้องอธิบาย decision path, tool call, signal evaluation และ action route ได้
R9 Human Oversight/Automation Bias	5	human-on-the-loop อาจกลายเป็น rubber-stamping หากไม่มี metrics
R10 Regulatory/Conduct Exposure	5	เกี่ยวข้องกับ market conduct, record-keeping, best execution และ supervisory scrutiny

Overall Risk Tier: High

2. ตัวอย่าง Detailed AI Inventory (Use Case Detail)

กรณีศึกษา: Agentic AI for Automated Trading Support

Use Case ID: AI-UC-TRD-AGT-01

ชื่อ Use Case: Agentic AI Trading Support System

หน่วยงานเจ้าของ: Investment/Trading Desk

วันที่จัดทำ/ทบทวน:

Overall Risk Tier: High

หมายเหตุ: ตารางนี้ใช้ template เดียวกับ AI Inventory แบบ full record สำหรับ High-risk Use Case โดยข้อมูลในส่วน Agentic AI ให้กรอกเฉพาะกรณีที่ระบบมีความสามารถในการใช้เครื่องมือ เชื่อมต่อระบบ หรือดำเนินการแทนมนุษย์

ตารางที่ 14 ตัวอย่างการจัดทำ AI Inventory ของ Automated Trading Support

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
1) Context & Strategic Importance - AI นี้ถูกใช้เพื่อบรรลุวัตถุประสงค์ใดบริบทไหน - ผลลัพธ์ของ AI ส่งผลกระทบต่อใครโดยตรงบ้าง (ผู้ลงทุน/ตลาด/องค์กร) - การใช้งานนี้มีนัยเชิง fiduciary หรือ conduct risk หรือไม่	- Purpose/วัตถุประสงค์สนับสนุนการวิเคราะห์ข้อมูลตลาดและ workflow การส่งคำสั่งซื้อขายภายใต้ขอบเขตที่กำหนด..... - Scope/Use Context ...ใช้โดยฝ่ายลงทุนหรือฝ่ายซื้อขายสำหรับ approved instruments และ approved strategies เท่านั้น.... - วัตถุประสงค์การใช้งานในตลาดทุน: ☒ Investment & Trading ☐ Advice & Services ☐ Risk & Compliance ☐ Internal Support - Client-Facing/Investor Impact: ☐ Yes ☒ No แต่อาจกระทบต่อสภาพตลาดโดยรวม	- AIMS Controls - - นโยบาย AI - โครงสร้างภายในองค์กร - การประเมินผลกระทบของระบบ AI
2) Role & Development Responsibility - องค์กรมีบทบาทเพียงผู้ใช้ หรือเข้าไปกำหนดพฤติกรรมระบบ AI - ระดับการพัฒนาส่งเสริมให้เกิดความรับผิดชอบ developer/ customizer เพียงใด	- Organizational AI Role: ☐ ระบบ AI ที่องค์กรพัฒนาเอง (developer) ☐ ระบบ AI ที่องค์กรนำมาใช้ (user) ☒ ระบบงาน AI แบบผสมผสานที่องค์กรใช้ software AI แต่มีการ customized เพื่อให้เหมาะสมกับบริบทหรือเนื้อหาภายในองค์กร (customizer) - Development Characteristics: ☐ None ☒ Prompt Eng./Orchestration ☐ RAG	- AIMS Controls - - โครงสร้างภายในองค์กร - ทรัพยากรสำหรับระบบ AI - วงจรชีวิตของระบบ AI

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	☒ Rule/Logic ☐ Fine-tuning ☐ Custom-trained ☐ Other: _____ โดยกรณี Agentic AI จำเป็นต้องพิจารณามุมมองด้านการพัฒนาอื่นประกอบด้วย ดังนี้ ☒ Tool Orchestration - การเชื่อม AI เข้ากับ tools หรือ workflow ต่าง ๆ ☒ Model Integration - การเชื่อมหลาย model หรือหลายบริการ AI เข้าด้วยกัน ☒ API Integration - การเชื่อม API หรือระบบภายนอก ☒ Multi-agent Workflow - มี AI หลาย agent ทำงานร่วมกัน ☐ Persistent Memory - มี memory หรือ context ข้าม session ☐ Self-modifying Capability - ระบบสามารถปรับเปลี่ยนพฤติกรรม/logic เอง • Deployment Model: ☐ In-house ☒ SaaS ☐ Hybrid ☒ Internal System Integration - เชื่อมกับระบบภายในองค์กร ☒ Market Data/External Data Tools - เชื่อมแหล่งข้อมูลตลาดหรือข้อมูลภายนอก • Supporting Roles: ☒ Risk ☒ Compliance ☒ IT ☒ Cybersecurity ☒ Internal Audit ☒ DPO	

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	Accountability ...ต้องระบุผู้รับผิดชอบ business decision, model/system risk, trading control และ kill-switch authority อย่างชัดเจน	
3) Data, System & Influence Characteristics - ข้อมูลที่ใช้มีความอ่อนไหวหรือสามารถระบุตัวบุคคลได้หรือไม่ - AI มีอิทธิพลต่อการตัดสินใจมากกว่ามนุษย์มากน้อยเพียงใด - มีการพึ่งพาข้อมูลหรือโมเดลจากภายนอกหรือไม่	- AI Technology/Model Type LLM-based agent + quantitative signal tools/ market data tools..... - Input Data Types (Internal/External/Market/Client) market data, order book data, approved signals, internal risk limits, historical trading data, approved market news/data source..... - PII/Sensitive Data: ☐ Low ☐ Medium ☒ High (กรณีข้อมูลเกี่ยวกับกลยุทธ์และ internal limits และ trading parameters) - Data Source: ☒ Internal ☒ External - Influence Level: ☐ Advisory ☒ Decision-Influencing with bounded autonomous ☐ Autonomous นอกจากนี้ กรณี Agentic AI อาจจะต้องพิจารณาอำนาจการตัดสินใจในมิติต่าง ๆ ตามที่ระบุไว้ก่อนหน้านี้ - Agentic Capability [เพิ่มเติมสำหรับ Agentic AI] ☒ Tool use ☒ Market data retrieval ☒ Multi-step signal evaluation ☒ Order proposal ☒ Bounded execution ☒ Pre-execution declaration - Action-space [เพิ่มเติมสำหรับ Agentic AI] ☒ Retrieve approved data ☒ Evaluate predefined signals ☒ Generate order proposal	- AIMS Controls - 3. ทรัพยากรสำหรับระบบ AI 5. วงจรชีวิตของระบบ AI 6. ข้อมูลสำหรับระบบ AI - IMDA, IBM – Agentic AI ต้องเพิ่มมิติ action-space, autonomy, tools, memory/contexts และ agent identity เนื่องจาก agent สามารถลงมือกระทำในระบบได้มากกว่า GenAI ทั่วไป

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	☒ Route order through pre-trade checks ☒ Execute only within approved hard limits ☒ Machine-readable mandate ☒ Multi-step independence • Actions Not Permitted ☐ Trade unapproved instruments ☐ Override trading limits ☐ Modify strategy by itself ☐ Execute outside approved mandate ☐ Self-extend delegated authority ☐ Bypass pre-execution checkpoint • Autonomy Level and Pre-execution Disposition [เพิ่มเติมสำหรับ Agentic AI] A. Autonomy Level ☒ Human-on-the-loop with real-time monitoring ☒ Bounded autonomous execution within hard limits ☐ Fully autonomous execution without monitoring B. Pre-execution Disposition ☒ Auto-Execute ☒ Escalate ☒ Deny ☐ Observe (ทางเลือก) • Agent Identity [เพิ่มเติมสำหรับ Agentic AI] ☒ Dedicated agent identity ☒ Least-privilege access ☒ Non-transferable credentials ☒ Tool-call and order-action logging ☒ Identity verified before every action	
4) Decision Path & Human Oversight • มนุษย์เข้ามาในขั้นตอนใดของ	• Human-in-the-Loop Position: ☐ Pre ☐ Post	• AIMS Controls - 2. โครงสร้างภายในองค์กร 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
การตัดสินใจ (ก่อน/หลัง ผลกระทบ) • มีอำนาจ override หรือหยุดการทำงานของ AI จริงหรือไม่ • Oversight ที่ออกแบบไว้เพียงพอหรือไม่เมื่อความเสี่ยงเพิ่ม	☐ Advisory กรณี Agentic AI อาจจะต้องพิจารณา ☒ Human-on-the-loop สำหรับ real-time monitoring ☒ Human approval สำหรับ strategy change, limit change, new instrument approval หรือ exceptional order • Human Authority: ☐ Review ☒ Override with suspend, activate kill switch, escalate incident และสิ่ง fallback ไป manual workflow ☐ Approve only • Oversight Effectiveness: ☒ High และต้องวัดผลได้ ไม่ใช่ระบุว่ามีอาการกำกวม และมีการดูแลอนุมัติหรือตรวจสอบคำแนะนำของ AI อย่างถี่ถ้วน ☐ Medium ☐ Low (อธิบายสั้น) 	• IBM - เน้นย้ำ Agentic AI ต้องเน้นการควบคุมและติดตาม action ไม่ใช่ตรวจเพียงคำตอบ ดังนั้น human oversight ต้องมี metrics เพื่อวัด effectiveness AI Security Guideline สกมช. – หมวด Logging, Traceability & Accountability
5) Impact & Risk (เชื่อม 10 Rubrics) • ใครได้รับผลกระทบจากการใช้งาน AI บ้าง • ความเสี่ยงหลักเกิดจาก model, data, conduct หรือ systemic effect ไດ • มาตรการที่ใช้จัดการความเสี่ยงเพียงพอหรือไม่ อย่างไร • ใครเป็นผู้รับผิดชอบและยอมรับ residual risks	• Overall Risk Tier: ☐ Low ☐ Medium ☒ High • Primary Impact Areas: ☐ Client/Investor ☒ Market Integrity ☐ Conduct/Fair Treatment ☐ Operational Efficiency ☒ Operational Resilience/Business Continuity ☒ Reputational ☒ Regulatory/Legal Compliance ☐ Privacy/Confidentiality ☐ Financial/Revenue Impact ☒ Systemic/Contagion Risk ☒ Third-party/Outsourcing Dependency ☐ Other: _____	• AIMS Controls - 4. การประเมินผลกระทบของระบบ AI 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI • IMDA, IBM – Agentic AI อาจสร้างความเสี่ยงใหม่จาก autonomous actions, tool use, open-endedness และ non-reversible actions ซึ่งต้องถูกประเมินในระดับ use case

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	• Key Risk Drivers (ศึกษา Risk Rubrics เพิ่มเติม และจากผลการประเมินก่อนหน้า) ☒ R1 Market Impact/Market Integrity Risk ☒ R2 Investor Harm/Client Impact Risk ☒ R3 Systemic/Contagion Risk ☒ R4 Model Risk: Bias/Drift/Hallucination/Accuracy ☒ R5 Data Risk: Quality/Privacy/Confidentiality/IP ☒ R6 Cyber/Adversarial Risk ☒ R7 Third Party and Concentration Risk ☒ R8 Explainability/Transparency Gap ☒ R9 Human Oversight/Automation Bias Risk ☒ R10 Regulatory/Legal/Conduct Exposure • Primary Risk Owner Head of Trading/Head of Investment..... • Residual Risk Acceptance Body (Committee) BoD..... • Residual Risk Position อาจยังอยู่ระดับ Medium-high แม้มี controls เพราะเกี่ยวข้องกับการซื้อขายและพฤติกรรมตลาด.....	
6) Deployment, Monitoring & Resilience • หาก AI ทำงานผิดพลาด จะตรวจจับได้เร็วเพียงใด • มีมาตรการบริหารจัดการเหตุการณ์ฉุกเฉินเพียงพอหรือไม่ • สามารถหยุดหรือถอยกลับไปใช้โมเดลสำรอง (fallback) ได้หรือไม่ • ใครเป็นผู้ตัดสินใจเมื่อเกิด incident	• Key Controls in Place (High-level) pre-trade controls, approved instruments, hard trading thresholds, kill switch, post-trade surveillance, model validation, stress testing, agent identity, tool-call logging, order-action logging..... • Review Frequency: ☐ Monthly ☒ Quarterly เป็นอย่างน้อย ☐ Annually ☒ Event-driven • Monitoring Focus: ☐ Bias/Fairness ☒ Model Drift/Performance Degradation ☐ Hallucination/Inaccurate Output ☐ Prompt Injection/Misuse ☐ Data Leakage/Sensitive Data Exposure	• AIMS Controls - 5. วงจรชีวิตของระบบ AI 7. การใช้งานระบบ AI 8. ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก • IMDA - เน้นการ monitor และ test ต่อเนื่องสำหรับ Agentic AI • IBM - เน้น action monitoring, accountability และ agent identity management

องค์ประกอบสำคัญของ AI Inventory (6 Elements)	ข้อมูลที่ต้องบันทึก (Record)	AIMS/Security Controls reference
	☐ Complaints/User Feedback ☐ Suitability/Appropriateness of Advice ☒ Market Abuse/Abnormal Trading Pattern ☐ Cybersecurity/Adversarial Events ☐ Third-party/Vendor Performance ☐ Availability/Service Disruption ☐ Regulatory/Compliance Breach ☒ Other: ___order anomalies, limit breach, signal decay, data latency, tool-call anomaly, unauthorized action attempt, override rate, approval latency, kill-switch activation, post-trade exceptions, vendor/API outage_____ • Incident/Kill Switch: ☒ Yes ต้องสามารถ suspend agent execution ทันที และ fallback ไป manual trading workflow ☐ No • Fallback Process manual trading/existing approved trading process/restricted operation mode..... • Third-party Provider (if any) ... market data provider, AI model provider, cloud/API provider, trading infrastructure provider (ถ้ามี).... • Tool-call/Action Logs [เพิ่มเติมสำหรับ Agentic AI] ต้องบันทึกทุก agent-initiated หรือ agent-recommended action • Tamper-evident action log [เพิ่มเติมสำหรับ Agentic AI] บันทึกทุก action ในรูปแบบ append-only ที่แก้ไขไม่ได้ เพื่อ reconstruct เหตุการณ์ได้เสมอ	

หมายเหตุ

ตารางนี้มีวัตถุประสงค์เพื่อสนับสนุนการพิจารณาและบันทึกเหตุผลเชิงกำกับดูแล สำหรับกรณีการใช้งาน AI ที่มีความเสี่ยงสูง และเป็น Agentic AI โดยได้กำหนดตัวอย่างและดำเนินการมาตรการควบคุมในทางปฏิบัติให้สอดคล้องกับกรอบระบบการจัดการ AI (AIMS) ที่กำหนดไว้ในเนื้อหาหลักและเครื่องมือที่นำเสนอ โดยจะสังเกตได้ว่า ในระดับความเสี่ยงที่ต่างกัน และคุณลักษณะหรือพฤติกรรมของ AI ในระดับที่ต่างกันทำให้รายละเอียดหรือความลึกของการจัดทำ AI Inventory นั้นต้องพิจารณาให้ครอบคลุมและสอดคล้องกับพฤติกรรมหรือการใช้งานตามวัตถุประสงค์ดังกล่าวด้วย ทั้งนี้ ตารางข้างต้นทำหน้าที่เป็นกลไกเชื่อมการประเมินความเสี่ยงกับการเลือกใช้มาตรการควบคุมอย่างได้สัดส่วน โดยไม่ซ้ำซ้อนกับ AIMS ซึ่งองค์กรสามารถประยุกต์ใช้ได้ตามความเหมาะสม สอดคล้องกับความเสี่ยงและบริบทการใช้งาน

3. ตัวอย่าง AIMS Controls (Full version)

วัตถุประสงค์ของส่วนนี้

ตาราง AIMS Controls ในเอกสาร พิจารณา controls ตาม Agentic AI-specific Controls สำหรับ Developer/Customizer และ User โดยครอบคลุม 8 ขั้นตอนหลักของ AIMS ได้แก่ นโยบาย โครงสร้างองค์กร ทรัพยากร การประเมินผลกระทบ วงจรชีวิต AI ข้อมูล การใช้งาน และความสัมพันธ์ภายนอก

ตารางที่ 15 ตัวอย่างการเลือกมาตรการควบคุม AIMS Controls ของ Automated Trading Support

หมวด AIMS Controls	วัตถุประสงค์การเลือกใช้มาตรการควบคุม Agentic AI for Automated Trading Support	การพิจารณาในส่วนที่เกี่ยวข้องกับ Agentic AI	Minimum Evidence
1. นโยบาย AI	กำหนดชัดว่า Agentic AI trading use case ต้องอยู่ภายใต้ approved strategy, risk appetite, prohibited conduct boundary และห้าม self-modify strategy โดยไม่มี approval	ระบุ permitted actions และ prohibited actions เช่น ห้ามซื้อขายหลักทรัพย์ที่ไม่ได้อนุมัติ, ห้ามฝ่าฝืน limits, ห้ามเปลี่ยน strategy เอง เป็นต้น	AI usage policy, trading AI policy, prohibited-use statement, approved strategy scope
2. โครงสร้างภายในองค์กร	กำหนด Business Owner, Risk Owner, Model/System Owner, Compliance Reviewer, IT/Cyber Owner, kill-switch authority และ escalation path	ต้องระบุว่าใครมีอำนาจ approve, override, suspend, rollback และ accept residual risk	RACI, approval matrix, escalation contact, committee minutes, residual risk acceptance memo
3. ทรัพยากรสำหรับระบบ AI	จำกัดสิทธิ์ของ agent ตามหลัก least privilege และแยก agent identity จาก human user หรือ service account ทั่วไป	[เพิ่มเติมสำหรับ Agentic AI] ใช้ dedicated agent identity, non-transferable credentials, tool-specific permissions และ permission boundary	Agent identity record, access list, permission matrix, API credential control, access review log
4. การประเมินผลกระทบของระบบ AI	ประเมิน market impact, investor harm, systemic risk, action-space, autonomy, reversibility และ residual risk	[เพิ่มเติมสำหรับ Agentic AI] ต้องประเมิน read/write/execute access, autonomous action, irreversible action และ blast radius ของ agent	Full assessment memo, 10 rubrics scoring, action-space assessment, autonomy assessment, risk acceptance record
5. วงจรชีวิตของระบบ AI	ควบคุมตั้งแต่ design, testing, validation, release, monitoring, rollback และ retirement	ทบทวนเมื่อเพิ่ม tool, connector, model version, data source, memory capability, autonomy level หรือ strategy scope	Model/system documentation, validation report, backtesting result, stress test result, release/rollback record, version log

หมวด AIMS Controls	วัตถุประสงค์การเลือกใช้มาตรการควบคุม Agentic AI for Automated Trading Support	การพิจารณาในส่วนที่เกี่ยวข้องกับ Agentic AI	Minimum Evidence
6. ข้อมูลสำหรับระบบ AI	ควบคุม data quality, source approval, latency, lineage, input integrity และ confidentiality ของกลยุทธ์/limits	ระวัง indirect prompt injection จากข้อมูลภายนอกหรือ market commentary ที่ agent ingest ผ่าน tools/API	Approved data source list, data quality report, latency monitoring, data lineage record, input validation record, data incident log
7. การใช้งานระบบ AI	ใช้ pre-trade controls, hard trading limits, approved instruments, real-time monitoring, human-on-the-loop, kill switch และ post-trade surveillance โดยจัดให้มี "จุดตัดสินใจก่อนส่งคำสั่ง (pre-execution checkpoint)" ที่ประเมินคำสั่งทุกคำสั่งก่อนไปถึงระบบซื้อขาย	[เพิ่มเติมสำหรับ Agentic AI] ใช้ deterministic tool-layer safeguards, checkpoint-based approval, tool-call logs, order-action logs, override-rate monitoring พร้อมกำหนดผลลัพธ์การตัดสินใจก่อนส่งคำสั่งอย่างชัดเจน (Runtime Disposition) และจัดเก็บบันทึกการตัดสินใจในรูปแบบที่แก้ไขไม่ได้	Trading limit configuration, pre-trade control evidence, surveillance dashboard, tool-call log, order-action log, override rate, approval latency, kill-switch log, pre-execution decision record (action + outcome + rule triggered + timestamp), append-only audit log configuration
8. ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก	สอบทาน market data providers, AI model providers, cloud/API providers, trading infrastructure providers, SLA, incident notice และ fallback/exit plan	ประเมิน dependency จาก provider, API change, outage, model update, data feed failure และ concentration risk	Vendor due diligence, SLA, incident notice procedure, provider change log, exit plan, fallback process

หมายเหตุ

มาตรการควบคุมขั้นต่ำในตัวอย่างนี้ถูกเลือกจากข้อมูลที่บันทึกใน AI Inventory และจากตาราง AIMS Controls ตามบทบาทขององค์กรในผู้กำกับดูแล Agentic AI ทั้งนี้ มาตรการดังกล่าวไม่ใช่รายการ controls ทั้งหมดของ AIMS แต่เป็นตัวอย่างการเลือก controls ในระดับที่ได้สัดส่วนกับ use case ความเสี่ยงสูง และยกตัวอย่างการพิจารณาเพิ่มเติมในส่วนที่เกี่ยวข้องกับ Agentic เพิ่มเติม อย่างไรก็ตาม หากข้อมูล ลักษณะการใช้งาน ระดับอิทธิพลของ AI หรือผลกระทบต่อผู้ลงทุนเปลี่ยนแปลง ผู้ประกอบธุรกิจสามารถทบทวนแนวทางการประเมินความเสี่ยงตามกระบวนการ 12 ขั้นตอน การจัดทำและปรับปรุง AI Inventory และเลือก controls ที่เข้มข้นขึ้นตามความเหมาะสม

4. ตัวอย่าง Trigger: เมื่อใด Trigger for Immediate Review/Suspension

ระบบควรถูกทบทวนหรือระงับการใช้งานทันทีเมื่อเกิดเหตุการณ์ต่อไปนี้

ตารางที่ 16 ตัวอย่าง trigger หรือปัจจัยที่ต้องระงับการใช้งานหรือทบทวนระบบ Agentic AI

Trigger	เหตุผลที่ต้องระงับการใช้งานหรือทบทวน
Agent พยายามส่งคำสั่งเกิน approved trading limits	อาจเกิด market impact หรือ breach risk appetite
Agent พยายาม trade unapproved instruments	ใช้งานเกิน mandate และ policy boundary
เกิด unexplained order spike หรือ abnormal order pattern	อาจเป็น model failure, cyber event หรือ market conduct risk
Market data feed failure หรือ severe latency	ระบบอาจตัดสินใจจากข้อมูลผิด/ล่าช้า
Model drift หรือ performance degradation เกิน threshold	controls เดิมอาจไม่เพียงพอ
Kill-switch ถูก activate	ต้องทบทวน root cause และ residual risk
Post-trade surveillance พบ pattern ผิดปกติ	อาจเกี่ยวข้องกับ market abuse/conduct concern
Override rate ต่ำหรือสูงผิดปกติ	อาจสะท้อน automation bias หรือ model instability
Unauthorized tool call หรือ API call	อาจสะท้อน permission breach หรือ adversarial event
เปลี่ยน model, strategy scope, data provider, tool connector หรือ autonomy level	risk profile เปลี่ยน ต้อง reassess
มี audit finding, incident หรือ regulatory concern	ต้อง review controls และ governance decision

หมายเหตุ

รายละเอียดเชิงเทคนิคของ controls เช่น agent identity management, runtime enforcement, deterministic tool-layer safeguards, tool-call logging, indirect prompt injection protection, continuous agent monitoring และ containment ควรอ้างอิงแนวทางเฉพาะด้าน Agentic AI เพิ่มเติมนอกเหนือจาก AIMS Controls เช่น IMDA Model AI Governance Framework for Agentic AI, IBM Agentic AI Governance Playbook/Agentic IAM, MAS' BuildFin.ai initiative Safeguards for Agentic Finance at Runtime และแนวปฏิบัติด้าน AI Security ของ สกมช. ตามความเหมาะสมของระบบและระดับความเสี่ยง

ภาคผนวก 8 ความเสี่ยงสำคัญจากการใช้งาน AI ในภาคตลาดทุน

ในช่วงหลายปีที่ผ่านมา ภาคตลาดทุนทั่วโลกต้องเผชิญความท้าทายจากการนำ AI มาใช้งานอย่างกว้างขวาง โดยเฉพาะเทคโนโลยี Generative AI และ Large Language Models (LLMs) ซึ่งแม้จะเพิ่มศักยภาพด้านประสิทธิภาพและนวัตกรรม แต่ขณะเดียวกันก็สร้างความเสี่ยงใหม่ทั้งในมิติการคุ้มครองผู้ลงทุน ความมั่นคงและโปร่งใสของตลาด และเสถียรภาพของระบบการเงิน หน่วยงานกำกับดูแลในหลายประเทศจึงเร่งพัฒนากฎกำกับดูแลที่เน้นการบริหารความเสี่ยงตามระดับความรุนแรงและผลกระทบ โดยภาพรวม คาดหวังให้องค์กรสามารถนำ AI มาใช้อย่างมีความรับผิดชอบตลอดวงจรชีวิตของระบบ (AI lifecycle) ตั้งแต่การกำหนดวัตถุประสงค์และบริบทการใช้งานอย่างชัดเจนว่า AI ถูกนำมาใช้เพื่อสนับสนุนการตัดสินใจหรือการให้บริการใด และมีขอบเขตที่ไม่ก่อให้เกิดอันตรายต่อผู้ลงทุนหรือบิดเบือนกลไกตลาด ตั้งแต่ในระยะออกแบบและพัฒนา พร้อมประเมินการพึ่งพาผู้ให้บริการเทคโนโลยีภายนอกอย่างเป็นระบบ ตลอดจนการนำไปใช้งานจริง ต้องมีการเปิดเผยและสื่อสารอย่างโปร่งใส กำหนด guardrails และ human oversight ไปตลอดจนถึงการติดตาม ทบทวนและปรับปรุงการควบคุมอย่างสม่ำเสมอ เป้าหมายสูงสุดคือ การใช้งาน AI เพื่อเพิ่มประสิทธิภาพและนวัตกรรม โดยไม่บั่นทอนการคุ้มครองผู้ลงทุน ความเป็นธรรมและความโปร่งใสของตลาด รวมถึงเสถียรภาพของระบบการเงินในภาพรวม

ตารางที่ 17 ประเภทความเสี่ยงสำคัญจากการใช้งาน AI ในภาคตลาดทุนและตัวอย่างมาตรการควบคุมความเสี่ยง

ประเภท	ความเสี่ยง	ผลกระทบ	ตัวอย่างแนวทาง/มาตรการควบคุมความเสี่ยง
RS-1:Firm-use-of-AI Risk leading to Market Abuse/Misconduct	การใช้ AI/GenAI ภายในองค์กร เช่น โมเดลเทรดอัตโนมัติ หรือระบบสื่อสารอัตโนมัติ อาจก่อให้เกิดรูปแบบคำสั่งซื้อขายที่คล้ายพฤติกรรมต้องห้าม (เช่น spoofing/layering) โดยไม่ตั้งใจ, การสื่อสารหรือการตลาดที่ทำให้ผู้ลงทุนเข้าใจผิด (AI-washing), การขาดการกำกับดูแลและทดสอบโมเดลที่เพียงพอ	ความเสี่ยงต่อการละเมิดกฎหมาย/ข้อบังคับด้าน Market Abuse	วัตถุประสงค์หลัก: “กำกับดูแลการพัฒนา-ใช้งาน AI/GenAI ภายในองค์กรอย่างถูกกฎ ระเบียบวิธีแบบจำลอง เพื่อป้องกันการกระทำ/การสื่อสารที่อาจก่อ market abuse หรือทำให้ผู้ลงทุนเข้าใจผิด เช่น AI-washing” 1. AI Governance & Model Risk Management - จัดตั้งกรอบการกำกับดูแล AI ครอบคลุมการออกแบบ ทดสอบ และอนุมัติก่อนใช้งาน - ใช้ stress testing และ behavioral constraints เพื่อป้องกัน pattern ที่คล้าย market abuse - กำหนด human-in-the-loop สำหรับการตัดสินใจสำคัญ 2. Surveillance & Monitoring - เฝ้าระวังธุรกรรมที่เกิดจากโมเดล AI และตรวจจับ pattern ผิดปกติ - ทำ post-trade analysis เพื่อตรวจสอบพฤติกรรมที่อาจเข้าข่าย market abuse - ติดตามรูปแบบหรือพฤติกรรมของตลาดและกำหนดกลไก anomaly detection alerts เพื่อระบุพฤติกรรมผิดปกติ เช่น คำสั่งซื้อขายที่มีรูปแบบผิดปกติ ความผันผวนที่ผิดปกติที่สัมพันธ์กับการใช้ AI หรือพฤติกรรมที่อาจกระทบต่อ market integrity 3. Communication & Marketing Controls - กำหนดนโยบายการขออนุมัติล่วงหน้า pre-clearance สำหรับข้อความ/สื่อที่สร้างโดย AI - จัดเก็บหลักฐานรองรับการกล่าวอ้างเกี่ยวกับ AI เช่น ผลการทดสอบ ประสิทธิภาพ ข้อจำกัดของระบบ ขอบเขตการใช้งาน

ประเภท	ความเสี่ยง	ผลกระทบ	ตัวอย่างแนวทาง/มาตรการควบคุมความเสี่ยง
			- กำหนดให้ข้อความที่กล่าวถึงความสามารถของ AI ต้องตรวจสอบได้ ไม่เกินจริง และสอดคล้องกับการใช้งานจริง - บันทึก version ของข้อความ สื่อโฆษณา หรือ communication ที่สร้างหรือช่วยสร้างโดย AI เพื่อใช้ตรวจสอบย้อนหลัง 4.Training & Awareness - อบรมทีมพัฒนา เทสต์ และการตลาดเกี่ยวกับความเสี่ยงด้านกฎระเบียบและจริยธรรม - สื่อสารแนวทางปฏิบัติที่ถูกต้องและบทลงโทษหากละเมิด
RS-2: Model and Data Risk (Bias/Drift/ Hallucination/ Prompt Injection)	การตัดสินใจที่อาศัยโมเดลอาจผิดพลาด/ไม่เป็นธรรม เนื่องจากข้อจำกัดของโมเดลและข้อมูลที่ใช้: เช่น ข้อมูลไม่ครบถ้วน ไม่เป็นปัจจุบัน ไม่เหมาะสมกับวัตถุประสงค์ แหล่งที่มาไม่ชัดเจน สิทธิการใช้ ไม่เพียงพอ หรือคุณภาพข้อมูล เปลี่ยนแปลงไปตามสภาพตลาด Bias: อคติจากข้อมูลที่น่ามาฝึกโมเดล, Drift: โมเดลเสื่อมคุณภาพเมื่อสภาพตลาดเปลี่ยน, Hallucination: LLM สร้างข้อมูลเท็จ, Prompt Injection: การโจมตีด้วยคำสั่งแฝงหรือ chain-of-tools	คำแนะนำ สัญญาณการ ลงทุนผิดพลาด ความเสียหายต่อผู้ลงทุน และชื่อเสียงองค์กร ความเสี่ยงด้านกฎหมาย และการกำกับดูแล	วัตถุประสงค์หลัก: ลดความผิดพลาดของโมเดล, ป้องกันการโจมตี, ทำให้ระบบโปร่งใส ตรวจสอบได้ และมีจุดควบคุมโดยมนุษย์ 1. Data Governance for AI - ตรวจสอบแหล่งที่มา คุณภาพ ความครบถ้วน ความเป็นปัจจุบัน ระดับชั้น สิทธิการใช้ ข้อจำกัดด้านกฎหมาย/สัญญา และ ความเหมาะสมของข้อมูลต่อวัตถุประสงค์การใช้ AI รวมถึงเชื่อมโยงกับ Data Catalog, Metadata, ROPA, Data Flow หรือ Data Owner ตามความเหมาะสม 2. Model Governance & Inventory - จัดทำ Model Inventory ครอบคลุมทุกโมเดล พร้อม Criticality Tiering เพื่อให้รู้ว่ามีโมเดลใดบ้าง ความเสี่ยงระดับไหน และใครเป็นเจ้าของ 3. Lifecycle Control (NIST AI RMF) - ใช้กรอบ MAP-MEASURE-MANAGE-GOVERN ครอบคลุม ตั้งแต่การออกแบบจนถึงการเลิกใช้งาน ให้มั่นใจว่ามีการประเมิน ความเสี่ยงและควบคุมตลอดวงจรชีวิต 4. Validation & Testing - จัดทำ Pre-Deployment Validation และ Post-Deployment Monitoring - Adversarial Testing/Red-Teaming เพื่อตรวจ prompt injection, jailbreak - Bias & Fairness Testing ตรวจสอบอคติในผลลัพธ์ - Robustness & Stress Testing ทดสอบความทนทานต่อข้อมูล เพื่อลดโอกาสเกิดข้อผิดพลาดก่อนใช้งานจริง และเฝ้าระวัง หลังใช้งาน - กำหนดตัวชี้วัดประสิทธิภาพของโมเดลเพิ่มเติม เช่น prediction accuracy, error rates, model drift and degradation over time และกำหนด risk thresholds หรือเกณฑ์ยกระดับเมื่อพบผลการติดตามที่เบี่ยงเบนจากระดับที่ยอมรับได้ - จัดให้มีนโยบายและกระบวนการทดสอบหรือทวนสอบโมเดล AI สำหรับ use case ที่มีความสำคัญ ทั้งก่อนนำไปใช้งานจริงและ ระหว่างการใช้งาน พร้อมจัดเก็บรายงานผลการทดสอบ ผลการ ทวนสอบ และข้อสรุปการอนุมัติใช้งาน เพื่อให้สามารถตรวจสอบย้อนหลังได้ 5. Guardrails & Safe Interaction - ตั้ง Policy & Guardrails สำหรับ input/output

ประเภท	ความเสี่ยง	ผลกระทบ	ตัวอย่างแนวทาง/มาตรการควบคุมความเสี่ยง
			- ใช้ Retrieval-Augmented Generation (RAG) เพื่อลด hallucination - ติดตั้ง Content Filters และ Prompt Management เพื่อป้องกันการโจมตีและลดความเสี่ยงจากข้อมูลเท็จ 6. Human-in-the-Loop (HITL) & Kill-Switch - กำหนดจุด HITL ในการตัดสินใจสำคัญ เช่น ก่อนส่งคำแนะนำลูกค้า หรือก่อนส่งคำสั่งซื้อขาย - จัดทำ Kill-Switch สำหรับหยุดระบบเมื่อพบความผิดปกติ เพื่อให้มนุษย์สามารถควบคุมและยับยั้งความเสี่ยงได้ทันที
RS-3: Third-Party/Concentration Risk (Foundation Models & Vendors)	การพึ่งพาผู้ให้บริการโมเดลหรือโครงสร้างพื้นฐาน AI ภายนอกมากเกินไป อาจทำให้เกิด single point of failure, ปัญหาการปฏิบัติตามกฎหมาย (compliance/sovereignty) หรือการละเมิดสัญญาใช้ข้อมูล	บริการหยุดชะงักเมื่อผู้ให้บริการล้ม ความเสี่ยงด้านกฎหมายและชื่อเสียงจากการละเมิดสัญญาหรือข้อกำหนดข้อมูล ความเสี่ยงเชิงระบบหากผู้ให้บริการรายใหญ่มีปัญหา	วัตถุประสงค์หลัก: ลดความเสี่ยงจากการพึ่งพาผู้ให้บริการรายเดียว และสร้างความยืดหยุ่นในการดำเนินงาน 1. Third-Party Due Diligence - ตรวจสอบผู้ให้บริการ AI ครอบคลุม data lineage, security controls, IP rights, และ evaluation reports เพื่อมั่นใจว่าผู้ให้บริการมีมาตรฐานความปลอดภัยและปฏิบัติตามกฎหมาย 2. Contractual Safeguards - กำหนด SLA/VOLA เฉพาะ AI, เงื่อนไขการแจ้งเหตุ, สิทธิการตรวจสอบ และข้อกำหนดด้าน data sovereignty เพื่อลดความเสี่ยงจากการละเมิดสัญญาและเพิ่มความโปร่งใส 3. Dependencies, Concentrations & Oversight - จัดทำทะเบียนหรือรายการผู้ให้บริการภายนอกที่เกี่ยวข้องกับระบบ AI โดยระบุว่า use case ใดพึ่งพาผู้ให้บริการรายใด เช่น ผู้ให้บริการโมเดล AI ระบบคลาวด์ API ข้อมูล โครงสร้างพื้นฐาน รวมถึงประเมินว่าการพึ่งพาดังกล่าวมีนัยสำคัญต่อองค์กรเพียงใด - กำหนดกระบวนการติดตามและกำกับดูแลผู้ให้บริการ AI ที่มีนัยสำคัญอย่างต่อเนื่อง เช่น ตัวชี้วัดการให้บริการ เงื่อนไขแจ้งเหตุ กระบวนการยกระดับปัญหา เงื่อนไขการออกจากสัญญา หรือแผนรองรับกรณีต้องเปลี่ยนผู้ให้บริการ เพื่อให้การดำเนินงานยังคงต่อเนื่องและควบคุมความเสี่ยงได้ 4. Exit & Portability Plan - จัดทำแผน Exit Strategy และ Portability สำหรับย้ายโมเดล/ข้อมูลไปยังผู้ให้บริการรายอื่น เพื่อป้องกัน vendor lock-in และลด downtime เมื่อเกิดเหตุฉุกเฉิน 5. Cross-Model/Multi-Vendor Redundancy - ใช้ multi-model readiness หรือ multi-cloud strategy สำหรับ use case สำคัญ เพื่อลด single point of failure และเพิ่มความต่อเนื่องทางธุรกิจ
RS-4: Conduct/Client Best Interest & Record-Keeping	องค์กรขาดการคำนึงถึงหลัก best interest ของลูกค้า และการสื่อสารไม่ชัดเจน โปร่งใส อธิบายไม่ได้ และไม่มีการเก็บหลักฐานการตัดสินใจเพื่อตรวจสอบย้อนหลัง อันเนื่อง	ความเสียหายต่อผู้ลงทุน จากคำแนะนำที่ไม่เหมาะสม ความเสี่ยงด้านกฎหมายและค่าปรับจากการละเมิดกฎ	วัตถุประสงค์หลัก: ให้การใช้ AI ในการแนะนำลูกค้าเป็นไปตามกฎหมาย ตรวจสอบได้ และป้องกันการใช้งานที่ไม่ได้รับอนุญาต 1. Governance & Policy Alignment - จัดทำ AI Governance Framework สำหรับงานคำแนะนำและจัดพอร์ต โดยอิงแนวทางสากล เพื่อให้มั่นใจว่าการใช้ AI ไม่ละเมิดข้อกำหนดด้าน conduct และความเหมาะสม

ประเภท	ความเสี่ยง	ผลกระทบ	ตัวอย่างแนวทาง/มาตรการควบคุมความเสี่ยง
	มาจากการใช้ AI ช่วยให้คำแนะนำหรือจัดพอร์ตลงทุน	สูญเสียความเชื่อมั่นและชื่อเสียงองค์กร	2. Standardized Disclosure - กำหนด รูปแบบการเปิดเผย (Disclosure) ให้ลูกค้าทราบว่าใช้ AI ในกระบวนการ พร้อมข้อจำกัดและวิธีติดต่อมนุษย์ เพื่อสร้างความโปร่งใสและลดความเข้าใจผิด 3. Audit Trail & Record-Keeping - บันทึก Prompt, Context, Output และเหตุผลประกอบการตัดสินใจในระบบที่ตรวจสอบได้ เพื่อพิสูจน์ความเหมาะสมและรองรับการตรวจสอบจากหน่วยงานกำกับ - กรณี use case ที่มีผลต่อผู้ลงทุนหรือตลาด ควรจัดเก็บเอกสารและบันทึกที่ทำให้ตรวจสอบย้อนหลังได้ว่า AI มีบทบาทอย่างไรต่อการตัดสินใจหรือการดำเนินการ เช่น ผลลัพธ์ที่ AI สร้างขึ้น การดำเนินการสุดท้ายที่เกิดขึ้น การทบทวนหรือแทรกแซงโดยมนุษย์ และหลักฐานที่เชื่อมโยงระหว่างผลลัพธ์ของ AI กับ การตัดสินใจที่ส่งผลกระทบต่อลูกค้าหรือตลาด 4. Human-in-the-Loop (HITL) - กำหนดจุด HITL ให้มนุษย์ตรวจสอบและอนุมัติคำแนะนำก่อนส่งถึงลูกค้า เพื่อลด automation bias และป้องกันความผิดพลาดที่กระทบลูกค้า 5. Shadow AI Prevention - ออกนโยบายห้ามใช้เครื่องมือ AI ที่ไม่ได้รับอนุญาต และติดตั้งระบบตรวจจับ เพื่อป้องกันการใช้ AI นอกกรอบควบคุมที่อาจละเมิดกฎหมายหรือทำให้ข้อมูลรั่วไหล
RS-5: Data/Privacy/IP & Cross-Border	ข้อมูลลูกค้าหรือข้อมูลภายในรั่วไหลจากการใช้งาน AI, การใช้ข้อมูลฝึกหรือป้อนให้ LLM โดยไม่มีสิทธิหรือฐานกฎหมายที่ถูกต้อง, ภาระตามกฎหมายเมื่อมีการโอนข้อมูลข้ามประเทศ, ความเสี่ยงด้านลิขสิทธิ์และทรัพย์สินทางปัญญาของเนื้อหาที่ AI สร้าง	การละเมิดกฎหมายคุ้มครองข้อมูล (PDPA, GDPR) ค่าปรับและความเสียหายทางชื่อเสียง ข้อพิพาทด้านลิขสิทธิ์ และความรับผิดทางกฎหมาย	วัตถุประสงค์หลัก: ป้องกันการรั่วไหลของข้อมูล คุ้มครองสิทธิส่วนบุคคล และลดความเสี่ยงด้านกฎหมาย/ลิขสิทธิ์ 1. Data Minimization & Masking - จำกัดการใช้ข้อมูลส่วนบุคคลให้น้อยที่สุด และทำการ masking ข้อมูลสำคัญ เพื่อลดโอกาสรั่วไหลและป้องกันการใช้อ้างอิงเกินความจำเป็น 2. Prompt Hygiene Policy - ออกนโยบาย “ห้ามใส่ข้อมูลส่วนบุคคลหรือข้อมูลลับใน prompt” และตรวจสอบการปฏิบัติ เพื่อป้องกันการรั่วไหลผ่านการโต้ตอบกับ AI 3. Content Provenance & Watermarking - ใช้เทคนิค provenance และ watermarking สำหรับเนื้อหาที่ AI สร้าง เพื่อยืนยันแหล่งที่มาและลดความเสี่ยงด้าน IP 4. DPIA/TRA สำหรับ Cross-Border Transfers - ทำ Data Protection Impact Assessment (DPIA) และ Transfer Risk Assessment (TRA) สำหรับการโอนข้อมูลข้ามประเทศ เพื่อให้มั่นใจว่าการโอนข้อมูลสอดคล้องกับกฎหมายและมาตรฐานสากล 5. Contractual IP & Indemnity Clauses - เพิ่มเงื่อนไขในสัญญากับผู้ให้บริการ AI เกี่ยวกับ สิทธิในข้อมูล/โมเดล และ การชดเชยความเสียหาย (indemnity) เพื่อลดความเสี่ยงข้อพิพาทด้านลิขสิทธิ์และความรับผิด

ประเภท	ความเสี่ยง	ผลกระทบ	ตัวอย่างแนวทาง/มาตรการควบคุมความเสี่ยง
			6. Fairness & Transparency Assessment - ใช้หลักการ FEAT หรือเทียบเท่าเพื่อตรวจสอบความเท่าเทียม ความโปร่งใสของการใช้ข้อมูล 7. Data Governance, Quality & Traceability - จัดให้มีแนวทางกำกับดูแลข้อมูลสำหรับระบบ AI ที่มีนัยสำคัญ โดยข้อมูลที่ใช้ควรมีความครบถ้วน ถูกต้อง ทันเวลา สอดคล้อง ตรงตามวัตถุประสงค์ - จัดเก็บหลักฐานการประเมินคุณภาพข้อมูล การตรวจสอบ ลักษณะของข้อมูล การแก้ไขข้อบกพร่อง และการตรวจสอบ แหล่งที่มาหรือเส้นทางการใช้ข้อมูล รวมถึงการประเมินข้อมูลจาก แหล่งภายนอกหรือผู้ให้บริการภายนอกตามความเหมาะสม
RS-6: Operational Resilience & Cyber	ระบบ AI หยุดทำงานเพราะผู้ให้บริการโมเดลหรือคลาวด์ล่ม, การโจมตีแบบ data/model poisoning ทำให้โมเดลให้ผลลัพธ์ผิด, ความเสี่ยงจาก supply-chain compromise (ช่องโหว่ใน library หรือ API), ความผิดพลาดของ agentic AI ที่ทำงานอัตโนมัติโดยไม่มีการควบคุม	การหยุดชะงักของธุรกิจและบริการลูกค้า ความเสียหายต่อความปลอดภัยของข้อมูลและชื่อเสียงองค์กร ความเสี่ยงเชิงระบบ หากเกิดพร้อมกันหลายจุด	วัตถุประสงค์หลัก: สร้างความทนทานของระบบ AI ให้พร้อมรับมือ เหตุขัดข้องและการโจมตี และป้องกันความเสียหายต่อธุรกิจ 1. AI-Specific Threat Modeling & Red-Teaming - ทำ Threat Modeling สำหรับ AI และทดสอบด้วย Red-Teaming เพื่อตรวจหาช่องโหว่ เช่น prompt injection, model poisoning เพื่อระบุและปิดช่องโหว่ก่อนถูกโจมตีจริง 2. SBOM/MBOM (Software & Model Bill of Materials) - จัดทำ SBOM/MBOM สำหรับ library, dependency และ โมเดลที่ใช้งาน เพื่อตรวจสอบความปลอดภัยของซัพพลายเชนและรองรับการแจ้งเตือนช่องโหว่ 3. Chaos Testing & Failover Simulation - ทำ Chaos Testing และจำลองเหตุการณ์ failover ระหว่าง โมเดลหรือภูมิภาค เพื่อทดสอบความพร้อมของระบบเมื่อเกิดเหตุ 4. DR/BCP Playbook for AI - จัดทำ Disaster Recovery (DR) และ Business Continuity Plan (BCP) ครอบคลุมหลายกรณี เพื่อลด downtime และรักษา ความต่อเนื่องของบริการ 5. Incident Logs, Monitoring Records & Reportable Incidents - จัดให้มีบันทึกและติดตามสำหรับเหตุการณ์ที่เกี่ยวข้องกับระบบ AI เช่น ระบบทำงานผิดพลาด การให้ผลลัพธ์ไม่ถูกต้อง เหตุด้าน ความมั่นคงปลอดภัยไซเบอร์ ข้อมูลรั่วไหล หรือเหตุการณ์ที่อาจ กระทบต่อผู้ลงทุน บริการสำคัญ หรือความเป็นธรรมของตลาด - กำหนดเกณฑ์ว่าเหตุการณ์ใดเกี่ยวกับ AI ที่ต้องรายงานหรือ ยกระดับต่อผู้บริหาร หรือพิจารณาแจ้งหน่วยงานกำกับดูแล เช่น เหตุการณ์ที่มีนัยสำคัญต่อผู้ลงทุน บริการสำคัญ ความมั่นคง ปลอดภัยไซเบอร์ ข้อมูลรั่วไหล ความเอนเอียง (Bias) รุนแรง หรือ ผลลัพธ์ผิดพลาดที่อาจสร้างความเสียหาย

ภาคผนวก 9 คำถามที่พบบ่อย (FAQs)

ส่วนนี้จัดทำขึ้นเพื่อช่วยให้องค์กร ผู้ประเมิน และผู้บริหาร สามารถทำความเข้าใจหลักการและการนำคู่มือไปประยุกต์ใช้ได้อย่างสอดคล้อง คู่มือฉบับนี้มีลักษณะเป็น guideline เชิงหลักการ สามารถปรับปรุงเพิ่มเติมโดยเน้นการใช้ดุลยพินิจตามบริบท ความเสี่ยง และระดับความพร้อมขององค์กร ทั้งนี้ องค์กรยังคงสามารถศึกษาและประยุกต์ใช้หลักการหรือแนวทางเชิงเทคนิคจากมาตรฐานสากลที่เกี่ยวข้องได้ตามความเหมาะสม

หมวด A: ความยืดหยุ่น และหลัก Proportionality

Q1: หากองค์กรมีงบประมาณหรือทรัพยากรจำกัด จำเป็นต้องปฏิบัติตามคู่มือนี้ครบทุกข้อหรือไม่?

A: องค์กรสามารถพิจารณาเลือกนำคู่มือฉบับนี้ไปใช้ตามหลัก Proportionality โดยให้ความสำคัญกับระบบ AI ที่มีความเสี่ยงหรือผลกระทบสูงก่อน เช่น ระบบที่ส่งผลกระทบต่อผู้ลงทุนหรือมีผลต่อการดำเนินงานสำคัญ ขณะที่ AI ที่มีความเสี่ยงต่ำสามารถใช้มาตรการควบคุมขั้นพื้นฐานได้ เพื่อรักษาสมดุลระหว่างการกำกับดูแลและความคล่องตัวในการดำเนินงาน

Principle reference: (Risk-based & proportional governance; OECD AI Principles; NIST AI RMF)

Q2: มาตรการควบคุมขั้นต่ำ (Minimum Governance Baseline) สามารถปรับลดหรือปรับเพิ่มได้หรือไม่?

A: ได้ มาตรการควบคุมที่เสนอในคู่มือเป็นเพียงเกณฑ์อ้างอิงขั้นต่ำ องค์กรสามารถปรับระดับความเข้มได้ตามความพร้อมของบุคลากร ความเข้มแข็งของระบบบริหารความเสี่ยงเดิม (ERM/IT Governance) และลักษณะการใช้งาน AI ทั้งนี้ การปรับมาตรการควรมีเหตุผลและสามารถอธิบายได้

Principle reference: (Adaptive governance; ISO/IEC 42001; NIST AI RMF – Govern Function)

หมวด B: ขอบเขตของความเสี่ยงและการประเมินตาม PIA

Q3: หาก AI ที่ใช้งานไม่มีข้อมูลส่วนบุคคลเลย ยังจำเป็นต้องทำกระบวนการประเมิน (12 ขั้นตอน) หรือไม่?

A: ยังจำเป็น เนื่องจากความเสี่ยงของ AI ไม่ได้จำกัดอยู่ที่ด้านข้อมูลส่วนบุคคลเท่านั้น แต่ยังคงครอบคลุมความเสี่ยงด้านความมั่นคงทางเทคนิค ความต่อเนื่องในการดำเนินงาน และความยืดหยุ่นของระบบ ทั้งนี้ ระดับความละเอียดของการประเมินสามารถปรับให้เหมาะสมกับความเสี่ยงที่แท้จริงได้

Principle reference: (Broad AI risk conception beyond privacy; ISO/IEC 23894; BIS operational resilience thinking)

Q4: การดำเนินการตาม 12 ขั้นตอนต้องใช้เวลานานเพียงใดต่อหนึ่งโครงการ?

A: ระยะเวลาขึ้นอยู่กับระดับความเสี่ยงและผลกระทบของ AI ระบบที่มีความเสี่ยงต่ำอาจใช้การทบทวนในระดับสั้น ขณะที่ระบบที่มีผลกระทบต่อผู้ลงทุนหรือเสถียรภาพของตลาดทุนอาจต้องใช้การวิเคราะห์เชิงลึกและขั้นตอนการอนุมัติที่รอบคอบมากขึ้น

Principle reference: (Risk-tailored assessment depth; proportionality principle in global AI governance frameworks)

Q5: หากผลการทดสอบในขั้น Sandbox ไม่เป็นไปตามที่คาด จำเป็นต้องเริ่มกระบวนการใหม่ทั้งหมดหรือไม่?

A: ไม่จำเป็น กระบวนการประเมิน AI ถูกออกแบบให้เป็นวงจรวนรอบ (iterative lifecycle) องค์กรสามารถย้อนกลับไปปรับปรุงเฉพาะขั้นตอนที่เกี่ยวข้อง เช่น ข้อมูล โมเดล หรือสมมติฐานการใช้งาน โดยไม่ต้องเริ่มใหม่ทั้งหมด

Principle reference: (Iterative lifecycle management; NIST AI RMF; SDLC-aligned AI governance)

Q6: องค์กรควรเริ่มจาก use case ไດก่อนเมื่อต้องการนำรอบการกำกับดูแล AI ไปใช้?

A: หากองค์กรมีการเริ่มใช้งาน AI ไปบ้างแล้ว และต้องการนำรอบการกำกับดูแลมาประยุกต์ใช้ ควรเริ่มจาก Think เพื่อระบุ use case วัตถุประสงค์ บทบาทของ AI ต่อการตัดสินใจ และผลกระทบที่อาจเกิดขึ้นต่อผู้ลงทุน ลูกค้า ตลาด ข้อมูลสำคัญ หรือบริการหลัก จากนั้นเข้าสู่ Record โดยบันทึก use case ลงใน AI Inventory และประเมินระดับความเสี่ยงเบื้องต้น ก่อนเข้าสู่ Act โดยใช้ผลการประเมินและ Risk Rubrics เพื่อจัดลำดับความสำคัญและเลือกมาตรการควบคุมที่เหมาะสม โดยควรเริ่มจาก use case ที่มีผลกระทบสูงหรือมีความเสี่ยงสูงก่อน

Principle reference: Think > Record > Act; AI Inventory; Risk Rubrics; IOSCO Supervisory Toolkit; NIST AI RMF

หมวด C: Human Oversight และความรับผิดชอบ

Q7: ขั้นตอน Human-in-the-loop จำเป็นต้องมีการตรวจสอบโดยมนุษย์ทุกกรณีหรือไม่?

A: ไม่เสมอไป ระดับและรูปแบบของ human oversight ควรพิจารณาจากระดับความเสี่ยงสุทธิและอิทธิพลของ AI ต่อการตัดสินใจ หากผลกระทบต่ำ อาจใช้การตรวจสอบแบบสุ่มหรือย้อนหลัง อย่างไรก็ตาม หากกระบวนการ AI มีความซับซ้อนและมีความเสี่ยงที่สูงขึ้น การมีมนุษย์เข้ามา oversight แค่นี้เพียงบางส่วน อาจไม่เพียงพอ ดังนั้น ระดับของ human oversight จึงขึ้นอยู่กับระดับความเสี่ยง ผลกระทบ และขนาดของโครงการ

Principle reference: (Meaningful human oversight; EU AI Act; MAS FEAT Principles)

Q8: ใครคือผู้รับผิดชอบหลักในแต่ละช่วงของกระบวนการประเมิน AI?

A: โดยทั่วไป Business Owner เป็นผู้รับผิดชอบกำหนดวัตถุประสงค์และบริบท ฝ่ายเทคนิคหรือข้อมูลรับผิดชอบด้านการออกแบบและข้อมูล ขณะที่ฝ่ายกำกับดูแลและคณะกรรมการที่เกี่ยวข้องมีบทบาทในการพิจารณาความเสี่ยงและให้ความเห็น ทั้งนี้ การแบ่งบทบาทควรสอดคล้องกับโครงสร้างองค์กรและหลัก accountability

Principle reference: (Clear accountability & role clarity; ISO/IEC 42001; IOSCO sound governance principles)

Q9: Use case ของ Agentic AI ควรพิจารณาเพิ่มเติมจาก GenAI ทั่วไปในประเด็นใด?

A: Agentic AI ควรถูกพิจารณาเพิ่มเติมจาก GenAI ทั่วไป เนื่องจากอาจมีความสามารถในการวางแผน เรียกใช้เครื่องมือ เชื่อมต่อระบบภายนอก หรือเสนอ/ดำเนินการแทนมนุษย์ภายใต้ขอบเขตที่กำหนด องค์กรจึงควรพิจารณาเรื่องขอบเขตการกระทำ ความเป็นอิสระในการตัดสินใจ สิทธิเข้าถึงข้อมูลและเครื่องมือ ความสามารถในการย้อนกลับผลการกระทำ การบันทึก tool-use/action logs การเฝ้าติดตามอย่างใกล้ชิด และมาตรการ kill switch หรือการควบคุมอื่นที่เทียบเท่า โดยเฉพาะ use case ที่อาจกระทบผู้ลงทุน ตลาด หรือบริการสำคัญ

Principle reference: (Human Oversight; Agentic AI; AI Lifecycle; IOSCO Supervisory Toolkit)

หมวด D: Third party, Security และการติดตามหลังใช้งาน

Q10: ระบบ AI ที่เป็นซอฟต์แวร์สำเร็จรูป (Commercial Off-the-Shelf: COTS/SaaS) จำเป็นต้องทำกระบวนการประเมินนี้หรือไม่?

A: จำเป็น เนื่องจากองค์กรยังคงมีความรับผิดชอบต่อผลกระทบจากการใช้งาน AI โดยควรเน้นการประเมินความเสี่ยงจากผู้ให้บริการภายนอก การตั้งค่าความปลอดภัย และการเชื่อมต่อกับระบบภายใน

Principle reference: (Third-party risk accountability; ISO/IEC 42001; TPRM best practices in financial sector)

Q11: กระบวนการนี้เชื่อมโยงกับมาตรฐานด้านความมั่นคงปลอดภัยไซเบอร์อย่างไร?

A: กระบวนการกำกับดูแล AI สามารถเชื่อมกับกรอบความมั่นคงปลอดภัยไซเบอร์เดิมขององค์กร เช่น ISO/IEC 27001 ได้ โดยใช้ controls เดิมเป็นฐาน เช่น การควบคุมสิทธิ์การเข้าถึง การจัดเก็บ log การบริหารเหตุการณ์ผิดปกติ และการบริหารผู้ให้บริการภายนอก แล้วขยายให้ครอบคลุมความเสี่ยงเฉพาะของ AI เช่น prompt/output log, model version, model drift, hallucination, tool-call log และ AI incident scenario ทั้งนี้ ISO/IEC 27001 ช่วยวางฐานด้านความมั่นคงปลอดภัยของข้อมูลและระบบ ขณะที่ AI governance ต้องเพิ่มการติดตามพฤติกรรม ผลลัพธ์ และผลกระทบของ AI ตลอดการใช้งาน เช่น เดิมองค์กรเก็บ log ว่าใครเข้าใช้ระบบเมื่อใด เมื่อใช้ GenAI ควรเพิ่ม log ของ prompt/output หรือหลักฐานว่าผลลัพธ์ของ AI ถูกนำไปใช้ในกระบวนการใด หรือเดิม incident response ครอบคลุม malware หรือระบบล่ม แต่หากใช้ AI ควรเพิ่ม scenario เช่น AI ให้คำตอบผิด, model drift, hallucination หรือ AI agent เรียกใช้ tool ผิดขอบเขต และสร้างความเสียหายให้แก่องค์กร และ/หรือผู้ให้บริการ

Principle reference: (Integrated risk management; ISO/IEC 27001; cyber-AI risk convergence)

Q12: หากองค์กรใช้โมเดลหรือบริการจากบุคคลที่ 3 เป็นหลัก ยังต้องรับผิดชอบต่อผลลัพธ์ของ AI หรือไม่?

A: องค์กรยังคงต้องรับผิดชอบต่อการนำระบบ AI ไปใช้งาน การกำหนดขอบเขตการใช้ การควบคุมความเสี่ยง การเปิดเผยข้อมูลที่เกี่ยวข้อง และผลกระทบที่อาจเกิดขึ้นต่อผู้ให้บริการ ลูกค้า ผู้ลงทุน หรือการดำเนินงานขององค์กร แม้ระบบ โมเดล หรือบริการ AI จะพัฒนาโดยบุคคลที่ 3 ก็ตาม ดังนั้น ควรมีการประเมินและกำกับดูแลผู้ให้บริการตามระดับความเสี่ยง เช่น due diligence สัญญาที่เหมาะสม การแจ้งเหตุผิดปกติ การติดตามผล และแผนรองรับกรณีผู้ให้บริการมีปัญหา

Principle reference: (AIMS 2.4.11; Third-party & Concentration Risk; IOSCO Supervisory Toolkit)

Q13: เมื่อใดที่องค์กรควรทบทวน use case หรือ model ใหม่ แม้จะยังไม่เกิดเหตุการณ์ผิดปกติ?

A: องค์กรควรทบทวน use case หรือ model ใหม่เมื่อมีการเปลี่ยนแปลงที่อาจกระทบต่อวัตถุประสงค์ ระดับความเสี่ยง ประสิทธิภาพ หรือความเหมาะสมของมาตรการควบคุมเดิม เช่น การเปลี่ยนโมเดล แหล่งข้อมูล ผู้ให้บริการ การเชื่อมต่อระบบ ลักษณะบริการ กลุ่มลูกค้า บริบทตลาด หรือข้อกำหนดด้านกฎหมายและการกำกับดูแล นอกจากนี้ ควรทบทวนเมื่อพบสัญญาณว่า performance ลดลง มี drift มี near miss มี incident หรือมีการ override โดยมนุษย์บ่อยผิดปกติ แม้ยังไม่เกิดเหตุการณ์เสียหายจริง

Principle reference: (AI Lifecycle; AIMS 2.4.9-2.4.10; Continuous Review; IOSCO Supervisory Toolkit)

Q14: หาก AI เกิด Model Drift หรือประสิทธิภาพลดลงหลังนำไปใช้งาน ควรดำเนินการอย่างไร?

A: ควรมีกกลไกติดตามผลอย่างต่อเนื่อง (continuous monitoring) เพื่อให้ตรวจพบว่าโมเดลเริ่มให้ผลลัพธ์ผิดไปจากที่คาดไว้หรือไม่ เมื่อพบ drift หรือประสิทธิภาพลดลง องค์กรควรวิเคราะห์สาเหตุ ประเมินผลกระทบ และตัดสินใจว่าจะปรับปรุง จำกัดขอบเขต หยุดใช้ชั่วคราว หรือประเมิน use case ใหม่ ทั้งนี้ ควรอัปเดต AI Inventory, risk tier และ controls ให้สะท้อนความเสี่ยงที่เปลี่ยนไป

Principle reference: (Continuous monitoring & resilience; NIST AI RMF - Manage & Measure Functions)

Q15: ISO/IEC 27001 และ ISO/IEC 42001 แตกต่างกันอย่างไรในบริบทของการกำกับดูแล AI?

A: ISO/IEC 27001 มุ่งเน้นการบริหารความมั่นคงปลอดภัยของข้อมูลและระบบสารสนเทศ เช่น ความลับ ความถูกต้องครบถ้วน และความพร้อมใช้งาน ส่วน ISO/IEC 42001 มุ่งเน้นการบริหารจัดการระบบ AI โดยรวม หรือ AIMS ซึ่งครอบคลุมเรื่องนโยบาย AI บทบาทและความรับผิดชอบ การประเมินผลกระทบ วงจรชีวิตของระบบ AI ข้อมูลสำหรับ AI การใช้งานระบบ AI และความสัมพันธ์กับบุคคลภายนอก ดังนั้น ISO/IEC 42001 จึงขยายจากการทำให้ “ระบบปลอดภัย” ไปสู่การทำให้ “AI ถูกใช้และ

ถูกกำกับดูแลอย่างรับผิดชอบ” เช่น ในส่วนของ ISO/IEC 27001 ดู cyber incident แต่ ISO/IEC 42001 อาจต้องพิจารณา AI incident เช่น hallucination, model drift, unfair output หรือ AI agent ทำ action เกินขอบเขต

Principle reference: (ISO/IEC 42001 vs ISO/IEC 27001 distinction; AI lifecycle governance; accountability, transparency and human oversight)

Q16: หากมี ISO/IEC 27001 อยู่แล้ว ควรต่อยอด controls อย่างไรให้สอดคล้องกับ ISO/IEC 42001: AIMS?

A: องค์กรไม่จำเป็นต้องสร้างระบบควบคุมใหม่ทั้งหมด แต่ควรใช้แนวทาง Reuse > Extend > Align กล่าวคือ Reuse controls เดิมจาก ISO/IEC 27001 เช่น access control, logging, incident response, change management และ supplier management; Extend controls เหล่านั้นให้ครอบคลุมความเสี่ยงเฉพาะของ AI เช่น AI Inventory, impact assessment, model monitoring, data suitability, explainability และ human oversight; และ Align controls ทั้งหมดกับหมวด AIMS ตามระดับความเสี่ยงของแต่ละ use case ทั้งนี้ ISO/IEC 42001 ใช้แนวทาง risk-based control selection และสามารถเลือก controls ตามความเหมาะสมของบริบทและความเสี่ยง ไม่จำเป็นต้องใช้ทุก controls เหมือนกันทุกกรณี

ตารางที่ 18 ตัวอย่างรายละเอียดการพิจารณาปรับปรุงการควบคุมตามแนวทาง Reuse-Extend-Align

มิติ	Reuse	Extend	Align
access control	ใช้ access control เดิมขององค์กร	เพิ่มสิทธิ์เฉพาะสำหรับ AI tools, model endpoint, API key หรือ agent identity	บันทึกใน AI Inventory ว่า AI หรือ agent เข้าถึงระบบใดและทำ action อะไรได้
logging	ใช้ logging framework เดิม	เพิ่ม prompt log, output log, model version log, tool-call log หรือ action log สำหรับ Agentic AI	ใช้ log เหล่านี้เป็น evidence สำหรับ monitoring, incident review และ auditability
change management	ใช้ change management เดิม	เพิ่มการทบทวนเมื่อเปลี่ยนโมเดล เปลี่ยน prompt template เปลี่ยน data source เพิ่ม tool connector หรือเพิ่ม autonomy level	map เข้าหมวด AI Lifecycle/ System Lifecycle ของ AIMS
Vendor management	ใช้ supplier management เดิม	เพิ่มการประเมิน AI vendor เช่น model update, data usage, data opt-out, SLA, API latency และ incident notification	map เข้าหมวด Provider/External Relationship ของ AIMS

Principle reference: (Risk-based control selection; ISO/IEC 42001 SoA logic; reuse and extension of ISO/IEC 27001 controls; proportional governance)

หมวด E: แนวทางการใช้ข้อมูลส่วนบุคคลในระบบ AI

Q17: นอกเหนือจากความเสี่ยงที่รอบการกำกับดูแลฉบับนี้ครอบคลุมเกี่ยวกับข้อมูลส่วนบุคคลแล้ว ผู้ประกอบธุรกิจควรพิจารณาปัจจัยด้านการคุ้มครองข้อมูลส่วนบุคคลเพิ่มเติมอย่างไร เมื่อใช้งาน AI โดยเฉพาะ Generative AI?

A: กรอบการกำกับดูแลฉบับนี้ได้กำหนดหลักการและมาตรการควบคุมที่ครอบคลุมความเสี่ยงด้านข้อมูลส่วนบุคคลไว้ในระดับหนึ่ง ทั้งในส่วนของ AIMS Controls หมวด 6 (ข้อมูลสำหรับระบบ AI) หมวด 8 (ความสัมพันธ์และการสื่อสารกับบุคคลภายนอก) และปัจจัยส่งเสริมด้านความเป็นส่วนตัวและข้อมูลส่วนบุคคล (Enabler 2.5.3)

อย่างไรก็ดี การใช้งาน Generative AI อาจก่อให้เกิดประเด็นด้านการคุ้มครองข้อมูลส่วนบุคคลที่มีความเฉพาะเจาะจงมากกว่าระบบสารสนเทศทั่วไป เนื่องจากข้อมูลส่วนบุคคลอาจเข้าไปเกี่ยวข้องกับระบบ AI ได้ในหลายจุด ตั้งแต่ข้อมูลที่ใส่ฝึกโมเดล ข้อมูลที่ผู้ใช้กรอกใน prompt ข้อมูลในฐานความรู้องค์กร (RAG) ไปจนถึง output ที่ระบบสร้างขึ้นและ log ของการทำงานของ AI agent ผู้ประกอบธุรกิจจึงอาจต้องพิจารณาปัจจัยเพิ่มเติม เช่น

- การแจ้งวัตถุประสงค์เฉพาะเมื่อใช้ข้อมูลส่วนบุคคลกับระบบ AI (AI-specific notification)
- การวิเคราะห์ฐานกฎหมายที่เหมาะสมตาม PDPA สำหรับแต่ละ AI use case
- การจำแนกว่า prompt, input, output และ log ของระบบ AI อาจเป็นข้อมูลส่วนบุคคลได้
- การกำหนด retention policy ที่แยกตามชั้นข้อมูลของระบบ AI
- การเตรียมกระบวนการรองรับสิทธิของเจ้าของข้อมูลในบริบท AI เช่น การเข้าถึง แก้ไข ลบ หรือถอนความยินยอม
- ข้อจำกัดทางเทคนิคในการลบข้อมูลจากโมเดลที่ผ่านการฝึกแล้ว (machine unlearning)

ทั้งนี้ ผู้ประกอบธุรกิจสามารถศึกษาแนวทางเพิ่มเติมได้จากสำนักงานคณะกรรมการคุ้มครองข้อมูลส่วนบุคคล (สคส.) หรือหน่วยงานที่เกี่ยวข้อง เพื่อใช้ประกอบการพิจารณาควบคู่กับมาตรการควบคุมตามกรอบการกำกับดูแลฉบับนี้ตามความเหมาะสม
Principle reference: (PDPA พ.ศ. 2562; Singapore PDPC Advisory Guidelines on Use of Personal Data in Generative AI; ISO/IEC 42001 — Data for AI systems; AI-specific privacy governance)

หมวด F: สถานะของมาตรการควบคุมและความรับผิดชอบ

Q18: จำเป็นต้องนำมาตรการควบคุมทั้งหมดในภาคผนวก (ทั้ง developer, user และ customizer) ไปใช้ครบทุกข้อหรือไม่?

A: องค์กรสามารถพิจารณาเลือกนำมาตรการควบคุมในภาคผนวกที่จัดทำขึ้นเพื่อสนับสนุน มาตรฐานขั้นต่ำด้านการกำกับดูแล (minimum governance baseline) มิได้มีเจตนากำหนดระดับความสมบูรณ์หรือความเป็นเลิศด้าน AI governance ให้เท่ากันทุกกรณี องค์กรอาจต้องยกระดับมาตรการเพิ่มเติมเมื่อ AI มีอิทธิพลต่อการตัดสินใจหรือมีผลกระทบต่อผู้ใช้งานและตลาดทุนในระดับสูง

Principle reference: (Baseline vs. maturity distinction; ISO/IEC 42001; BIS & IOSCO technology risk discussions)

Q19: อะไรคือความแตกต่างระหว่างการเปิดเผยข้อมูลอย่างเหมาะสมกับ AI-washing?

A: การเปิดเผยข้อมูลอย่างเหมาะสม คือการอธิบายบทบาท ขอบเขต ข้อจำกัด และความเสี่ยงของการใช้ AI ตามข้อเท็จจริงด้วยภาษาที่เข้าใจได้ และมีหลักฐานรองรับอย่างเพียงพอ ขณะที่ AI-washing คือการกล่าวอ้างเกี่ยวกับ AI ในลักษณะที่เกินจริง หลอกลวง ไม่ครบถ้วน หรือทำให้เข้าใจผิดเกี่ยวกับความสามารถ ระดับการใช้งานจริง ประสิทธิภาพ หรือผลลัพธ์ของระบบ AI

Principle reference: (Transparency; AIMS 2.4.12; IOSCO Supervisory Toolkit)

หมวด G: ความคาดหวังด้านเอกสารและหลักฐานขั้นต่ำ

Q20: คู่มือฉบับนี้มีความคาดหวังขั้นต่ำต่อเอกสารหรือหลักฐานที่องค์กรควรจัดให้มีอย่างไร สำหรับ AI ที่มีระดับความเสี่ยงแตกต่างกัน?

A: คู่มือฉบับนี้มีความคาดหวังให้องค์กรสามารถจัดให้มี เอกสารหรือหลักฐานขั้นต่ำ (minimum evidence set) ที่เหมาะสมกับระดับความเสี่ยง วัตถุประสงค์ของการใช้งานและผลกระทบของการใช้งาน AI โดยมีได้กำหนดให้ทุกองค์กรหรือทุก use case ต้องมีเอกสารในระดับเดียวกัน ทั้งนี้ สามารถสรุปความคาดหวังเชิงหลักการได้ดังนี้

AI ความเสี่ยงต่ำ (เช่น การสนับสนุนการทำงานภายใน):

ควรมีเอกสารหรือหลักฐานพื้นฐาน เช่น AI policy ระดับองค์กร บัญชีรายชื่อ AI (AI inventory) แบบสรุปกติกาการใช้งาน (AI usage rules) และหลักฐานการให้ความรู้หรือการสื่อสารกับผู้ใช้งานภายใน

AI ความเสี่ยงปานกลาง (เช่น ระบบที่มีอิทธิพลต่อการตัดสินใจของผู้ใช้บริการ):

ควรเพิ่มเติมหลักฐานด้านการประเมินความเสี่ยง การกำหนด human oversight ที่เหมาะสม เอกสารอธิบายการทำงานหรือข้อจำกัดของระบบ (เช่น system description หรือ model card ในระดับที่เหมาะสม) รวมถึงแนวทางรับมือเหตุผิดปกติ

AI ความเสี่ยงสูงหรือมีผลกระทบต่อตลาดทุน (เช่น การให้คำแนะนำการลงทุนหรือการซื้อขายอัตโนมัติ):

ควรมีหลักฐานที่เข้มแข็งขึ้น เช่น การประเมินความเสี่ยงเชิงลึก การกำหนดบทบาทและความรับผิดชอบอย่างชัดเจน เอกสารอธิบายโมเดลและเหตุผลของการควบคุม การติดตามผลอย่างต่อเนื่อง และบันทึกการตัดสินใจเชิงกำกับดูแลเพื่อรองรับการตรวจสอบย้อนหลัง

ทั้งนี้ เอกสารหรือหลักฐานดังกล่าว ไม่จำเป็นต้องอยู่ในรูปแบบเดียวกันทุกองค์กร และองค์กรสามารถอ้างอิงเอกสารหรือกระบวนการที่มีอยู่เดิม หากสามารถแสดงให้เห็นได้ว่าครอบคลุมวัตถุประสงค์ด้านการกำกับดูแลและการจัดการความเสี่ยงตามระดับความเสี่ยงของ AI นั้นอย่างเหมาะสม

Principle reference: (Minimum governance baseline; proportional documentation; ISO/IEC 42001; risk-based evidence concept in NIST AI RMF และ IOSCO technology risk discussions)

Q21: ผู้ประกอบธุรกิจควรพิจารณาเรื่องตัวชี้วัดและการติดตามผลการทำงานของระบบ AI อย่างไร และสามารถเตรียมความพร้อมไว้หากมีการกำกับดูแลหรือตรวจสอบโดยหน่วยงานกำกับดูแลในอนาคต?

A: แนวทางการกำกับดูแลระดับสากล เช่น IOSCO Supervisory Toolkit for AI Use in Capital Markets (2026) ได้ระบุว่าการติดตามผลการทำงานของระบบ AI อย่างต่อเนื่องเป็นองค์ประกอบสำคัญของการกำกับดูแลที่มีประสิทธิภาพ โดยครอบคลุมทั้งตัวชี้วัดด้านคุณภาพผลลัพธ์ ด้านเหตุการณ์ผิดปกติ ด้านการพึ่งพาบุคคลที่ 3 และด้านระดับการกำกับดูแลโดยมนุษย์ โดยมีวัตถุประสงค์ให้หน่วยงานกำกับดูแลใช้เครื่องมือ คำถาม และข้อมูลตามระดับความเสี่ยงและความมีนัยสำคัญของ AI use case ทั้งนี้ กรอบการกำกับดูแลฉบับนี้ได้ออกแบบเครื่องมือที่รองรับการติดตามผลไว้แล้วในหลายส่วน ได้แก่

- AI Inventory (องค์ประกอบที่ 6) ที่มีช่องระบุ Review Frequency, Monitoring Focus และ Incident/Kill Switch
- AIMS Controls หมวด 5 (วงจรชีวิต) ที่กำหนดให้มี KPIs/KRIs เป็น Minimum Evidence
- AIMS Controls หมวด 7 (การใช้งาน) ที่กำหนดให้มี Usage dashboards & alerts
- ตัวอย่างการประยุกต์ใช้ในภาคผนวก 6 และ 7 ที่แสดงตัวอย่างตัวชี้วัดสำหรับ use case ทั้งความเสี่ยงต่ำและสูง

ผู้ประกอบธุรกิจจึงสามารถใช้เครื่องมือข้างต้นเป็นจุดเริ่มต้น โดยอาจไม่จำเป็นต้องกำหนดตัวชี้วัดทุกประเภทตั้งแต่เริ่มต้น แต่ควรเริ่มจากตัวชี้วัดพื้นฐานแล้วขยายตามความพร้อมขององค์กร เช่น ควรกำหนดอย่างน้อยตัวชี้วัดขั้นต่ำสำหรับทุก use case ตามระดับความเสี่ยง โดยกรณีระดับความเสี่ยงต่ำ อาจใช้ตัวชี้วัดพื้นฐาน ส่วนกรณีความเสี่ยงสูงควรกำหนดตัวชี้วัดเชิงลึก threshold และ trigger สำหรับการยกระดับ ทบทวนหรือหยุดใช้งานชั่วคราว

นอกจากนี้ ผู้ประกอบธุรกิจควรตระหนักว่าหน่วยงานกำกับดูแลอาจพิจารณาขอข้อมูลหรือหลักฐานเกี่ยวกับการใช้งาน AI ในลักษณะต่อไปนี้ เพื่อประกอบการกำกับดูแลหรือตรวจสอบ

- ภาพรวมของ use case ที่ใช้งานอยู่ รวมถึงสถานะ (ทดลอง/ใช้งานจริง/ยุติ) และระดับความเสี่ยง
- ระดับการพึ่งพาผู้ให้บริการภายนอก รวมถึง foundation model, cloud หรือ API ที่สำคัญ
- ตัวชี้วัดด้านคุณภาพผลลัพธ์ เหตุการณ์ผิดปกติ และระดับการแทรกแซงโดยมนุษย์
- หลักฐานการเปิดเผยข้อมูลต่อผู้ให้บริการ และกระบวนการทบทวนข้อความที่อ้างอิงถึง AI

อย่างไรก็ดี ข้อมูลส่วนใหญ่ข้างต้นสามารถจัดเตรียมได้จากเครื่องมือที่กรอบการกำกับดูแลฉบับนี้กำหนดไว้แล้ว

Principle reference: (IOSCO Supervisory Toolkit for AI Use in Capital Markets, 2026 - Table 7 and Section 3.5; NIST AI RMF - Measure Function)

Q22: องค์กรควรพิจารณาเรื่องข้อมูลอย่างไรเมื่อนำ AI มาใช้งาน?

A: ผู้ประกอบธุรกิจสามารถใช้กลไก Data Governance ที่มีอยู่ เช่น Data Inventory, Data Catalog, Metadata, ROPA, Data Flow, Data Classification และ Data Owner เป็นฐานในการกำกับดูแลข้อมูลที่ใช้กับ AI ได้ โดยไม่จำเป็นต้องสร้างทะเบียนหรือกระบวนการใหม่ซ้ำทั้งหมด อย่างไรก็ตาม การใช้ข้อมูลกับ AI ควรพิจารณาเพิ่มเติมว่าข้อมูลนั้นเหมาะสมกับวัตถุประสงค์ของ use case หรือไม่ มีคุณภาพเพียงพอหรือไม่ ใช้ตามสิทธิและข้อจำกัดที่เกี่ยวข้องหรือไม่ และมีมาตรการควบคุมเพียงพอสำหรับ input, prompt, output, log, RAG, training/testing data หรือ agentic workflow หรือไม่ ดังนั้น Data Governance ช่วยให้องค์กรรู้ว่าข้อมูลคืออะไร อยู่ที่ไหน ใช้อย่างไร และใครรับผิดชอบ ส่วน AI Governance ต้องทำให้องค์กรควบคุมการนำข้อมูลนั้นไปใช้โดย AI และรับผิดชอบต่อผลลัพธ์ที่ AI สร้างขึ้นได้ด้วย

Principle reference: (Data governance; data quality; purpose limitation; accountability; traceability; ISO/IEC 42001 - Data for AI systems)

Q23: หากใช้ AI ผ่าน cloud, model provider, AI platform หรือ API ภายนอก ควรพิจารณาอะไรเพิ่มเติม?

A: ผู้ประกอบธุรกิจควรพิจารณาความเสี่ยงจากการพึ่งพาผู้ให้บริการหรือ platform ภายนอกตามระดับความสำคัญของ use case โดยเฉพาะกรณีที่เกี่ยวข้องกับข้อมูลลูกค้า ข้อมูลอ่อนไหว ระบบงานสำคัญ การให้คำแนะนำ การซื้อขาย หรือกระบวนการที่มีผลกระทบสูง ประเด็นที่ควรพิจารณาอาจรวมถึงตำแหน่งการจัดเก็บหรือประมวลผลข้อมูล เขตอำนาจศาลที่เกี่ยวข้อง สิทธิการเข้าถึงของผู้ให้บริการ การควบคุมกุญแจหรือ credential การเปลี่ยนแปลง model หรือ service เงื่อนไขการใช้ข้อมูล การแจ้งเหตุผิดปกติ การยุติการให้บริการ การคืนหรือลบข้อมูล และกระบวนการ fallback หากบริการไม่พร้อมใช้งาน ซึ่งองค์กรอาจไม่ได้ควบคุมโมเดลหรือโครงสร้างพื้นฐานทั้งหมด แต่ยังคงควบคุมบริบท ข้อมูลที่ใช้ ผลลัพธ์ที่นำไปใช้ และความเสี่ยงจากการพึ่งพาผู้ให้บริการ

Principle reference: (Third-party risk management; operational resilience; data residency; access control; accountability; minimum sufficient control)

Q24: ผู้ประกอบธุรกิจควรพิจารณาองค์ประกอบใดบ้างในการจัดให้มี runtime governance สำหรับ Agentic AI ที่มีผลกระทบสูง?

A: ผู้ประกอบธุรกิจควรพิจารณา 4 องค์ประกอบสำคัญ ในการจัดให้มี runtime governance สำหรับ Agentic AI ที่ริเริ่ม action ซึ่งกระทบต่อผู้ลงทุน ตลาด หรือบริการสำคัญ ได้แก่

- กลไกยืนยันตัวตนของระบบ AI (Agent Identity Verification) - ระบบ AI ต้องมีรหัสตัวตนที่ตรวจสอบได้ ผูกกับผู้รับผิดชอบและขอบเขตอำนาจดำเนินการ ก่อน execute action ทุกครั้ง
- การบันทึกข้อมูลประกอบก่อนดำเนินการ (Pre-execution Record) - ระบบต้องบันทึก action ที่จะทำ ขั้นตอนการวิเคราะห์ เครื่องมือและข้อมูลที่ใช้ พร้อม context ที่เกี่ยวข้อง ก่อนส่งไปยังระบบภายนอก
- การตัดสินใจก่อนดำเนินการ (Pre-execution Disposition) - ทุก action ต้องได้ผลตัดสินใจชัดเจน เช่น ดำเนินการอัตโนมัติ/เฝ้าติดตาม/ส่งให้มนุษย์อนุมัติ/ปฏิเสธ ตามระดับความเสี่ยงของแต่ละ action โดยองค์กรสามารถเลือกใช้ตามความเหมาะสมของ design และ risk profile ของ use case
- การบันทึกเหตุการณ์ที่ตรวจสอบย้อนหลังได้ (Runtime Audit Log) - บันทึกการตัดสินใจและการดำเนินการทุก action ในรูปแบบที่แก้ไขไม่ได้ (append-only) และเก็บในระบบที่แยกจากการควบคุมของ agent เอง เพื่อรองรับ compliance, audit และ incident review

ทั้งนี้ องค์ประกอบดังกล่าวเป็น เพียงส่วนหนึ่งของการกำกับดูแลทั้งหมด ผู้ประกอบธุรกิจยังคงต้องมีมาตรการอื่นตามกรอบการกำกับดูแลฉบับนี้ ซึ่ง runtime governance เป็นเพียง layer เสริมที่ทำงานระหว่าง pre-deployment governance และ post-deployment monitoring

Principle reference: (MAS SAFR (Agent Identity, Governance Envelope, Disposition Engine, Audit Log); Runtime governance layer; Pre-execution assurance; Tamper-evident logging)

ภาคผนวก 10 ตัวอย่างตัวชี้วัดและข้อมูลสำหรับการกำกับติดตามการใช้งาน AI

ภาคผนวกนี้มีวัตถุประสงค์เพื่อช่วยให้ผู้ประกอบธุรกิจสามารถกำหนดตัวชี้วัดขั้นต่ำสำหรับการติดตาม ประสิทธิภาพ ความเสี่ยง และความพร้อมใช้งานของระบบ AI ได้อย่างสอดคล้องกับลักษณะของ use case โดยผู้ประกอบธุรกิจอาจปรับตัวชี้วัดตามบริบทของตนได้ แต่ควรจัดให้มีตัวชี้วัดพื้นฐานเพียงพอที่จะสนับสนุน การติดตามอย่างต่อเนื่อง การยกระดับการทบทวน และการรายงานต่อผู้บริหารหรือหน่วยงานกำกับดูแลเมื่อจำเป็น

ตารางที่ 19 ตัวอย่างตัวชี้วัดและข้อมูลสำหรับการกำกับติดตามการใช้งาน AI

หมวดตัวชี้วัด	ตัวอย่างตัวชี้วัด	วัตถุประสงค์ในการติดตาม
การนำไปใช้งาน	จำนวน use case ที่อยู่ระหว่างทดลอง/ใช้งานจริง/ ยุติการใช้งาน	ใช้ติดตามระดับการใช้งานและความพร้อมขององค์กร
การพึ่งพาบุคคลที่ 3	สัดส่วน use case ที่ใช้ผู้ให้บริการภายนอก, foundation model, cloud หรือ external API	ใช้ประเมินความเสี่ยงจาก outsourcing และ concentration
คุณภาพผลลัพธ์	accuracy, error rate, false positive/false negative หรือ quality score อื่นที่เหมาะสม	ใช้ติดตามความเหมาะสมของผลลัพธ์และการเชื่อมโยงประสิทธิภาพ
ความเสี่ยงเชิงพฤติกรรม	drift rate, override rate, escalation rate, จำนวนครั้งที่เกิน threshold	ใช้ติดตามสัญญาณเสี่ยงที่ต้องทบทวน use case
เหตุการณ์ผิดปกติ	จำนวน AI incident, near miss, data leakage, system outage, model failure	ใช้ติดตามความรุนแรงและความถี่ของเหตุการณ์
ผลกระทบต่อผู้ใช้บริการ	จำนวนข้อร้องเรียน, จำนวนลูกค้าที่ได้รับผลกระทบ, ระยะเวลาฟื้นฟูบริการ	ใช้ประเมินผลกระทบและจัดลำดับการแก้ไข
ความเป็นธรรมและความโปร่งใส	bias indicator, explainability score, จำนวนคำขอคำอธิบายจากลูกค้า	ใช้ติดตามประเด็น fairness, accountability และ transparency

แนวทางการใช้งานตัวชี้วัด

- 1) ควรกำหนดเจ้าของตัวชี้วัด (metric owner) และกำหนดรอบการรายงานที่เหมาะสม เช่น รายเดือน รายไตรมาส หรือ เมื่อเกิดเหตุการณ์
- 2) use case ที่มีความเสี่ยงสูงควรกำหนด threshold และเกณฑ์ trigger สำหรับการยกระดับการทบทวน หรือการหยุดใช้งานชั่วคราว
- 3) ตัวชี้วัดควรเชื่อมโยงกับ AI Inventory, incident management และรายงานต่อผู้บริหาร เพื่อให้เห็นภาพรวมขององค์กรอย่างต่อเนื่อง

สำหรับผู้ประกอบธุรกิจแต่ละประเภท เช่น ตัวแทนซื้อขายหลักทรัพย์ ผู้จัดการกองทุน และตลาดหลักทรัพย์ อาจพิจารณา กำหนดตัวชี้วัดเพิ่มเติมที่สะท้อนความเสี่ยงเฉพาะของกลุ่มธุรกิจ เช่น ตัวชี้วัดด้านความเหมาะสมของคำแนะนำ (suitability) ความเอนเอียงของผลลัพธ์ (bias/fairness indicators) ความถูกต้องของการประเมินมูลค่าสินทรัพย์ การเบี่ยงเบนจากนโยบายการลงทุน (mandate drift) คุณภาพตลาดและสภาพคล่อง (market quality measures) ประสิทธิภาพของระบบเฝ้าระวัง (surveillance effectiveness) หรือความล้มเหลวในกระบวนการหลังการซื้อขาย (settlement failures) ตามความเหมาะสม

ภาคผนวก 11 ตัวอย่างแบบรายงานเหตุการณ์ที่เกี่ยวข้องกับ AI

ภาคผนวกนี้จัดทำเป็นตัวอย่างสำหรับใช้ภายในองค์กรในการรวบรวมข้อเท็จจริง เมื่อเกิดเหตุการณ์ผิดปกติที่เกี่ยวข้องกับ AI และยังสามารถใช้เป็นข้อมูลพื้นฐานในการรายงานต่อสำนักงาน ก.ล.ต. ตามหลักเกณฑ์และช่องทางที่สำนักงาน ก.ล.ต. กำหนดได้

ตารางที่ 20 ตัวอย่างรายละเอียดรายงานเหตุการณ์ที่เกี่ยวข้องกับ AI และข้อมูลที่ต้องจัดเก็บ

ส่วนของรายงาน	ข้อมูลที่ต้องจัดเก็บ
ส่วนที่ 1 ข้อมูลทั่วไป	ชื่อองค์กร, วันและเวลาที่พบเหตุการณ์, ผู้รายงาน, ชื่อ use case/ระบบ AI, ระดับความสำคัญของบริการ
ส่วนที่ 2 ลักษณะเหตุการณ์	ประเภทเหตุการณ์ เช่น hallucination, data leakage, model drift, outage, prompt injection, unauthorized action, third-party failure
ส่วนที่ 3 ผลกระทบ	จำนวนผู้ใช้บริการหรือธุรกรรมที่ได้รับผลกระทบ, ผลกระทบต่อข้อมูลส่วนบุคคล, ผลกระทบต่อการซื้อขาย/การให้บริการ, ผลกระทบต่อความต่อเนื่องทางธุรกิจ
ส่วนที่ 4 สาเหตุเบื้องต้น	สาเหตุที่คาดหมาย แหล่งที่มาของปัญหา ความเกี่ยวข้องกับบุคคลที่ 3/ผู้ให้บริการ
ส่วนที่ 5 มาตรการตอบสนอง	การหยุดหรือจำกัดวงการใช้งาน, fallback/rollback, การแจ้งลูกค้า, การประสานผู้ให้บริการ, การเก็บพยานหลักฐาน
ส่วนที่ 6 แผนฟื้นฟูและป้องกันซ้ำ	เวลาที่คาดว่าจะฟื้นฟูเสร็จ มาตรการระยะสั้น มาตรการระยะยาว เจ้าของแผน และกำหนดแล้วเสร็จ

เกณฑ์ที่องค์กรควรพิจารณาว่าเป็นเหตุการณ์ที่มีนัยสำคัญ

- 1) มีผลกระทบต่อผู้ลงทุนหรือผู้ใช้บริการจำนวนมาก หรือมีผลกระทบต่อสิทธิ/ผลประโยชน์ของลูกค้าอย่างมีนัยสำคัญ
- 2) ทำให้เกิดการให้คำแนะนำ การเปิดเผยข้อมูล หรือการดำเนินการ ที่ไม่ถูกต้องอันอาจก่อให้เกิดความเสียหาย หรือความเข้าใจผิด
- 3) เกี่ยวข้องกับข้อมูลส่วนบุคคล ข้อมูลสำคัญ หรือความมั่นคงปลอดภัยไซเบอร์
- 4) เกิดจาก หรือเชื่อมโยงกับความขัดข้องของบุคคลที่ 3 หรือผู้ให้บริการภายนอกที่สำคัญ ผู้ให้บริการคลาวด์ หรือ foundation model ที่องค์กรพึ่งพา
- 5) เหตุการณ์ดังกล่าวมีแนวโน้มกระทบต่อความต่อเนื่องของบริการ สภาพคล่องของตลาด หรือความเชื่อมั่นในวงกว้าง

ภาคผนวก 12 ตัวอย่างแบบประเมินผู้ให้บริการ AI และความเสี่ยงจากการกระจุกตัว

ภาคผนวกนี้มีวัตถุประสงค์เพื่อใช้เป็นตัวอย่างรายการตรวจสอบขั้นต่ำในการคัดเลือก และติดตามผู้ให้บริการ AI รวมถึงการประเมินความเสี่ยงจากการพึ่งพาผู้ให้บริการรายเดียวหรือกลุ่มเดียวกันมากเกินไป

ตารางที่ 21 ตัวอย่างรายละเอียดการประเมินผู้ให้บริการ AI และความเสี่ยงจากการกระจุกตัว

หัวข้อ	รายการประเมิน
ความโปร่งใสของบริการ	ผู้ให้บริการอธิบายขอบเขตการทำงาน ข้อจำกัด วิธีอัปเดต และการเปลี่ยนแปลงโมเดลได้เพียงใด
ความมั่นคงปลอดภัยไซเบอร์	มีมาตรการ IAM, encryption, logging, vulnerability management, incident response หรือไม่
ข้อมูลและสิทธิทางกฎหมาย	ข้อมูลขององค์กรจะถูกนำไป train model หรือไม่, มีข้อกำหนดด้าน data residency, IP, confidentiality และ deletion อย่างไร
การตรวจสอบและการแจ้งเหตุการณ์	องค์กรมี audit right หรือได้รับผลการประเมินอิสระหรือไม่, ผู้ให้บริการมี SLA และข้อกำหนดการแจ้งเหตุการณ์อย่างไร
ความยืดหยุ่นและการย้ายระบบ	มีทางเลือกทดแทนหรือไม่, export/portability ทำได้เพียงใด, มี exit plan หรือไม่
ความเสี่ยงจากการกระจุกตัว	มี use case สำคัญที่รายการที่พึ่งพา ผู้ให้บริการรายนี้หรือรายในกลุ่มเดียวกัน, หากหยุดให้บริการจะส่งผลอย่างไร

แนวทางการพิจารณาความเสี่ยงจากการกระจุกตัว

- 1) ให้พิจารณาทั้ง direct dependency และ indirect dependency เช่น การใช้ผู้ให้บริการแอปพลิเคชันหลายราย แต่ทั้งหมด operate หรืออยู่บนโครงสร้างพื้นฐานคลาวด์รายเดียวกัน
- 2) สำหรับ use case ที่มีความสำคัญสูง ควรมีแผนย้ายระบบ หรือแผนปฏิบัติการเมื่อผู้ให้บริการหยุดให้บริการ หรือเกิดเหตุขัดข้องเป็นเวลานาน
- 3) ควรทบทวนความเสี่ยงจากการกระจุกตัวอย่างน้อยปีละ 1 ครั้ง หรือเมื่อมีการเพิ่ม use case สำคัญใหม่อย่างมีนัยสำคัญ

ภาคผนวก 13 ตัวอย่างแนวทางการเปิดเผยการใช้ AI และการป้องกัน AI-washing

ภาคผนวกนี้จัดทำขึ้นเพื่อเป็นแนวทางให้ผู้ประกอบธุรกิจสามารถเปิดเผยข้อมูลเกี่ยวกับการใช้ AI ได้อย่างเหมาะสมกับบริบทของบริการและป้องกันการสื่อสารเกินจริง หรือทำให้เข้าใจผิด

หลักการเปิดเผยข้อมูลขั้นต่ำ

- 1) เมื่อลูกค้าหรือผู้ใช้บริการกำลังโต้ตอบกับระบบ AI ควรเปิดเผยข้อมูลอย่างเหมาะสมว่ามีการใช้ AI ในกระบวนการหรือบริการนั้น
- 2) การเปิดเผยควรอธิบายบทบาทของ AI ในระดับที่เพียงพอ เช่น AI ทำหน้าที่สร้างข้อเสนอเบื้องต้น สนับสนุนการวิเคราะห์ หรือมีผลต่อการตัดสินใจระดับใด
- 3) ควรเปิดเผยข้อจำกัด ความเสี่ยง หรือความไม่แน่นอนที่มีนัยสำคัญ เช่น hallucination, ความคลาดเคลื่อนของผลลัพธ์, การพึ่งพาบุคคลที่ 3 หรือข้อจำกัดของข้อมูล
- 4) หากผู้ใช้บริการสามารถขอให้บุคคลเข้ามาทบทวน หรือให้คำอธิบายเพิ่มเติมได้ ควรเปิดเผยช่องทางดังกล่าวอย่างชัดเจน

ตัวอย่างพฤติกรรมที่ควรหลีกเลี่ยงเพื่อป้องกัน AI-washing

- 1) กล่าวอ้างว่าใช้ “AI ขั้นสูง” หรือ “AI อัจฉริยะ” โดยไม่สามารถอธิบายขอบเขตการใช้งานที่แท้จริง หรือไม่มีหลักฐานรองรับว่าระบบดังกล่าวถูกใช้งานจริง
- 2) สื่อสารในลักษณะที่ทำให้ผู้ลงทุนเข้าใจว่าการตัดสินใจทั้งหมดเกิดจาก AI ทั้งที่ในทางปฏิบัติ เป็นเพียงระบบสนับสนุนข้อมูลเบื้องต้น
- 3) เน้นเฉพาะประโยชน์ของ AI โดยไม่เปิดเผยข้อจำกัด ความเสี่ยง หรือเงื่อนไขสำคัญ ที่มีผลต่อความเข้าใจของลูกค้า
- 4) อ้างผลการทดสอบ ประสิทธิภาพ หรือความแม่นยำโดยไม่มีวิธีการวัดที่ตรวจสอบได้ หรือไม่สอดคล้องกับบริบทการใช้งาน

ภาคผนวก 14 ตัวอย่างคำถามสำหรับการกำกับติดตาม การทบทวนภายใน และการเตรียมความพร้อมสำหรับ supervisory dialogue

ภาคผนวกนี้สามารถใช้เป็นรายการคำถามเบื้องต้นให้แก่ผู้บริหาร ฝ่ายกำกับดูแล ฝ่ายบริหารความเสี่ยง ฝ่ายตรวจสอบภายใน และเจ้าของ use case เพื่อทบทวนความพร้อมขององค์กร และใช้เตรียมเอกสารเมื่อมีการหารือกับหน่วยงานกำกับดูแล

ตารางที่ 22 ตัวอย่างรายละเอียดประเด็นคำถามสำหรับทบทวนความพร้อมสำหรับ supervisory dialogue

มิติ	ประเด็นคำถาม
ด้านการกำกับดูแล	• องค์กรได้กำหนด AI strategy, risk appetite และเกณฑ์การยกระดับ use case ที่มีความเสี่ยงสูงไว้อย่างชัดเจน หรือไม่ • มีการรายงาน use case สำคัญ เหตุการณ์ผิดปกติ และความเสี่ยงจากบุคคลที่ 3 ต่อผู้บริหารระดับสูง และคณะกรรมการอย่างสม่ำเสมอ หรือไม่ • มีการกำหนด RACI หรือเอกสารเทียบเท่าสำหรับ Developer, Customizer, User และผู้มีส่วนได้เสียที่เกี่ยวข้อง ครบถ้วนเพียงใด
ด้านข้อมูลและโมเดล	• องค์กรสามารถอธิบายแหล่งที่มาของข้อมูล การเตรียมข้อมูล การเลือกโมเดล และข้อจำกัดที่สำคัญของระบบได้เพียงใด • มีการทดสอบ bias, hallucination, drift, robustness และความเหมาะสมก่อนใช้งานจริง และระหว่างการใช้งานอย่างไร • มีเกณฑ์สำหรับ rollback, suspension หรือ retraining เมื่อผลลัพธ์ต่ำกว่าเกณฑ์ หรือเกิดเหตุการณ์ผิดปกติ หรือไม่
ด้านบุคคลที่ 3 หรือผู้ให้บริการภายนอก	• มี due diligence ที่ครอบคลุมประเด็น transparency, security, subcontracting, audit right และ exit plan หรือไม่ • องค์กรได้ประเมินความเสี่ยงจากการกระจุกตัว (concentration) ของผู้ให้บริการในระดับ use case และระดับองค์กรแล้ว หรือยัง • เมื่อผู้ให้บริการมีการเปลี่ยนโมเดลหรือเงื่อนไขบริการ องค์กรมีกระบวนการประเมินผลกระทบ และขอดำเนินการอนุมัติใหม่ หรือไม่
ด้านการเปิดเผยและการดูแลผู้ให้บริการ	• ผู้ใช้บริการทราบหรือไม่ว่ากำลังได้พบกับระบบ AI และเข้าใจขอบเขตของบทบาท AI อย่างเพียงพอหรือไม่ • องค์กรมีหลักเกณฑ์ทบทวนข้อความทางการตลาด หรือข้อความสาธารณะที่อ้างถึง AI เพื่อป้องกัน AI-washing หรือไม่ • ลูกค้ามีช่องทางขอคำอธิบาย ร้องเรียน หรือขอให้บุคคลทบทวนผลลัพธ์จาก AI ได้เพียงใด
ด้านเหตุการณ์และการติดตาม	• องค์กรได้กำหนดนิยาม AI incident, significant AI incident และเกณฑ์การรายงานต่อหน่วยงานกำกับดูแลไว้ หรือไม่ • มีตัวชี้วัดขั้นต่ำและ threshold สำหรับ use case สำคัญ พร้อมหลักฐานการติดตามอย่างต่อเนื่องหรือไม่ • มีการจัดทำ post-incident review และนำบทเรียนที่ได้ไปปรับปรุง controls, SOP และ training หรือไม่

บรรณานุกรม

- Bank for International Settlements (BIS). Sound Practices on the Implications of AI and Digital Innovation for Financial Services. 2024. <https://www.bis.org>
- De Nederlandsche Bank. General principles for the use of Artificial Intelligence in the financial sector. 2019. <https://www.dnb.nl/media/jkbip2jc/general-principles-for-the-use-of-artificialintelligence-in-the-financial-sector.pdf>
- Infocomm Media Development Authority (IMDA) & AI Verify Foundation. Model AI Governance Framework for Generative AI: Fostering a Trusted Ecosystem - Companion Guide on Agentic AI. Singapore, 2026. <https://aiverifyfoundation.sg>
- International Business Machines Corporation (IBM). Agentic AI Governance Playbook: Managing Identity, Action, and Accountability for Autonomous AI Agents. IBM Institute for Business Value, 2026. <https://www.ibm.com/thought-leadership/institute-business-value>
- International Organization for Standardization/International Electrotechnical Commission. "Information technology - Artificial intelligence - Management system," ISO/IEC 42001:2023.
- International Organization for Standardization/International Electrotechnical Commission, "Information technology - Artificial intelligence - Guidance on risk management," ISO/IEC 23894:2023.
- J. Gama, I. Žliobaite, A. Bifet, M. Pechenizkiy, and A. Bouchachia, "A Survey on Concept Drift Adaptation," ACM Computing Surveys, 2014, vol. 46, no. 4, pp. 1–37.
- National Institute of Standards and Technology (NIST), "Artificial Intelligence Risk Management Framework (AI RMF 1.0)," NIST AI 100-1. 2023. <https://doi.org/10.6028/NIST.AI.100-1>
- M. Mitchell et al., "Model Cards for Model Reporting," in Proc. Conf. Fairness, Accountability, and Transparency (FAT '19), Atlanta, GA, USA, 2019, pp. 220–229.
- Monetary Authority of Singapore. Principles to Promote Fairness, Ethics, Accountability and Transparency (FEAT) in the Use of Artificial Intelligence and Data Analytics in Singapore's Financial Sector. 2018. <https://www.mas.gov.sg/publications/monographs-or-information-paper/2018/FEAT>
- Monetary Authority of Singapore. Safeguards for Agentic Finance at Runtime (SAFR) [White paper]. 2026. <https://www.mas.gov.sg/-/media/mas-media-library/development/fintech/ai-safr/safr.pdf>
- Office of National Higher Education Science Research and Innovation Policy Council (NXPO), AIS Academy, & IRIS Consulting. Thailand AI Readiness Index (TARI). 2026. <https://thailand-tari.ai>

- Organisation for Economic Co-operation and Development (OECD). OECD Principles on Artificial Intelligence. 2019 (updated 2024). <https://oecd.ai/en/ai-principles>
- Privacy Impact Assessment: Claude Opus 4.7 - A principles-based assessment under GDPR, the EU AI Act, and the AI Governance Stack framework. Digital 520, 2026. <https://www.digital520.com>
- The High-Level Expert Group on Artificial Intelligence (set up by the European Commission). Ethics guidelines for trustworthy AI. 2019. <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>
- The International Organization of Securities Commissions (IOSCO). The use of artificial intelligence and machine learning by market intermediaries and asset managers. 2021. <https://www.iosco.org/library/pubdocs/pdf/IOSCOPD684.pdf>
- The International Organization of Securities Commissions (IOSCO). Operational Resilience, Outsourcing and Third-Party Risks in Financial Market Infrastructures. 2024. <https://www.iosco.org>
- The International Organization of Securities Commissions (IOSCO), "Artificial Intelligence in Capital Markets: Use Cases, Risks, and Challenges (CR/01/2025)," 2025. <https://www.iosco.org/library/pubdocs/pdf/IOSCOPD788.pdf>
- The International Organization of Securities Commissions (IOSCO). Supervisory Toolkit for AI Use in Capital Markets. Final Report (FR/02/2026), 2026. <https://www.iosco.org/library/pubdocs/pdf/IOSCOPD822.pdf>
- ธนาคารแห่งประเทศไทย. แนวนโยบายธนาคารแห่งประเทศไทย เรื่อง การบริหารจัดการความเสี่ยงของการใช้งานระบบปัญญาประดิษฐ์. 2568. <https://www.bot.or.th/content/dam/bot/fipcs/documents/FOG/2568/ThaiPDF/25680178.pdf>
- สำนักงานคณะกรรมการการรักษาความมั่นคงปลอดภัยไซเบอร์แห่งชาติ. แนวปฏิบัติการใช้ปัญญาประดิษฐ์อย่างมั่นคงปลอดภัย (AI Security Guidelines). 2568. <https://www.ncsa.or.th/news/28559b1b363037fb2d000158?category=news>
- สำนักงานคณะกรรมการกำกับหลักทรัพย์และตลาดหลักทรัพย์. หลักและแนวปฏิบัติตามนโยบายธรรมาภิบาลข้อมูล (Principles and Guidelines for Data Governance). 2569.
- สำนักงานคณะกรรมการกำกับและส่งเสริมการประกอบธุรกิจประกันภัย. แนวปฏิบัติ เรื่องการกำกับดูแลการใช้งานปัญญาประดิษฐ์ สำหรับบริษัทประกันภัย (AI Governance Guideline) พ.ศ. 2568. 2568. <https://www.oic.or.th/th/rm-survey/>
- สำนักงานคณะกรรมการดิจิทัลเพื่อเศรษฐกิจและสังคมแห่งชาติ. เอกสารแนวปฏิบัติจริยธรรม ปัญญาประดิษฐ์. 2564. <https://bact.cc/f/2022/11/202012-thailand-ai-ethics-guideline-mdes.pdf>

สำนักงานพัฒนาธุรกรรมทางอิเล็กทรอนิกส์. แนวทางประยุกต์ใช้ Generative AI อย่างมีธรรมาภิบาลสำหรับองค์กร (Generative AI Governance Guideline for Organizations). 2567.

https://www.etda.or.th/th/Useful-Resource/GovernanceGuideline_v1.aspx

หมายเหตุเกี่ยวกับการใช้ AI ในการจัดทำรอบการกำกับดูแลฉบับนี้

ในการจัดทำรอบการกำกับดูแลฉบับนี้ อาจมีการใช้เครื่องมือปัญญาประดิษฐ์เพื่อสนับสนุนการวิเคราะห์ การสังเคราะห์ข้อมูล การเปรียบเทียบแนวทางสากล และการตรวจสอบความสอดคล้องของโครงสร้างเนื้อหา อย่างไรก็ตาม การกำหนดทิศทางเชิงนโยบาย การเลือกแนวคิด การตัดสินใจเชิงดุลยพินิจ การประเมินความเหมาะสมกับบริบทตลาดทุนไทย และการรับรองเนื้อหาขั้นสุดท้ายยังคงอยู่ภายใต้การพิจารณาและความรับผิดชอบของมนุษย์

การใช้ AI ในลักษณะดังกล่าวจึงเป็นการใช้เพื่อสนับสนุนกระบวนการทำงาน ไม่ใช่การมอบอำนาจให้ AI เป็นผู้ตัดสินใจเชิงนโยบาย ทั้งนี้ แนวทางดังกล่าวสอดคล้องกับหลักการของกรอบฉบับนี้ที่เน้นการใช้ AI โดยมีมนุษย์กำกับดูแล มีการบันทึกเหตุผลตรวจสอบย้อนหลังได้ และคำนึงถึงความเสี่ยงจาก automation bias หรือการพึ่งพา AI โดยไม่ใช้วิจารณญาณอย่างเพียงพอ



สำนักงานคณะกรรมการกำกับหลักทรัพย์และตลาดหลักทรัพย์
333/3 ถนนวิภาวดีรังสิต แขวงจอมพล เขตจตุจักร กรุงเทพมหานคร 10900
www.sec.or.th