

รู้เท่าทัน "AI Deepfake" ด้วยหลักโยนิโสมนสิการ

AWARENESS OF "AI DEEPPAKES" USING THE PRINCIPLES OF YONISOMANASIKARA

ศุภกาญจน์ ชยาพัฒน์¹, และ พุไทย วันหากิจ

Supakanjana Jayapattana¹ and Puthai Wanhakit

คณะมนุษยศาสตร์และสังคมศาสตร์ มหาวิทยาลัยราชภัฏพระนคร

Faculty of Humanities and Social Sciences Phranakhon Rajabhat University,

Email: supakanjana.v@pnru.ac.th¹

Received 16 April 2024; Revised 7 August 2024; Accepted 7 August 2024.



บทคัดย่อ

บทความวิชาการนี้มีวัตถุประสงค์เพื่อศึกษาธรรมชาติของ AI Deepfakes และเพื่อเสนอแนวทางคำอธิบายหลักโยนิโสมนสิการประยุกต์ใช้ให้รู้เท่าทันกับ AI Deepfakes จากการศึกษาพบว่า AI Deepfakes เป็น Generative AI รูปแบบหนึ่งที่สามารถสร้างสื่อสังเคราะห์ ทั้งภาพนิ่ง เสียง ภาพเคลื่อนไหวพร้อมลอกเลียนเสียงจากมนุษย์จริงมาพูดคุยจริง ๆ จนแทบแยกไม่ออกซึ่งปัญญาประดิษฐ์สามารถมีระบบการเรียนรู้ หรือ Deep Learning จากลักษณะภายนอกของบุคคล เช่น สีมืด ตา ปาก จมูก รูปลักษณ์ต่าง ๆ ด้วยความฉลาดล้ำจากการเรียนรู้ข้อมูลมหาศาลของ Generative AI

AI Deepfakes มีทั้งข้อดีและข้อเสีย หรือ ดาบสองคม จึงมีความจำเป็นต้องรู้เท่าทันเพื่อนำไปใช้ประโยชน์อย่างแท้จริงซึ่งไม่มีผลเสียตามมา โดยเครื่องมือที่รองรับ AI Deepfakes คือ โยนิโสมนสิการหรือการทำในใจโดยแยบคาย เป็นหลักพุทธธรรมที่สามารถถ่วงถ่วงและตรวจสอบให้มนุษย์ใช้ทักษะการคิดด้วยกระบวนการของโยนิโสมนสิการ 10 รูปแบบ ได้แก่ 1.วิธีคิดแบบสืบสาวเหตุปัจจัย 2. วิธีคิดแบบแยกแยะส่วนประกอบ 3.วิธีคิดแบบสามัญลักษณะ 4.วิธีคิดแบบอริยสัจ/คิดแบบแก้ปัญหา 5. วิธีคิดแบบบรรณธรรมสัมพันธ์ 6.วิธีคิดแบบเห็นคุณค่าและทางออก 7. วิธีคิดแบบรู้คุณค่าแท้ – คุณค่าเทียม 8. วิธีคิดแบบรู้คุณธรรม 9. วิธีคิดแบบอยู่กับปัจจุบัน และ 10. วิธีคิดแบบวิวิธภาพ ก่อนที่จะตัดสินใจนำไปใช้ให้ถูกต้องตามทำนองคลองธรรม

โยนิโสมนสิการเป็นปัจจัยภายในของสัมมาทิฐิจำเป็นต้องอาศัยเหตุปัจจัยให้เกิดขึ้นพร้อม ได้แก่ สติสัมปชัญญะ บุรณการกบองศ์ธรรม 4 (เหตุผล แยกแยะ ความจริง และคุณค่า) จึงจะสามารถรู้เท่าทัน AI Deepfakes ได้ ทั้งนี้โยนิโสมนสิการเป็นเพียงระเบียบความคิดให้เป็นระบบ เพื่อให้ตัวปัญญาได้ตัดให้เด็ดขาดตกลงใจโดยคุชฎี หรือให้ปัญญาได้แก้ไขปัญหา AI Deepfakes ให้เด็ดขาดแน่นอน จึงจะสามารถได้คำตอบถูกต้องเหมาะสมและเกิดประโยชน์ได้อย่างแท้จริง

คำสำคัญ : AI Deepfakes; โยนิโสมนสิการ; รู้เท่าทัน

Abstract

The purpose of this academic article is to study the nature of AI Deepfakes and to propose a guideline to explain the principles of Yonisomanasikāra and its application to be aware of AI Deepfakes. From the study, it was found that AI Deepfakes is a form of Generative AI that can create synthetic media, including still images, sound, and moving images, along with imitating the voices of real humans to talk to the point that they are almost indistinguishable. Artificial intelligence can have a learning system or Deep Learning from a person's external characteristics such as skin colour, eyes, mouth, nose, and various appearances with the intelligence of learning from the enormous data of Generative AI.

AI Deepfakes have both advantages and disadvantages or a double-edged sword. Therefore, it is necessary to be aware of them in order to put them to good use without any negative consequences truly. The tool that supports AI Deepfakes is Yonisomanasikāra or doing things in the mind in a subtle manner. It is a principle of Buddhism that can be distilled and examined for humans to use their thinking skills using the 10 Yonisomanasikāra processes, namely: (1) the method of thinking by investigating causes and factors, (2) the method of thinking by separating components, (3) common thinking method, (4) Four Noble Truths/problem-solving thinking method, (5) thematic relationship thinking method, (6) a way of thinking that sees benefits and solutions, (7) a way of thinking that knows real value - artificial value, (8) a way of thinking that encourages morality, (9) the present-based thinking method, and (10) a way of thinking that is Vibhajjavāda before deciding to use it correctly according to morals.

Yonisomanasikāra is an internal factor of the right view. It requires factors to occur together, including mindfulness and clear comprehension integrating with the 4 Dharma elements (reason, discrimination, truth, and value) in order to be able to understand AI Deepfakes. However, Yonisomanasikāra is just a method of thinking system so that the wisdom can cut it off completely, decide it with pleasure, or let wisdom solve the AI Deepfakes problem for sure. So you will be able to get correct, appropriate, and truly beneficial answers.

Keywords: AI Deepfake Yonisomanasikāra Awareness

บทนำ

Deepfake เกิดจากคำ 2 คำรวมกัน ได้แก่ Deep Learning และ Fake โดย Deep Learning คือการเรียนรู้ ชุดตรรกะของระบบปัญญาประดิษฐ์หรือ AI ที่สามารถเรียนรู้จากสิ่งเล็กๆ ๑ เข้าใจง่ายไปจนถึงเรื่องที่ซับซ้อน หรือ ยากที่จะเข้าใจ ส่วนคำว่า Fake แปลว่าปลอม หรือ เหมือนจริง ดังนั้น วิเคราะห์ได้ว่า Deepfake จัดอยู่ในประเภทของปัญญาประดิษฐ์หรือ AI ระดับ Machine Learning ซึ่งเป็น ปัญญาประดิษฐ์ที่เรียนรู้เชิงลึกแบบเหมือนจริง (Reinforcement Learning) ซึ่งถือได้ว่า ระดับ Deepfakes มีการเรียนรู้และพัฒนาเทคโนโลยีปัญญาประดิษฐ์ ระดับสูงกว่าระดับพื้นฐาน โดยมีลำดับพัฒนาการประเภทปัญญาประดิษฐ์ ได้แก่

ปัญญาประดิษฐ์แบบแขนงคอมพิวเตอร์ทั่วไป (General Artificial Intelligence) ปัญญาประดิษฐ์แบบแขนง (Narrow Artificial Intelligence) ปัญญาประดิษฐ์แบบแขนงเฉพาะ (Specialized Artificial Intelligence) และ ปัญญาประดิษฐ์ที่เรียนรู้ (Machine Learning) ตามลำดับ

AI Deepfakes เป็นระบบกระตุ้นเพื่อให้ AI เรียนรู้ว่าการกระทำใดมีผลลัพธ์ดีหรือเสีย และปรับปรุงสามารถทำประโยชน์ได้มาก ซึ่งแรกเริ่ม AI Deepfake ถูกนำมาใช้ประโยชน์ในการตรวจจับใบหน้าคนของ AI เข้ามา และทำการสลับและแทนที่ใบหน้าของคนที่เราเลือก ในตอนแรกเทคโนโลยีนี้ดูเหมือนจะมีข้อจำกัดมากมาย เช่น วิดีโอที่ได้ออกมาจะมีคุณภาพต่ำและมีท่าทางของการแสดงออกที่ไม่เป็นธรรมชาติ แต่ด้วยเวลาที่ผ่านไปอย่างรวดเร็วทำให้เทคโนโลยี Deepfakes ได้รับการพัฒนาอย่างต่อเนื่องทำให้ข้อจำกัดต่าง ๆ ทำให้ทุกวันนี้ วิดีโอที่ออกมาโดยใช้เทคโนโลยี Deepfakes แทบจะแยกด้วยตาเปล่าไม่ได้เลย เช่น วิดีโอที่แสดงใบหน้าและท่าทางของ Mark Zuckerberg เจ้าของและผู้บริหารของ Facebook ที่ถูกทำขึ้นโดยเทคโนโลยี Deepfake ทั้งหมด (Gamingdose, 2021)

จุดเริ่มต้นความน่ากลัวของ Deepfake เกิดขึ้นในปี 2017 เพราะมีผู้ใช้อินเทอร์เน็ต ตั้งชื่อว่า Deepfake ใช้ภาพของดารานักแสดงดังระดับโลกมากมาย มาสลับหน้ากับภาพนักแสดงหนึ่งปี ทำให้ดาราคอนนั้นเกิดความเสียหาย เกิดเป็นปัญหาใหญ่ที่มีผู้คนโจมตีและกล่าวถึงความน่ากลัวของ AI หนาหูขึ้นเรื่อย ๆ ซึ่งถึงแม้ว่าในช่วงแรก ๆ นั้น Deepfake เองจะมีจุดอ่อนอยู่ตรงที่ภาพที่ออกมานั้นมีความละเอียดต่ำ ไม่ชัดรวมถึงยังสามารถมองเห็นและรู้ได้ว่าเป็นของเลียนแบบและตัดต่อ ไม่ใช่ของจริง แต่เมื่อกาลเวลาผ่านไปเทคโนโลยีต่าง ๆ ถูกพัฒนาขึ้นกว่าเดิม จึงทำให้ AI และ Deepfake มีประสิทธิภาพมากขึ้น จนเข้าใกล้คำว่า “อันตราย” แบบที่ทุกคนเคยหวาดกลัว ดังปรากฏผลเสียที่เป็นข่าวใหญ่ทั่วโลก ได้แก่ มิจนาซีฟใช้ Deepfakes ปลอมเป็นหัวหน้าฝ่ายการเงิน หลอกเงิน 900 ล้านบาท จากฮ่องกง เมื่อเจ้าหน้าที่การเงินของบริษัทข้ามชาติแห่งหนึ่งถูกหลอกให้โอนเงิน 25 ล้านดอลลาร์สหรัฐ ให้กับมิจนาซีฟที่ใช้เทคโนโลยี “Deepfakes” สวมรอยเป็นประธานเจ้าหน้าที่ฝ่ายการเงิน หรือซีเอฟโอของบริษัท ซึ่งปรากฏตัวในการประชุมทางวิดีโอ (PPTV online, 2024) สอดคล้องกับเมื่อปลายเดือนมกราคม 2024 ภาพลามกอนาจารของนักร้องชื่อดัง เทย์เลอร์ สวิฟต์ ที่สร้างโดยเทคโนโลยี Deepfakes ได้แพร่กระจายไปทั่วโซเชียลมีเดีย และสอดคล้องกับประเด็นข่าวเรื่องเกือบเนียนแล้ว มิจนาซีฟโชว์ล้ำวิดีโอคอลเป็นตำรวจ แต่นั่งนิ่ง ขยับได้แค่ปาก ซึ่งเป็นความก้าวหน้าของมิจนาซีฟ คอลเซ็นเตอร์ ล่าสุดเป็นวิดีโอแชต ขวนเปิดกล้อง เปิดมาเป็นภาพตำรวจเลยจ้า แต่ตำรวจนั่งนิ่งเลยนะ ขยับแต่ปากพะงาบๆ มันคือ deepfake หรือ deepfake dubs ที่ทำให้ปากขยับตามเสียงพูดได้ (Thairath online, 2023) พบว่า ข้อมูลข่าวสารดังกล่าวเป็นการตอกย้ำถึงศักยภาพในการทำลายล้างที่เกิดจากเทคโนโลยี Deepfake อย่างแท้จริง สอดคล้องกับผลกระทบของ AI deepfake ครอบคลุมไปมากกว่ากรณีส่วนบุคคลและนำไปสู่วงกว้าง คือ ผลกระทบระดับสังคมที่กว้างขึ้น รวมถึงภัยคุกคามต่อความมั่นคงแห่งชาติและการเสื่อมทรุดโทรม ของความไว้วางใจในสื่อและแหล่งข้อมูล (Rattamangam, Thitipasitthikorn, Yusamran, Wichai, 2022)

AI Deepfakes จึงเป็นเทคโนโลยีที่มีความเสี่ยงสูงที่ให้คุณและโทษกับผู้นำไปใช้ จึงจำเป็นต้องมีแนวทางที่ชัดเจนเพื่อรับมือกับเทคโนโลยีนี้อย่างมีประสิทธิภาพ อย่างรู้เท่าทันซึ่งเป็นท่าทีภายในจิตใจที่มีต่อ AI Deepfakes การรู้เท่าทันได้ต้องอาศัยโยนิโสมนสิการอย่างมีวิจารณญาณเพื่อใช้เป็นเครื่องมือตรวจจับ AI Deepfakes การสร้างความตระหนักรู้ การสนับสนุนการออกกฎหมาย และการพัฒนาเทคโนโลยีต่อต้าน Deepfakes ทั้งหมดนี้ล้วนเป็นแนวทางสำคัญในการลดความเสี่ยงจากเทคโนโลยีนี้อย่างยิ่ง

กรอบในการวิเคราะห์

1. AI Deepfake อยู่ในลำดับพัฒนาการปัญญาประดิษฐ์ระดับใด มีลักษณะการทำงานอย่างไรและผลกระทบหรือความเสี่ยงมีอะไรบ้าง
2. โยนีโสมนสิการเป็นเครื่องมือรู้เท่าทัน AI Deepfake ได้อย่างไร

ทำความเข้าใจกับ AI Deepfake

ความจริง AI Deepfakes เป็นส่วนหนึ่งของ AI Deep learning ซึ่งเจตนาโดยธรรมชาติของ Deep learning อาศัยเครื่องมืออัลกอริทึม (Algorithm) เพื่อเรียนรู้ และจดจำรูปแบบซับซ้อนจากข้อมูลขนาดใหญ่ ประกอบด้วยเครือข่ายประสาทเทียมชั้นลึก (Deep Neural Networks) ซึ่งเป็นโครงสร้างคณิตศาสตร์ที่เชื่อมต่อกันเป็นลำดับชั้น ที่มีความสลับซับซ้อน คล้ายกับจำลองการทำงานของเซลล์ประสาทมนุษย์ที่รับข้อมูลแล้วประมวลผล และส่งผลลัพธ์ในขั้นถัดไป จนสามารถนำข้อมูลที่ได้จากการเรียนรู้ไปทำนาย หรือจำแนกข้อมูลได้ด้วยตนเอง เหมาะสำหรับการประมวลผลข้อมูลที่มีความซับซ้อนสูง เช่น ข้อมูลภาพ เสียง ข้อความ ฯลฯ ทำให้สามารถนำไปประยุกต์ใช้ในงานด้านต่างๆ เช่น การรู้จำภาพและวิดีโอ การแปลภาษา การสังเคราะห์เสียงพูด เป็นต้น ดังนั้น AI Deepfakes ไม่สามารถตีค่าหรือแปลเจตนาว่า เป็นกุศโลบายหรืออกุศโลบาย เป็นความดีหรือความไม่ดี เป็นเจตนาดีหรือเจตนาร้าย แต่ AI Deepfakes มีหน้าที่เรียนรู้จากข้อมูลขนาดใหญ่แล้วแสดงผลหรือออกมาเพื่อให้สมจริง จากนั้นก็เป็นอันหมดกรอบขอบเขตหน้าที่ของ AI Deepfakes ว่าจะนำผลลัพธ์ที่ได้จาก AI Deepfakes ไปทำอะไรต่อที่เป็นประโยชน์หรือโทษอย่างไร

อนึ่ง จุดเปลี่ยนสำคัญที่ทำให้ Deepfake มีความแนบเนียนมากขึ้น เกิดจากการผสมผสานของ 2 เทคโนโลยีหลัก คือ เครือข่ายประสาทเทียม (Artificial Neural Networks) และ Generative Adversarial Networks (GANs) โดยวิธีการหลักคือ การที่ AI ใช้เครือข่ายประสาทเทียม (Artificial Neural Networks) เก็บข้อมูลต้นฉบับที่ต้องการปลอมแปลงในปริมาณมาก แล้วมาประมวลผลเลียนแบบระบบประสาทของมนุษย์ จนสามารถเรียนรู้ได้เอง และผลิตสื่อสังเคราะห์นั้น ๆ และเพื่อความแนบเนียนและเป็นธรรมชาติของเนื้อหา จึงมีอีกหนึ่งตัวช่วยสำคัญคือ Generative Adversarial Networks (GANs) ซึ่งเป็นอัลกอริทึมที่จะคอยตรวจสอบและบอกว่า สื่อสังเคราะห์ดังกล่าวนั้นใกล้เคียงกับเป้าหมายหรือยัง ซึ่งหากยังก็จะส่งพีตแบ็กกลับไปยังผู้สร้างเพื่อให้ปรับแก้จนกระทั่งใกล้เคียงกับต้นฉบับจริงมากที่สุด (True Digital, 2022) ดังนั้นสามารถวิเคราะห์ได้ว่า AI Deepfakes ซึ่งมีรากฐานมาจากปัญญาประดิษฐ์ที่เรียกว่าการเรียนรู้เชิงลึก การเรียนรู้เชิงลึกใช้เครือข่ายประสาทเทียมที่ได้รับแรงบันดาลใจจากโครงสร้างและการทำงานของสมองมนุษย์ในการประมวลผลข้อมูลและจดจำรูปแบบในบริบทของ Deepfakes เครือข่ายประสาทเทียมเหล่านี้ได้รับการฝึกฝนเกี่ยวกับข้อมูลภาพและวิดีโอจำนวนมหาศาลของบุคคลใดบุคคลหนึ่ง ข้อมูลนี้สอนเครือข่ายถึงความซับซ้อนของรูปลักษณ์ของบุคคลนั้น รวมถึงลักษณะใบหน้า การแสดงออกทางสีหน้า พื้นผิว และแม้แต่สภาพแสงด้วยข้อมูลการฝึกอบรมที่เพียงพอ เครือข่ายประสาทเทียมจึงมีความเชี่ยวชาญในการสร้างภาพและวิดีโอที่สมจริงของบุคคลนั้น สามารถดัดแปลงพฤติกรรมที่อยู่หรือแม้แต่สร้างเนื้อหาใหม่ทั้งหมด ขณะเดียวกันก็ทำให้ดูเหมือนจริงอย่างน่าเชื่อ การพัฒนา AI Deepfakes สามารถย้อนกลับไปถึงความก้าวหน้าในการวิจัยการเรียนรู้เชิงลึกในช่วงทศวรรษที่ผ่านมา ต่อไปนี้เป็นลำดับเวลาสั้น ๆ ของเหตุการณ์สำคัญ

พัฒนาการเทคโนโลยี AI Deepfakes เกิดจากเทคนิคการเรียนรู้เชิงลึก (Deep Learning) ซึ่งเป็นสาขาหนึ่งของการเรียนรู้ของเครื่อง (Machine Learning) เทคนิคนี้ถูกพัฒนาขึ้นมาในปี 2014 โดยนักวิจัยชาวอเมริกันในสาขาวิศวกรรมคอมพิวเตอร์ แต่ได้รับความนิยมอย่างกว้างขวางในปี 2017 เมื่อมีการเผยแพร่วิดีโอปลอมของนักแสดงและบุคคลสำคัญ ทำให้เกิดความกังวลเรื่องการนำไปใช้ในทางที่ผิด เช่น การหลอกลวง การละเมิดสิทธิส่วนบุคคล และการสร้างข่าวปลอม โดยมีจุดเริ่มต้น Deepfake นั้นเกิดขึ้นในปีค.ศ.2017 เพราะมี

ผู้ใช้อินเทอร์เน็ต ตั้งชื่อว่า Deepfake ใช้ภาพของดารานักแสดงดังระดับโลกมากมาย มาสลับหน้ากับภาพนักแสดงหนังโป๊ ทำให้ดาราคานั้นเกิดความเสียหาย เกิดเป็นปัญหาใหญ่ที่มีผู้คนโจมตีและกล่าวถึงความน่ากลัวของ AI หนาหูขึ้นเรื่อย ๆ ซึ่งถึงแม้ว่าในช่วงแรก ๆ นั้น Deepfake เองจะมีจุดอ่อนอยู่ตรงที่ภาพที่ออกมานั้นมีความละเอียดต่ำ ไม่ชัด รวมถึงยังสามารถมองและรู้ได้ว่าเป็นของเลียนแบบและตัดต่อ ไม่ใช่ของจริง แต่เมื่อเวลาผ่านไป เทคโนโลยีต่าง ๆ ถูกพัฒนาขึ้นกว่าเดิม จึงทำให้ AI และ Deepfake นี้มีประสิทธิภาพมากขึ้น จนเข้าใกล้คำว่า “อันตราย” แบบที่ทุกคนเคยหวาดกลัว (Gamingdose, 2021) สอดคล้องกับสำนักข่าวไทยพีบีเอส (2023) รายงานว่า Deepfake เกิดขึ้นเมื่อปีค.ศ. 2017 คณะผู้วิจัยจากมหาวิทยาลัยวอชิงตัน ใช้ AI สร้างเค้าโครงปากของ “บารัค โอบามา” อดีตประธานาธิบดีสหรัฐอเมริกา และผสมเข้ากับภาพวิดีโอของโอบามาที่เผยแพร่ผ่านสื่อต่าง ๆ ก่อนนำเสียงของบุคคลอื่นใส่ผสมในใบหน้าได้อย่างแนบเนียน

ปัจจุบัน เทคโนโลยี AI deepfakes มีการพัฒนาอย่างต่อเนื่อง โดยนักวิจัยได้ผลักดันขอบเขตของความสมจริงและการเข้าถึงได้ มีแม้กระทั่งแอปที่เป็นมิตรต่อผู้ใช้ที่อนุญาตให้ทุกคนสร้างดีเฟคของตัวเองโดยมีความรู้ด้านเทคนิคเพียงเล็กน้อย ซึ่ง Deepfakes และ AI สามารถปรับเปลี่ยนการโฆษณาโดยการเปลี่ยนแปลงอย่างฉับพลันทั้งเนื้อหาเสียงและภาพโดยอัตโนมัติ (Colin, Campbell., Kirk, Plangger., Sean, Sands., Jan, Kietzmann., Kenneth, Bates, 2022)

กรณีตัวอย่างที่น่าสนใจและส่งผลกระทบต่อชีวิตผู้คนในสังคมอย่างหลีกเลี่ยงไม่ได้ เช่น กรณีของผู้ใช้ Reddit ที่มีชื่อว่า Deepfake โพสต์คลิปโป๊โดยปลอมใบหน้าของคนดังอย่าง กัล กาด็อต, เทย์เลอร์ สวิฟต์, สการ์เลตต์ โจแฮนส์สัน และคนอื่น ๆ จนนำมาสู่การใช้เพื่อก่อวินาศกรรม ช่มชู้ สร้างความเสียหายแก่บุคคลหรือองค์กร รวมไปถึงการทำลายความมั่นคงของชาติด้วยโดยคนดังที่ถูกนำมาทำ Deepfake เช่น “บารัค โอบามา” กล่าวถึง “โดนัลด์ ทรัมป์” ด้วยถ้อยคำหยาบคาย หรือ “มาร์ค ซักเคอร์เบิร์ก” ผู้ก่อตั้งเฟซบุ๊ก กล่าวว่า นโยบายของเฟซบุ๊กคือหาประโยชน์จากผู้ใช้งาน หรืออย่าง “ทอม ครูซ” ที่ถูกคนนำไปเลียนแบบใบหน้าเพื่อถ่ายวิดีโอลง TikTok จนทำให้เกิดความสับสน นอกจากนี้ยังมี “ประธานาธิบดีโวลโด้มิร์ เซเลนสกี” ผู้นำประเทศยูเครน ถูกนำไปสร้างโฆษณาชวนเชื่อหลอกให้ชาวยูเครนยอมจำนนในช่วงเวลาของความขัดแย้งระหว่างยูเครนกับรัสเซีย เป็นต้น

ลักษณะ AI Deepfakes มีอัตลักษณ์โดดเด่นด้านความสมจริง มีความหลากหลาย และ การเข้าถึงเทคโนโลยีดิจิทัล ซึ่งความสมจริงด้วยเทคโนโลยี Deepfake พัฒนาขึ้นอย่างรวดเร็ว ทำให้ภาพและเสียงที่สร้างขึ้นมีความสมจริงมากขึ้น ยากต่อการแยกแยะด้วยตาเปล่า ความหลากหลายด้วย Deepfake สามารถสร้างสื่อสังเคราะห์ได้หลากหลายรูปแบบ เช่น เปลี่ยนใบหน้าในวิดีโอ ใส่เสียงพูดใหม่ หรือสร้างภาพคนเสมือนจริง และการเข้าถึงด้วยเทคโนโลยี Deepfake มีการใช้งานแพร่หลาย เข้าถึงได้ง่ายผ่านโปรแกรมและแอปพลิเคชันต่างๆ

โดยสรุป AI Deepfakes เป็นส่วนหนึ่งของระดับการพัฒนาปัญญาประดิษฐ์ขั้น “ปัญญาประดิษฐ์ทั่วไป (AGI: Artificial General Intelligence)” ซึ่งเป็นรูปแบบหนึ่งที่สามารถสร้างสื่อสังเคราะห์ ทั้งภาพนิ่ง เสียง ภาพเคลื่อนไหวพร้อมลอกเลียนเสียงจากมนุษย์จริงมาพูดหรือคุยจริง ๆ จนแทบแยกไม่ออกซึ่งปัญญาประดิษฐ์สามารถมีระบบการเรียนรู้ หรือ Deep Learning จากลักษณะภายนอกของบุคคล เช่น สีผิว ตา ปาก จมูก รูปลักษณะต่าง ๆ ด้วยความฉลาดล้ำจากการเรียนรู้ข้อมูลมหาศาลของ Generative AI

AI Deepfake ทำอะไรได้บ้าง

แพลตฟอร์มดิจิทัลที่มีอยู่บนโลกออนไลน์ล้วนมีเหตุปัจจัยของปัญญาประดิษฐ์เข้ามามีส่วนร่วมอย่างหลีกเลี่ยงไม่ได้ โดยเฉพาะแพลตฟอร์มดิจิทัลที่ต้องอาศัยข้อมูลขนาดใหญ่แล้วประมวลผลเพื่อนำไปใช้งานตามต้องการ ในที่นี้ AI Deepfake จะเข้ามาทำงานโดยการเรียนรู้ข้อมูลจำนวนมหาศาลไม่ว่าจะมาจากแหล่งข้อมูล

ที่มาจากรูปภาพ วิดีโอ เสียง บทสัมภาษณ์ ข้อความ สถิติ แล้วนำมาสร้างสื่อสังเคราะห์ใหม่โดยใส่ใบหน้า เสียง ท่าทาง ของบุคคลไปยังวิดีโอหรือภาพอื่นเพื่อให้สมจริง ซึ่งการสร้างคลิปภาพและเสียงจาก AI Deepfake จะมี 2 แบบ คือ การปลอมแค่บางส่วน หรือ Face Wrap เช่น เอาหน้าที่อยากปลอม ไปปะใส่หน้าคนจริงที่ถ่าย แล้วพยายามเลียนแบบเสียง และพฤติกรรมให้เหมือนคนนั้น ๆ (Office of Information Technology Administration for Educational Development, 2024) โดยมีขั้นตอนดังนี้

1. การเรียนรู้และสร้างโมเดล

ระบบ AI จะถูกฝึกฝนด้วยข้อมูลจำนวนมาก เช่น ภาพ วิดีโอ ของใบหน้า เสียงพูด ท่าทาง ของบุคคล เป้าหมาย จากนั้นระบบจะสร้างโมเดล (Model) เพื่อจดจำและสกัดลักษณะเฉพาะของบุคคลนั้น ๆ

2. การสร้างภาพ/วิดีโอปลอม

โมเดลที่ได้จะนำมาใช้ในการสร้างภาพหรือวิดีโอปลอม โดยการแปลงใบหน้า เสียง หรือท่าทางของบุคคลอื่นให้เหมือนกับบุคคลเป้าหมาย AI Deepfakes สามารถสลับใบหน้าในภาพและวิดีโอ สร้างเนื้อหาที่สมจริงสูง ซึ่งก่อให้เกิดภัยคุกคาม เช่น ข่วปลอม หลอกหลวง และการฉ้อโกง ส่งผลกระทบต่อความเป็นส่วนตัว ประชาธิปไตย และความมั่นคงแห่งชาติ (Hina, Shahzad., etc, 2022) เทคนิคหนึ่งที่ใช้คือ Generative Adversarial Network (GAN) ซึ่งใช้สองเครือข่าย AI ต่อสู้กันเพื่อปรับปรุงคุณภาพให้ดูเหมือนจริงมากที่สุด AI Deepfake สามารถเปลี่ยนวิดีโอและรูปภาพเพื่อสร้างเนื้อหาที่ทำให้เข้าใจผิดการศึกษาที่ใช้ CNN-LSTM-FCNS แบบแคสเคดเพื่อตรวจจับวิดีโอที่เปลี่ยนแปลง AI ตามลำดับสถานะตา ทำให้ได้ความแม่นยำ 90.8% (Muhammad, Salihin, Saealal., etc, 2022)

3. การตัดต่อ/ประกอบภาพ

ภาพหรือวิดีโอปลอมที่สร้างขึ้นจะถูกนำมาตัดต่อ ประกอบเข้ากับสถานการณ์หรือเนื้อหาที่ต้องการ อาจมีการใช้เทคนิคอื่นๆ เช่น เสียงสังเคราะห์ ประกอบเข้าไปด้วย เพื่อให้ดูสมจริงมากขึ้น ไม่ว่าจะเป็นการสลับหน้า สลับใบหน้าระหว่างแหล่งวิดีโอสองแหล่งได้อย่างราบรื่น หรือจะเป็นการเจดหุ่นที่มีการถ่ายโอนการแสดงออกทางสีหน้าจากบุคคลหนึ่งไปยังอีกบุคคลหนึ่ง หรือจะเป็นการลิปซิงค์ซึ่งเป็นจัดการการเคลื่อนไหวของปากเพื่อซิงค์กับทางเลือกอื่น แทร็กเสียง และการโคลนเสียงเป็นเลียนแบบเสียงของบุคคลที่มีอยู่แล้ว (Damir, Yalov., Danil Myakin, 2023)

4. การแพร่กระจาย

เมื่อสร้างเสร็จ ภาพหรือวิดีโอ Deepfake จะถูกเผยแพร่ออกไปบนอินเทอร์เน็ตหรือสื่อต่าง ๆ เนื่องจากมีความสมจริงมาก จึงทำให้ผู้รับสารหลงเชื่อได้ง่าย

ฉากทัศน์ในภาพกว้างจะเห็นได้ชัดเจนว่า เทคโนโลยีขั้นสูงนี้ดึงความแข็งแกร่งจากกลไกการเปิดใช้งาน เช่น generative adversarial network (GANs) โปรแกรมเข้ารหัสอัตโนมัติ และการถ่ายโอนสไตล์ ส่วนประกอบสำคัญทั้งหมดที่สนับสนุนโมเดลการเรียนรู้เชิงลึกที่เติมพลังให้กับ deepfakes เครื่องมือ Deepfake ที่เป็นมิตรกับผู้ใช้ได้แทรกซึมเข้าสู่จิตสำนึกกระแสหลัก เครื่องมือเหล่านี้ก่อให้เกิดแอปพลิเคชันยอดนิยมมากมาย รวมถึงวิดีโอเลียนแบบคนดัง ฟิวเตอร์เปลี่ยนใบหน้า วิดีโอพากย์ และการสร้างมีมแบบไวรัล โดยสุดยอดเครื่องกำเนิด AI Deepfake ที่ล้ำสมัย 10 อันดับในปี 2023 ได้แก่ ซาโอ รีเฟช ดิฟเฟซแล็บ สร้างอวตาร คิดถึงอย่างลึกซึ้ง วอมโบ ซินริเซีย ดิฟเฟซไลฟ์ ผู้สร้าง MetaHuman แอปปลอม (Damir, Yalov., Danil Myakin, 2023) ซึ่งแน่นอนว่า Deepfake 10 อันดับแรกในต้นปี 2024 ที่สามารถ Deepfake Video Maker ได้แก่ แพลตฟอร์ม AI faceswap, Deepswap, Deepfakes Web, DeepFaceSwap.AI, FaceSwaponline, FaceMagic, Faceswapper.ai, DeepBrain.ai, Pixble, และ SwapStream.ai ตามลำดับ ซึ่งเครื่องมือสร้างวิดีโอ Deepfakes เหล่านี้ ได้เปลี่ยนแปลงการผลิตเนื้อหาดิจิทัลโดยให้ผู้ใช้

สามารถเข้าถึงเครื่องมือที่สร้างสรรค์สำหรับการสร้างวิดีโอที่สมจริงและน่าดึงดูด ตอบโจทย์ผู้บริโภค สามารถสื่อภาพที่กำลังถูกกำหนดโดยจุดเชื่อมต่อระหว่างความคิดสร้างสรรค์และเทคโนโลยี (Eddy, 2024)

ปัจจุบัน Deepfakes ได้รับการพัฒนาให้มีความสมจริงและซับซ้อนมากขึ้นตามลำดับความเจริญก้าวหน้าทางปัญญาประดิษฐ์ ทั้งภาพ เสียง และการเคลื่อนไหว จึงเป็นความท้าทายสำหรับการตรวจจับและป้องกัน ส่วนประโยชน์สามารถสร้างสรรค์ในอุตสาหกรรมบันเทิง Deepfakes โดยสามารถนำมาสร้างเทคนิคพิเศษ (Special Effects) และตัดต่อใบหน้าได้อย่างแนบเนียน ทำให้การทำงานตัดต่อภาพยนตร์และวิดีโอเกมง่ายขึ้น AI Deepfakes ก็มีประโยชน์หากใช้ในทางที่ดี เช่น ในด้านการศึกษาอาจนำมา สร้างครูในโลกเสมือนจริง หรือ Virtual Teacher อย่างการจำลอง “อัลเบิร์ต ไอน์สไตน์” ขึ้นมาเพื่อให้เป็น ครูสอนวิชาวิทยาศาสตร์ เป็นการเพิ่มอรรถรสในการเรียนได้เป็นอย่างดี เด็กนักเรียนอาจมีความรู้สึกสนุก กับการเรียนมากยิ่งขึ้น หรือแม้แต่การใช้เทคโนโลยีนี้สร้าง ศิลปิน นักร้องที่จากโลกนี้ไปแล้ว ให้กลับมาแสดงภาพยนตร์ หรือยืนร้องเพลง เสมือนตัวจริง ได้ ซึ่งก็มีการทำเช่นนี้แล้วทั้งในประเทศไทยและต่างประเทศ เป็นการนำเทคโนโลยี มาสร้างความสุขให้ แฟนๆอีกครั้ง เสมือนศิลปินผู้นั้นยังมีชีวิตและลมหายใจอยู่ (Uninet, 2023) อย่างไรก็ตาม การใช้ Deepfakes ยังมีความเสี่ยง และควรใช้ด้วยความระมัดระวัง เพื่อป้องกันความเสี่ยงที่อาจเกิดขึ้น

AI Deepfake สามารถสร้างผลกระทบและภัยคุกคามอะไรได้บ้าง

จากการศึกษาของ iProov 71% ของผู้ตอบแบบสอบถามทั่วโลก ไม่ทราบว่ Deepfakes คืออะไร และ 43% ยอมรับว่าไม่สามารถตรวจจับได้ เป็นที่มาให้อัตราของผู้ถูกหลอกลวงด้วย Deepfakes ค่อย ๆ เพิ่มขึ้นและอยู่ใกล้ตัวมากยิ่งขึ้น (Techsauce Team, 2024) DeepTrace เดือนกันยายน 2019 พบวิดีโอ Deepfake 15,000 รายการ เพิ่มขึ้นเกือบ 2 เท่าในช่วง 9 เดือนที่ผ่านมา โดยพบว่าเป็นวิดีโออาจารย์ถึง 96% ซึ่งการทำ Deepfake นั้น แม้เป็นคนที่ไม่เชี่ยวชาญก็สามารถสร้างวิดีโอปลอมด้วยภาพถ่ายเพียงไม่กี่ภาพได้ จึงมีแนวโน้มแพร่กระจายไปได้ไกลทั่วโลก (Thai PBS Sci & Tech, 2023) ก่อนหน้านี้นปี 2017 คณะผู้วิจัยจากมหาวิทยาลัยวอชิงตัน ใช้ AI สร้างเค้าโครงปากของ “บารัค โอบามา” อดีตประธานาธิบดีสหรัฐอเมริกา และผสมเข้ากับภาพวิดีโอของโอบามาที่เผยแพร่ผ่านสื่อต่าง ๆ ก่อนนำเสียงของบุคคลอื่นใส่ลงในใบหน้าได้อย่างแนบเนียน ผลการศึกษากรณีศึกษาแสดงให้เห็นว่าเป็นไปได้ที่จะโคลนเสียงของใครบางคนด้วยแค่ป๊อปมาตรฐานโดยไม่จำเป็นต้องใช้ทรัพยากรการประมวลผลประสิทธิภาพสูงและใช้เสียงอ้างอิงเพียงไม่กี่วินาทีซึ่งสร้างประสบการณ์ผู้ใช้ที่เหนือกว่า แต่ในขณะเดียวกันก็เผยให้เห็นว่าใครก็สามารถเข้าถึงการโคลนเสียงได้อย่างง่ายดายเพียงใด (Narora, Amezaga., Jeremy, Hajek, 2022) ซึ่งสถิติการเติบโตของวิดีโอติปเฟกจากช่วงสิ้นปี 2018 ที่ผ่านมามีอัตราการเติบโตขึ้นเกือบ 100% โดยมีจำนวนวิดีโอติปเฟกสูงถึง 14,678 วิดีโอ (96% เกี่ยวข้องกับหนังอนาจาร) จากรายงานข่าวดังกล่าวข้างต้น สามารถแจกแจงประเด็น AI Deepfakes สร้างผลกระทบและภัยคุกคาม ดังนี้

1. ด้านข่าวสารและข้อมูล สามารถเผยแพร่ข่าวปลอม สามารถสร้างวิดีโอหรือข้อความปลอมแปลงที่เหมือนจริง ซึ่งสามารถบิดเบือนความจริง ใส่ร้ายผู้อื่น หรือสร้างความสับสนวุ่นวายในสังคม มีการทำลายความน่าเชื่อถือของสื่อ เมื่อผู้คนเริ่มไม่แน่ใจว่าวิดีโอหรือข้อความใดเป็นจริง จะส่งผลต่อความน่าเชื่อถือของสื่อโดยรวม มีการบิดเบือนความคิดเห็นของประชาชน Deepfakes สามารถใช้เพื่อโฆษณาชวนเชื่อ หรือบิดเบือนความคิดเห็นของประชาชน ในประเด็นทางการเมืองหรือสังคม

2. ด้านบุคคล สามารถการทำลายชื่อเสียง สามารถใช้สร้างวิดีโอหรือข้อความปลอมแปลง เพื่อทำลายชื่อเสียง หรือสร้างความอับอายให้กับบุคคลอื่น มีการคุกคามทางไซเบอร์ สามารถใช้เพื่อคุกคาม หรือกลั่นแกล้ง

ผู้ผ่านทางออนไลน์ มีการลวงละเมิดความเป็นส่วนตัว Deepfakes สามารถใช้เพื่อขโมยข้อมูลส่วนตัว หรือปลอมแปลงตัวตน

3. ด้านความมั่นคง สามารถปลุกปั่นให้เกิดความขัดแย้ง สามารถใช้เพื่อปลุกปั่นให้เกิดความเกลียดชัง หรือความขัดแย้งในสังคม มีการโจมตีทางไซเบอร์ สามารถใช้เพื่อหลอกลวง หรือโจมตีระบบคอมพิวเตอร์ มีการบ่อนทำลายความมั่นคงของชาติ สามารถใช้เพื่อบ่อนทำลาย หรือสร้างความเสียหายให้กับชาติ

เพราะฉะนั้น ผลกระทบและความเสี่ยงด้านการทำลายหรือผลเสียอันเกิดจาก AI Deepfakes สามารถสร้างการละเมิดความน่าเชื่อถือ สามารถสร้างความสับสนและก่อความขัดแย้งตั้งแต่ระดับครอบครัวสู่สังคม และเกิดความเสี่ยงต่อการเลือกตั้งที่นำไปสู่การระดมพลปลุกปั่นและสร้างกระแสแรงกดดันในการเลือกตั้งส่งผลให้คนไม่สามารถแยกแยะข้อมูลที่ถูกรสร้างขึ้นกับข้อมูลจริงได้ จึงจำเป็นต้องระมัดระวังและไม่ประมาทในการใช้ AI Deepfakes ก่อนตัดสินใจนำไปใช้ซึ่งต้องคำนึงผลลัพธ์ที่ตามมา

โยนิโสมนสิการเป็นเครื่องมือรู้เท่าทัน AI Deepfake ได้อย่างไร

โยนิโสมนสิการเป็นธรรมที่มีอุปการะมาก เป็นยอดธรรมฝ่ายกุศล เป็นปัจจัยภายในที่ทำให้เกิดสัมมาทิฐิ เป็นต้น ดังปรากฏหลายแห่งในพระไตรปิฎก ดังมีรายละเอียด เกิดเรื่องที่ กรุงสาวัตถี พุทธพจน์ว่า

“ภิกษุทั้งหลาย เมื่อดวงอาทิตย์กำลังจะอุทัย ย่อมมีแสงอรุณขึ้นมาก่อน เป็นบุพนิมิต ฉันทโยนิโสมนสิการสัมปทา (ความถึงพร้อมด้วยการมนสิการโดยแยบคาย) ก็เป็นตัวนำ เป็นบุพนิมิตเพื่อความเกิดขึ้นแห่งอริยมรรคมีองค์ 8 ฉะนั้น” (Thai Tripitakas: 19/55/44-45)

พุทธพจน์ดังกล่าว สอดคล้องกับการรู้เท่าทัน AI Deepfakes ได้ย่อมแสดงถึงคำตอบของปัญหาหรือนิमितแห่งการนำ AI Deepfakes ตามอรรถธรรมสัมพันธ์ อีกทั้งโยนิโสมนสิการของผู้ใช้ AI Deepfakes ยังสามารถกำจัดนิรณ 5 ได้อีกด้วย สมดังพุทธพจน์ที่ว่า “เมื่อใช้โยนิโสมนสิการ นิรณ 5 ย่อมไม่เกิด ที่เกิดแล้วก็ถูกกำจัดได้ ในขณะที่เดียวกันก็เป็นเหตุให้โพฆณค 7 เกิดขึ้นและเจริญเต็มบริบูรณ์ ” (Thai Tripitakas: 19/446-7/122) นั่นหมายถึง ท่าทีภายในจิตใจของผู้ใช้ AI Deepfakes จะไม่เกิดอุปสรรคขัดขวางต่อการความดีหรือคุณประโยชน์โดยชอบธรรมเพราะมีเครื่องมืออาวุธสำคัญที่กำจัดนิรณ 5 ได้ นั่นคือ โยนิโสมนสิการ

คัมภีร์พระไตรปิฎกซึ่งเป็นหลักฐานชั้นปฐมภูมิของพระพุทธศาสนาเถรวาท ได้ใช้หลักโยนิโสมนสิการในลักษณะเชิงวิภาษวาทหรือขัดไปในเรื่องหาที่สามารถนำไปใช้เรื่องใดเรื่องหนึ่งได้ทันที ไม่ได้แบ่งเป็นประเภทรูปแบบ 10 ลักษณะดังที่ปราชญ์ชั้นหลังที่ได้นำมาวิเคราะห์ แต่ทั้งนี้รูปแบบทั้ง 10 ประการนั้นไม่ได้ทำให้หลักโยนิโสมนสิการเสียหลักการจากฐานปฐมภูมิ เป็นการดีที่รุ่นหลังได้ทำให้เห็นเด่นชัดโดยตกลึกจากโยนิโสมนสิการมาใช้ให้ถูกบริบทตามธรรมเนียมคลองธรรม ซึ่งการรู้เท่าทัน AI Deepfakes ที่จะเข้ามามีอิทธิพลต่อการบริโภคเทคโนโลยีให้ปลอดภัยและได้ประโยชน์อย่างแท้จริงตามวัตถุประสงค์แห่งความต้องการ

ขอบเขตเนื้อหาการนำรูปแบบโยนิโสมนสิการภายใต้ “การรู้เท่าทัน AI Deepfakes” จะวางกรอบด้วยรูปแบบโยนิโสมนสิการไม่ครบทั้ง 10 รูปแบบ เพราะเทคโนโลยีปัญญาประดิษฐ์ผ่านกิจกรรม Deep Learning แล้วสามารถสร้างสื่อสังเคราะห์เสมือนจริง มีสภาพการกำหนดรู้ว่าเป็นเทคโนโลยีเพื่อการบริโภคและตอบโต้ความต้องการของผู้เสพสื่อสังเคราะห์ในโลกเสมือนจริง ดังนั้นรูปแบบแรกของโยนิโสมนสิการที่สามารถวางกรอบแนวคิดเพื่อเชื่อมโยงรูปแบบอื่น ๆ ของโยนิโสมนสิการ คือ วิธีคิดแบบอริยสัจหรือวิธีคิดแบบแก้ปัญหา โดยจะใช้กิจในอริยสัจเป็นกรอบแนวทางขั้นตอนแล้วเสริมด้วยรูปแบบของโยนิโสมนสิการอื่น ๆ บุรณการตามขั้นตอนแล้วสำเร็จเป็นฐานให้ปัญญาได้ตัดสินใจเด็ดขาดต่อการนำ AI Deepfakes ไปใช้ให้ถูกต้องเหมาะสมและประโยชน์โดยธรรมอย่างแท้จริง ลำดับจากนี้ นำเสนอขั้นตอนของรูปแบบโยนิโสมนสิการ ดังต่อไปนี้

1. กำหนดรู้ตามสภาพจริง

เริ่มต้นด้วยโยนิโสมนสิการรูปแบบวิธีคิดแบบสืบสาวเหตุปัจจัย โดยผู้บริโภคร AI Deepfakes จำเป็นต้องกำหนดรู้ตามความเป็นจริง โดยตั้งสติด้วยความไม่ประมาท จุดใจคิดว่า สื่อสังเคราะห์ที่เกิดจาก AI Deepfakes มีเหตุปัจจัยใดที่สามารถสร้างผลงานสื่อสังเคราะห์ได้เสมือนจริง โดยวิธีสังเกตจากลักษณะทางกายภาพ และวิธีสังเกตจากลักษณะอื่น ๆ (Techsauce Team, 2024)

วิธีการสังเกตจากลักษณะทางกายภาพ

การกะพริบตา: การกะพริบตาที่มากเกินไป หรือไม่กะพริบตาเลย ถือเป็นจุดสังเกต Deepfake เพราะการเลียนแบบการเคลื่อนไหวของตาจริง ยังคงเป็นเรื่องที่ทำได้ยาก

ลักษณะปากและฟัน: สังเกตได้เวลาปากที่ขยับไม่ตรงเวลาพูด ช้ากว่าเสียง รูปปากเคลื่อนไหวไม่เป็นธรรมชาติ ไม่เห็นลักษณะของฟันที่ชัดเจน อนึ่งควรกำหนดรู้จังหวะการเว้นวรรค สิ่งที่ AI ยังผิดพลาดอยู่เป็น เรื่องของอารมณ์ จะไม่มีจังหวะหยุด จะพูดรัวยาวไปเลย หรือแม้กระทั่งโทนเสียงโมนোটोन เป็นเสียงเดียวราบเรียบ ไม่มีเสียงสูงเสียงต่ำ ไม่ได้เน้นความสำคัญของคำบางคำ ยกตัวอย่าง อุเบอร์ เวลาพูดจะเป็น “อุ-เบ้อ” แต่เมื่อให้ AI อ่านจะเป็น “อุเบอ” เนื่องจากสมองคนเรามี DATA อยู่จำนวนมาก จะทำให้แยกออกได้ว่า คำไหนที่ควรจะเป็นเสียงสูงหรือเสียงต่ำ เช่น อุเบอร์ กับ เบอร์โทรศัพท์ มีคำว่า “เบอร์” เหมือนกัน แต่ออกเสียงไม่เหมือนกัน ซึ่งหากผิดเพี้ยนไปก็ตั้งสมมติฐานได้ว่าอาจจะจะเป็น AI (Thai PBS Sci & Tech, 2023) ดังนั้น ผู้บริโภคใช้วิธีคิดแบบอยู่กับปัจจุบัน โดยดูเนื้อหาในบริบท หรือศึกษารายละเอียดคอมเมนต์ ตรวจสอบหลาย ๆ ช่องทาง หรือศึกษาจากข่าวว่าเทคโนโลยีไปถึงไหนแล้ว มีความธรรมชาติมากแค่ไหน เป็นต้น

การเคลื่อนไหวของใบหน้า: Deepfake มักประสบปัญหาการวางโครงสร้างใบหน้าผิดปกติดังเช่น ใบหน้าหันไปทางหนึ่งแต่จมูกไม่ได้ขยับตามไปด้วย นอกจากนี้อาจจะสังเกตจากใบหน้าที่ขาดอารมณ์ร่วม ไม่สอดคล้องกับเนื้อหาที่กำลังพูดอยู่

รายละเอียดอื่น ๆ: สังเกตรายละเอียดเล็ก ๆ น้อย ๆ เช่น การสะท้อนของแสงและเงาผิดที่ผิดทาง แสงสะท้อนเวลาใส่แว่นตา เส้นผมที่ต้านแรงโน้มถ่วงหรือชี้ฟูมากเกินไป ตาหนีบนูนหรือหนังตาห้อย รอยเหี่ยวย่นตามอายุ ว่าตรงกับความเป็นจริงหรือไม่

วิธีสังเกตจากลักษณะอื่น ๆ

ความขัดของวิดีโอ: สังเกตจากการเบลอเพียงบางจุด เช่น ระหว่างใบหน้าและลำคอ หรือ ระหว่างคอและช่วงลำตัว จะช่วยให้สังเกตถึงความไม่เป็นระนาบเดียวกันของวิดีโอได้

เสียงที่ผิดปกติ: ผู้สร้าง Deepfake ไม่ค่อยใส่ใจกับการใส่เสียงเท่ากับการทำวิดีโอให้แนบเนียน ดังนั้น จะสังเกตได้จากเสียงที่ไม่สอดคล้องกับการพูด เสียงเหมือนหุ่นยนต์ การออกเสียงบางคำที่ผิดปกติ

บริบทและแหล่งที่มา: พิจารณาแหล่งที่มาและบริบทของวิดีโอ ว่าสอดคล้องกับข้อมูลที่ทราบหรือไม่ และพิจารณาว่ามาจากแหล่งที่เชื่อถือได้หรือหน่วยงานที่ไม่รู้จัก

นอกจากวิธีการสังเกตดังกล่าวข้างต้น ควรตั้งสติซึ่งเป็นอาหารของโยนิโสมนสิการ ด้วยรูปแบบวิธีคิดแบบเห็นคุณโทษและทางออก ซึ่งในที่นี้กำหนดรู้ที่โทษของ AI Deepfakes ว่าการขโมยข้อมูลเพื่อเข้าถึงข้อมูลส่วนตัว เป็นสิ่งที่ผู้คนทั่วโลกกังวลเกี่ยวกับ Deep Fake มากที่สุด เนื่องจากในปัจจุบันผู้คนสามารถทำธุรกรรมต่าง ๆ ในรูปแบบออนไลน์ผ่านโทรศัพท์มือถือ หรือคอมพิวเตอร์ได้ ซึ่งทำให้มีการใช้เทคโนโลยี Deep Fake ในการสวมรอยบุคคลเหล่านั้นทำธุรกรรม จนสามารถเข้าถึงข้อมูลส่วนตัวเช่น เลขบัตรประชาชน หรือข้อมูลบัญชีธนาคารได้ การถูกชักจูงในสิ่งที่ไม่เป็นความจริง เป็นอีกหนึ่งสิ่งที่ผู้คนกังวลเกี่ยวกับ Deepfake เนื่องจากอาจมีผู้ไม่หวังดี นำใบหน้าของเรามีชื่อเสียง มาสร้างข่าวสารหลอกลวงผู้คนให้หลงเชื่อตาม อาจทำให้เจ้าของใบหน้าเสียชื่อเสียง และสร้างความเสียหายต่อผู้ที่หลงเชื่อ และการทำอนาจาร เป็นการนำ Deepfake ด้วยการนำ

ใบหน้าของคนที่มีชื่อเสียง มาใส่ลงบนสื่อลามก โดยส่วนใหญ่มักเกิดกับผู้หญิง จนทำให้เกิดความเข้าใจผิด และทำให้บุคคลเจ้าของใบหน้าเสื่อมเสียชื่อเสียง (Appman, 2023)

1. กำจัดเหตุภัยคุกคามจาก AI Deepfakes

ก่อนจะตัดสินใจเชื่อและลงมือปฏิบัติการใช้ AI Deepfakes เพื่อประโยชน์แก่ผู้บริโภคร จำเป็นต้องตรวจสอบแหล่งที่มาและความน่าเชื่อถือให้ถูกต้อง มีหน่วยงานภาครัฐและเอกชนได้ตระหนักถึงการนำไปใช้ในทางที่ผิดและส่งผลเสียต่อผู้บริโภคและสังคมในระดับกว้าง จึงได้มีช่องทางตรวจสอบ Deepfakes เช่น MIT (Massachusetts Institute of Technology) หรือ สถาบันเทคโนโลยีแมสซาชูเซตส์ ได้สร้างเว็บไซต์ “Detect Fakes” ที่แสดงวิดีโอ Deepfake และของจริงคุณภาพสูงหลายพันรายการจากชุดข้อมูล DFDC สู่อสาธารณะ เพื่อให้ผู้เข้าชมได้ลองแยกแยะว่าเนื้อหาแบบไหนเป็นของจริง หรือปลอม ลองไปฝึกเล่นกันได้ที่นี่ > <https://detectfakes.media.mit.edu/> หรือ เว็บไซต์สำหรับการตรวจสอบวิดีโอ Deepfake เบื้องต้น > <https://scanner.deepware.ai/> หรือโมเดล OpenVINO™ AI, Intel® Deep Learning Boost และ Intel® AVX-512 ซึ่งเร่งความเร็วปัญญาประดิษฐ์ และขีดความสามารถในการประมวลผลสื่อในระบบทั้งยังเป็นแพลตฟอร์มตรวจจับ Deepfake แบบเรียลไทม์ หรือ Microsoft, Adobe, DARPA ตรวจสอบความถูกต้องและตรวจจับภาพหรือวิดีโอที่เกิดจากเทคโนโลยี Deepfake ในขณะที่โซเชียลมีเดียแพลตฟอร์มหลายแห่ง เช่น Twitter, TikTok และ Reddit ก็มีนโยบายห้ามเผยแพร่เนื้อหาที่สร้างจากสื่อสังเคราะห์ลักษณะ Deepfake หรือแม้แต่ Meta (Facebook เดิม) ที่มีโปรเจกต์ Deepfake Detection Challenge เมื่อปี 2020 ที่ใช้ AI วิเคราะห์และแยกแยะวิดีโอมากกว่า 100,000 คลิป ว่าอันไหนเป็นสื่อจริงและอันไหนเป็นสื่อสังเคราะห์ ด้าน MIT ก็ได้สร้างเว็บไซต์ Detect Fakes ที่ประกอบด้วยสื่อจริงและสื่อสังเคราะห์ เพื่อให้ผู้เข้าใช้เว็บไซต์ได้ทดลองแยกแยะว่าสื่อไหนจริง สื่อไหนสังเคราะห์ เป็นต้น (True Digital, 2022) และเทคโนโลยี Model parsing (ตรวจสอบคลิปวิดีโอ ตรวจสอบต่อว่า Deepfake ได้รับการยินยอมหรือไม่) เมื่อได้ตรวจสอบแหล่งที่มาแล้วพบว่า เป็นสื่อปลอม ผู้บริโภคควรตัดทิ้งหรือลบหรือกำจัดทันที หลักโยนิโสมนสิการเมื่อรู้เท่าทันว่าสิ่งใดเป็นคุณค่าแท้หรือเทียม ให้ละคุณค่าเทียมหรือโทษนั้นทันที นี่คือนิยามหรือหน้าที่ต่อสมุทัย

2. ทำให้ประจักษ์แจ้ง

ผู้บริโภคต่อสู้ AI Deepfakes ต้องตั้งเป้าหมายให้ชัดถึงการนำไปใช้ประโยชน์ไม่มีโทษในภายหลัง ยกตัวอย่างประโยชน์ที่เกิดจาก AI Deepfakes การนำผู้ที่เสียชีวิตกลับมาฟื้นคืนชีพในจอภาพยนตร์เพื่อที่แฟนคลับให้หายคิดถึง สามารถไปใช้ประโยชน์ต่อพิพิธภัณฑ์สถานให้ดูเหมือนปลูกฟื้นคืนชีพขึ้นมา ทำให้ดูเหมือนเข้าไปอยู่ในเหตุการณ์ย้อนอดีตนั้น ๆ ได้เหมือนจริง ทำให้ได้ธรรมรสและความทรงจำความรู้ความเข้าใจที่ชัดประจักษ์แจ้งแก่ผู้บริโภคสื่อ AI Deepfakes หรือแม้กระทั่งทางจิตวิทยาต้องการบำบัดฟื้นฟูจิตใจผู้ที่พลัดพรากจากหายไปได้กลับมาพบบน Visual Studio อีกครั้ง AI Deepfake สามารถนำผู้เสียชีวิตกลับมาพบเจอกับคนรักหรือครอบครัวได้อีกครั้ง เพื่อช่วยบรรเทาความโศกเศร้า ความคิดถึง ดังในตัวอย่างของภาพยนตร์โฆษณาไก่ย่างห้าดาว ที่ได้ถ่ายทอดเรื่องราวชีวิตจริงของคุณสุพิชญา ณ สงขลา (โอ) ที่สูญเสียคุณแม่ไป ได้กลับมาเจอกับคุณแม่อีกครั้งผ่านเทคโนโลยี Deepfake และ VR โดยทีมงานได้ใช้รูปภาพกว่า 2,000 รูปและวิดีโอ ให้ระบบ AI ได้เรียนรู้ลักษณะหน้าตาและร่างกายของคุณแม่ เพื่อให้คุณโอได้สามารถกลับมาทานข้าวกับคุณแม่ได้อีกครั้ง หรือการสร้างบรรยากาศการเรียนการสอนบนโลกเสมือนจริงผ่านยุคก่อนประวัติศาสตร์ ยุคประวัติศาสตร์ทำให้ผู้เรียนเข้าถึงเข้าใจและพัฒนาทักษะการเรียนรู้ได้อย่างมีประสิทธิภาพบรรลุผลสัมฤทธิ์ทางการเรียนได้ดีกว่าแบบบรรยายปากเปล่าหรือท่องหนังสือ หรือการเข้ามาแทนอาชีพพนักงานต้อนรับเสมือน สามารถไปประยุกต์ใช้ใน

การพนักงานต้อนรับที่เป็น AI เหมือนคนเสมือน ใช้ในโรงพยาบาล สนามบิน หรือในห้างสรรพสินค้า ตลอดจนมาทำหน้าที่แทนอาชีพนักข่าว นักสื่อสารมวลชน นักออกแบบสื่อดิจิทัล เป็นต้น

เพราะฉะนั้นหน้าที่ต่อนิโรธ ผู้บริโภค AI Deepfakes ต้องนำวิธีคิดแบบบรรณธรรมสัมพันธ์ว่าจะนำเทคโนโลยี AI Deepfakes ไปใช้ประโยชน์ได้อย่างไร ปฏิบัติหรือทำไปเพื่ออะไร เมื่อปฏิบัติแล้วจะนำไปสู่ผลจุดหมายปลายทางอย่างไร ทั้งนี้ผู้บริโภคต้องคิดในใจโดยแยกแยะว่าหลักการและจุดมุ่งหมายของการกระทำทุกอย่างต้องเนื่องกันและสัมพันธ์กันหรือไม่ ถ้าไม่หลักการไม่สัมพันธ์กับเป้าหมายก็ถือว่าไม่เข้าหลักเกณฑ์แห่งวิธีคิดแบบบรรณธรรมสัมพันธ์ เพราะฉะนั้นหลักการกับความมุ่งหมาย เมื่อจะลงมือปฏิบัติธรรมหรือทำตามหลักการอย่างใดอย่างหนึ่งเพื่อให้ประสบผลตามความมุ่งหวัง ไม่คลาดเคลื่อน เลื่อนลอย หรือมกมาย

3. ทำให้เจริญขึ้นหรือเกิดขึ้น

เมื่อกำหนดเป้าหมายชัดเจนโดยในใจให้แยกแยะดีแล้ว ผู้บริโภค AI Deepfakes ต้องทำให้เกิดขึ้นหรือเจริญขึ้นด้วยวิธีคิดแบบเร้าคุณธรรม เป็นวิธีคิดที่มุ่งให้เราทำความดี เกิดความคิดที่ดี ๆ เป็นไปในทางสร้างสรรค์ว่าเทคโนโลยี AI Deepfakes ที่เกิดขึ้นควรจะมีการจัดการอย่างไรถึงจะถูกต้อง เหมาะสมตามหลักของศีลธรรมที่ดีงาม ไม่สร้างความเดือดร้อนหรือเป็นไปเพื่อเบียดเบียนผู้อื่นหรือผู้ที่เราไปล่วงละเมิดสิทธิ์ วิธีคิดแบบเร้าคุณธรรมนี้สามารถแบ่งออกได้เป็น 2 ประการ ได้แก่ (1) การคิดเร้าเร้าโดยมีบุคคลอื่นเป็นเป้าหมาย (หรือมีปัจจัยภายนอกเป็นแรงกระตุ้น) กล่าวคือ AI Deepfakes สามารถพัฒนาหรือต่อยอดนวัตกรรมการออกแบบสื่อสร้างสรรค์เพื่อประโยชน์และคุณค่าแท้ต่อสังคม เพราะ AI Deepfakes สามารถเรียนรู้ข้อมูลที่ใส่เข้ามาในระบบปัญญาประดิษฐ์แล้วสร้างสรรค์สิ่งใหม่ ๆ ตามความต้องการที่ผู้บริโภคได้ให้ข้อมูลเข้าไป การเรียนรู้ของ Deep learning อย่างไม่มีที่สิ้นสุดนี้ย่อมเป็นแรงผลักดันให้ผู้บริโภคสร้างสรรค์สิ่งที่ดีงามและเป็นประโยชน์ตนและผู้อื่นซึ่งถือได้ว่าเป็นความท้าทายและกระตุ้นให้พัฒนาอยู่เสมอ (2) การคิดเร้าเร้าเพื่อสกัดกั้นความประมาท (อกุศล) วิธีคิดที่เร้าเร้าคุณธรรม ใช้เพื่อสกัดกั้นอกุศลคือความประมาทเมื่ออกุศลอันเป็นเหตุของปัญหาต่าง ๆ เกิดขึ้น พระพุทธองค์ได้ตรัสบอกอุปายที่สกัดกั้น แก่ไขกำจัดอกุศลให้หมดไป ด้วยการนึกถึงนิมิต 5 ประการที่ภิกษุควรทำไว้ในใจโดยแยกแยะที่เหมาะสมกับเวลา โดยใช้เทคนิคของพระพุทธเจ้าในการกำจัดอกุศลวิตก 3 ดังปรากฏในวิตักกสัณฐานสูตร (Thai Tripitakas: 12/216-220/226-230) เป็นที่ชัดเจนว่า AI Deepfakes เป็นดาบสองคมมีทั้งคุณและโทษ เพราะฉะนั้นวิธีคิดแบบคุณโทษและทางออกย่อมตอบโจทย์ข้อนี้ว่า ผู้บริโภคไม่ควรประมาทเพราะปัญญาประดิษฐ์พัฒนาอยู่เสมอและไม่มีที่สิ้นสุด ภัยคุกคามในรูปแบบต่าง ๆ ย่อมเกิดขึ้นได้ ถ้าไม่มีระบบป้องกันหรือไม่รู้เท่าทันปัญญาประดิษฐ์ ย่อมเกิดความเสียหายที่ไม่สามารถประเมินได้

อนึ่ง มีหลักโยนิโสมนสิการหนึ่งที่ทำให้ทักษะวิจารณ์ญาณก่อนตัดสินใจ หรือทักษะรู้เท่าทันเทคโนโลยี AI หรือ AI Literacy และหรือทักษะการบริโภคสื่ออย่างสร้างสรรค์ คือ วิธีคิดแบบวิภาษวาท ซึ่งเป็นวิธีคิดที่แจกแจงรายละเอียดในสิ่งที่ประสบพบเห็นแต่ละด้านให้ครบทุกแง่มุม ลักษณะสำคัญ คือ การมองและแสดงความจริง โดยแยกแยะออกให้เห็นแต่ละแง่มุมด้านครบทุกแง่มุมด้าน ไม่ใช่จับเอาแง่มุมหนึ่งแง่เดียวหรือบางแง่มุมมาวินิจฉัยแล้วตีคลุมลงไปอย่างนั้นทั้งหมด หรือประเมินคุณค่าความดีความชั่ว เป็นต้น โดยถือเอาส่วนเดียวหรือบางส่วนเท่านั้น แล้วตัดสินโดยฉับพลัน วิธีคิดแบบวิภาษวาทจะใช้องค์ธรรม 4 เป็นเครื่องมือรู้เท่าทัน AI Deepfakes ได้แก่ เหตุผล แยกแยะ ความจริง และคุณค่า ยกตัวอย่างกรณีศึกษา “การใช้เทคโนโลยี Deepfake ในทางจิตบำบัด” เช่นการรักษาอาการ PTSD หรือ Post-Traumatic Stress Disorder ที่เกิดจากภาวะจิตใจของผู้ป่วยที่ได้รับการกระทบกระเทือนอย่างรุนแรงจากเหตุการณ์เลวร้าย วิภาษวาทแรกเริ่มจากคิดในใจโดยแยกแยะก่อนว่า ความจริงที่เกิดขึ้นกับผู้ป่วยมีอะไรบ้าง และความจริงที่น่าเทคโนโลยี AI Deepfakes ไปใช้แล้วเกิดคุณค่าหรือประโยชน์กับผู้ป่วยได้อย่างไร ถ้าเป็นความจริงแต่เมื่อนำไปใช้แล้วไม่เกิดประโยชน์ต้องตัดทิ้งไป แต่ถ้าเป็นความจริงแล้วเกิดประโยชน์ก็นำมาใช้ ลำดับต่อมาคือเหตุผลในแต่ละขั้นตอน

บ้ำบัต คือการให้ผู้รับการบ้ำบัตทดลองเผชิญหน้ากับความกลัวดังกล่าว ในสภาพแวดล้อมที่ปลอดภัย ภายใต้การดูแลของนักจิตบ้ำบัต ซึ่ง Deepfake จะเข้ามามีส่วนช่วยในการสร้างสื่อสังเคราะห์ในกรณีที่ความกลัวดังกล่าวมีบุคคลเกี่ยวข้อง เพื่อให้ผู้รับการบ้ำบัตสามารถเรียนรู้วิธีบรรเทาและก้าวข้ามผ่านเหตุการณ์ดังกล่าว ในสภาพแวดล้อมที่จะไม่เกิดอันตราย เป็นต้น นอกจากนี้ในมุมมองของการบ้ำบัตแล้ว Deepfake ยังสามารถเข้ามามีส่วนช่วยในการรักษาความเป็นส่วนตัวของผู้เข้ารับการบ้ำบัต ด้วยการนำใบหน้าที่ของบุคคลอื่น มาทดแทนในวิดีโอบันทึกการรักษา โดยยังคงสามารถรักษาน้ำเสียง อากัปกริยา และรายละเอียดอื่น ๆ ได้ เพื่อให้ นักจิตบ้ำบัตสามารถย้อนทวน ส่งต่อ หรือปรึกษาเคสร่วมกับนักจิตบ้ำบัตคนอื่น ๆ โดยไม่เผยแพร่ข้อมูลระบุตัวตนของผู้เข้ารับการบ้ำบัต (True Digital, 2022) อย่างไรก็ตามในกรณีนี้ยังเป็นประเด็นที่ยังต้องศึกษาต่อไป เนื่องจากเกี่ยวข้องกับข้อกฎหมายการขอความยินยอมการใช้ข้อมูลตัวตนในการสร้างสื่อสังเคราะห์ ทั้งนี้ต้องแยกแยะ AI Deepfakes มีผลกระทบต่อผู้บริโภคในภายหลังด้วยหรือไม่ เพราะการใส่ข้อมูลส่วนตัวของเราลงบนแพลตฟอร์ม AI Deepfakes ถือว่าเราได้นำข้อมูลส่วนตัวจำนวนมหาศาลให้ไปกับระบบหรือที่สาธารณะแล้ว ผู้บริโภคนั้นไม่สามารถควบคุมหรือสร้างหลักประกันได้ว่าข้อมูลของเราจะถูกนำไปใช้ในทางใดกับผู้แสวงหาประโยชน์จากภายนอก ดังนั้นวิธีคิดที่วิเศษหาได้ให้แนวทางคิดในใจโดยแยกคายก่อนตัดสินใจลงมือกระทำ โดยเริ่มต้นให้ตั้งสติดูใจคิดก่อนว่า โดยแยกแยะว่าสิ่งใดเป็นความจริงความเท็จ แล้วสืบหาเหตุปัจจัยว่าแหล่งใดที่สามารถคัดกรอง AI Deepfakes เมื่อแน่วแน่แล้วว่าสามารถนำไปใช้ประโยชน์ได้จริง ก็ต้องมาคิดว่า เหตุผลการนำไปใช้คืออะไร และคุณค่าแท้จริงเป็นอย่างไร ผลกระทบที่ตามมาจะมีอะไรบ้าง หมั่นตั้งคำถามเชิงสงสัยหรือตั้งข้อสังเกตและวิเคราะห์ข้อมูลอย่างรอบคอบก่อนตัดสินใจตัดสินใจกับ AI Deepfakes เพื่อไม่ให้ตกเป็นเหยื่อของการใช้เทคโนโลยีมาหลอกลวงในทางที่ผิดต่อไป

สรุป

"AI Deepfake" เป็นเทคโนโลยีที่ให้ทั้งคุณประโยชน์และโทษ AI Deepfakes เป็น Generative AI รูปแบบหนึ่งที่สามารถสร้างสื่อสังเคราะห์ ทั้งภาพนิ่ง เสียง ภาพเคลื่อนไหวพร้อมลอกเลียนเสียงจากมนุษย์จริง มาพูดคุยจริง ๆ จนแทบแยกไม่ออกซึ่งปัญญาประดิษฐ์สามารถมีระบบการเรียนรู้ หรือ Deep Learning จากลักษณะภายนอกของบุคคล เช่น สีผิว ตา ปาก จมูก รูปลักษณ์ต่าง ๆ ด้วยความฉลาดล้ำจากการเรียนรู้ข้อมูลมหาศาลของ Generative AI

การรู้เท่าทัน "AI Deepfake" ด้วยหลักโยนิโสมนสิการ ต้องเริ่มต้นที่สติ โดยไม่จำเป็นต้องใช้วิธีคิดทั้ง 10 รูปแบบทั้งนี้ขึ้นอยู่กับเหตุการณ์จริงตามเหตุปัจจัยที่เกี่ยวข้อง แต่ทั้งนี้มีแกนหลักในการนำไปใช้ให้รู้เท่าทัน "AI Deepfake" คือ ใช้กรอบกิจในอริยสัจแล้วนำรูปแบบวิธีคิดตามเหตุปัจจัยที่เกี่ยวข้องจริงเข้าไปในรายละเอียดที่เกิดขึ้น โดยเครื่องมือที่เป็นกลไกขับเคลื่อน คือ องค์ธรรม 4 ได้แก่ ความจริง เหตุผล แยกแยะ และคุณค่า

References

Appman. (2023). Dealing with AI weapons used in identity fraud with Deep Fake Detection. Retrieved September 4 , 2023 , from https://www.appman.co.th/how-to-deep-fake-detection-with-liveness-detection/#_brief/

- Colin, Campbell., Kirk, Plangger., Sean, Sands., Jan, Kietzmann., Kenneth, Bates. (2022). How Deepfakes and Artificial Intelligence Could Reshape the Advertising Industry. *Journal of Advertising Research*, doi: 10.2501/jar-2022-017
- Damir, Yalov., Danil Myakin. (2023). Top 10 AI Deepfake Generators for Images and Videos in 2023. Retrieved September 06, 2023, from <https://mpost.io/th/top-ai-deepfake-generators/>
- Eddy. (2024). *Top 10 Deepfake Video Creation Tools for Fun Don't use anything before you read*. Retrieved January 30, 2024, from <https://school.mangoanimate.com/th/video-creators-deepfake-10-top/>
- Gamingdose. (2021). *What is DeepFake? Why is it scarier than you think?*. Retrieved July 29, 2021, from <https://www.gamingdose.com/feature/deepfake/>
- Hina, Shahzad., Furqan, Rustam., Emmanuel, Soriano, Flores., Juan, Luís, Vidal, Mazón., Isabel, de la, Torre, Díez., Imran, Ashraf. (2022). *A Review of Image Processing Techniques for Deepfakes. Sensors*, doi: 10.3390/s22124556
- Mahachulalongkornrajavidyalaya University. (1996). *Thai Tripitakas*. Bangkok: Mahachulalongkornrajavidyalaya Press.
- Muhammad, Salihin, Saealal., M., Z., Ibrahim., David, Mulvaney., Mohd, Ibrahim, Shapiai., Norasyikin, Binti, Fadilah. (2022). *Using cascade CNN-LSTM-FCNs to identify AI-altered video based on eye state sequence*. PLOS ONE, doi: 10.1371/journal.pone.0278989
- Narora, Amezaga., Jeremy, Hajek. (2022). Availability of Voice Deepfake Technology and its Impact for Good and Evil. doi: 10.1145/3537674.3554742
- Office of Information Technology Administration for Educational Development. (2024). *Get to know AI Deepfake technology, if used incorrectly it becomes a double-edged sword!! Nowadays, artificial intelligence technology or AI (AI: Artificial Intelligence) has been developed and used. Used in many industries, especially the extremely intelligent Generative AI*. Retrieved March 29, 2024, from <https://www.uni.net.th/index.php/article/2513/>
- PPTV Online. (2024). *Scammers use Deepfake to impersonate the head of finance and scam them out of 900 million baht*. Retrieved February 5, 2024, from <https://www.pptvhd36.com/news/abroad/216306/>
- Rattamangam, S., Thitipasitthikorn, P., Yusamran, P., Wichai, P. (2022). Communication for Social Development in the Artificial Intelligence Era. *Journal of Buddhist Studies Vanam Dongrak*, 9(1), 179-189.
- Techsauce Team. (2024). *Teach you how to spot Deepfakes and prevent yourself from being tricked. Be aware of scams and false information*. Retrieved February 16, 2024, from <https://techsauce.co/tech-and-biz/how-to-counteract-deepfake>

Thai PBS Sci & Tech. (2023). *Be aware of "Deepfake" clips edited to deceive the world by the smoothest AI. Suggestions on how to capture landmarks.* Retrieved July 27, 2023, from <https://www.thaipbs.or.th/now/content/219/>

Thairath Online. (2023). *Almost smooth A criminal pretends to be a police officer on a video call, but sits still and can only move his mouth.* Retrieved March 10, 2023, from <https://www.thairath.co.th/news/society/2338077/>

True Digital. (2022). *A compilation of all stories about Deepfake, the villain of technology or the best helper of 2022.* Retrieved June 13, 2022, from [https://www.truedigital.com/post/ A compilation of all stories about Deepfake, the villain of technology or the best helper of 2022/](https://www.truedigital.com/post/A-compilation-of-all-stories-about-Deepfake-the-villain-of-technology-or-the-best-helper-of-2022/)

Uninet. (2023). *Get to know AI Deepfake technology, if used incorrectly it becomes a double-edged sword!! Nowadays, artificial intelligence technology or AI (AI: Artificial Intelligence) has been developed and used. Used in many industries, especially the extremely intelligent Generative AI.* Retrieved August 1, 2023, from <https://www.uni.net.th/index.php/article/2513/>

