

AI Ethics

จริยธรรมด้านปัญญาประดิษฐ์ ในประเทศไทย

ปัจจุบัน เทคโนโลยีปัญญาประดิษฐ์ (Artificial Intelligence: AI) ได้เข้ามามีบทบาทในกระบวนการทำงานหลากหลายมากขึ้น ทั้งด้านการแพทย์ การศึกษาวิจัย การตลาด ศิลปะ วรรณกรรม หรือแม้แต่กระบวนการยุติธรรม อย่างไรก็ตามเมื่อมีประโยชน์มหาศาล ก็เป็นที่รู้กันทั่วก็อาจกลายเป็นเครื่องมือที่ถูกนำไปใช้ในทางที่ผิดได้เช่นกัน และอาจส่งผลกระทบต่อร้ายแรงตามมาในอนาคตหากไม่มีการกำกับดูแล หรือควบคุมจริยธรรมการใช้ประโยชน์ให้ถูกต้องชอบธรรม

ดร.ไอแซค อสิมอฟ นักเขียนนิยายวิทยาศาสตร์ (Sci-Fi) และนักชีวเคมีชื่อดังชาวอเมริกันเชื้อสายรัสเซีย ได้กล่าวถึงโลกในอนาคต และ “กฎ 3 ข้อของหุ่นยนต์” (three laws of robotic) ไว้ในนิยายของเขา จนถูกนำมาอ้างอิงกลายเป็นกฎที่ใช้กันอย่างแพร่หลายในภาพยนตร์ Sci-Fi ต่อเนื่องมา มีกฎว่า

1

หุ่นยนต์มิอาจกระทำการอันตรายต่อผู้ที่เป็มนุษย์ หรือนิ่งเฉยปล่อยให้ผู้ที่เป็มนุษย์ตกอยู่ในอันตรายได้ (A robot may not harm a human being, or, through inaction, allow a human being to come to harm.)

2

หุ่นยนต์ต้องเชื่อฟังคำสั่งที่ได้รับจากผู้ที่เป็นมนุษย์ เว้นแต่คำสั่งนั้นๆ ขัดแย้งกับกฎข้อแรก (A robot must obey the orders given to it by human beings, except where such orders would conflict with the First Law.)

3

หุ่นยนต์ต้องปกป้องสถานะความมีตัวตนของตนเองไว้ トラบเท่าที่การกระทำนั้นมิได้ขัดแย้งต่อกฎข้อแรกหรือกฎข้อที่สอง (A robot must protect its own existence, as long as such protection does not conflict with the First or Second Law.)

ในปี 2016 สัตยา นาเทลลา (Satya Nadella) CEO ประธานเจ้าหน้าที่บริหาร บริษัท Microsoft เคยเขียนบทความลงตีพิมพ์เผยแพร่ในนิตยสาร Slate ถึงภาพอนาคตความร่วมมือระหว่างมนุษย์กับปัญญาประดิษฐ์ (The Partnership of the Future) ไว้ว่า ต้องเริ่มจากการออกแบบ AI ที่ดี อันประกอบด้วยหลัก 6 ประการ ได้แก่

1. AI ต้องช่วยเหลือมนุษยชาติ คือ ได้รับการออกแบบเพื่อทำงานร่วมกับมนุษย์ ช่วยทำงานที่เสี่ยงเพื่อปกป้องคนทำงาน

2. AI ต้องมีความโปร่งใส คือ มนุษย์ต้องเข้าใจการทำงานของ AI และรู้วิธีที่ AI ใช้ประมวลผลและวิเคราะห์ข้อมูล

3. AI ต้องรักษาคำมั่นศีลธรรม คือ ต้องช่วยส่งเสริมความหลากหลายทางวัฒนธรรมและความเป็นมนุษย์

4. AI ต้องฉลาดรู้การปกป้องข้อมูลส่วนบุคคล คือ ต้องรักษาความปลอดภัยข้อมูลส่วนบุคคล ข้อมูลกลุ่มผู้ใช้ ได้อย่างชาญฉลาด

5. AI ต้องมีความรับผิดชอบทางอัลกอริทึม คือ ต้องสามารถแก้ไขความผิดพลาดที่เกิดขึ้นได้เอง

6. AI ต้องป้องกันความลำเอียงหรืออคติ คือ ต้องได้รับการออกแบบเพื่อป้องกันการเลือกปฏิบัติ และให้ผลการทำงานที่เป็นธรรม

จะเห็นได้ว่าในบริษัทขนาดใหญ่ที่พัฒนาลงทุนด้านเทคโนโลยีจริงๆ หรือแม้แต่ในนิยายที่จินตนาการถึงเทคโนโลยี

แห่งโลกอนาคต ก็ยังต้องพูดถึงกฎเกณฑ์กำกับการใช้ปัญญาประดิษฐ์ไว้ “จริยธรรมด้าน AI” จึงมาเป็นประเด็นสำคัญที่ต้องให้ความสนใจและสร้างความตระหนักด้านการใช้ปัญญาประดิษฐ์ทั้งต่อตัวผู้พัฒนาและผู้ใช้งานควบคู่กันไปด้วย เมื่อโลกทุกวันนี้มีการใช้ AI ช่วยทำงานกันอย่างแพร่หลายในวงกว้างเพิ่มขึ้นมากและรวดเร็ว จนอาจก่อให้เกิดปัญหาด้านจริยธรรมขึ้นมาได้หากไม่มีการกำกับดูแล ตัวอย่างกรณีศึกษาปัญหจริยธรรมจากการใช้ AI เช่น

1. ใช้ AI ในการวิเคราะห์พฤติกรรมของผู้บริโภค

การเก็บข้อมูลพฤติกรรมของผู้ใช้งานผ่านอินเทอร์เน็ตและโซเชียลมีเดีย เพื่อใช้ในการโฆษณาหรือแนะนำสินค้า แม้จะช่วยให้การตลาดมีประสิทธิภาพมากขึ้น แต่ก็อาจละเมิดความเป็นส่วนตัวและทำให้ผู้คนรู้สึกถูกคุกคามขึ้นได้

2. ใช้ AI ในระบบการตัดสินใจ

การใช้ AI ในการพิจารณาหรือประเมินคดีความ หรือวิเคราะห์ความเสี่ยงในการกระทำผิดซ้ำ หาก AI ถูกฝึกด้วยข้อมูลที่มีความลำเอียง อาจทำให้เกิดการตัดสินใจที่ไม่ยุติธรรม และส่งผลกระทบต่อชีวิตของผู้ถูกตัดสินขึ้นได้

3. ใช้ AI ในการประเมินและคัดเลือกผู้สมัครงาน

หาก AI นั้นถูกฝึกด้วยข้อมูลที่มีการกีดกันแบ่งแยกทางเพศ เชื้อชาติ หรืออายุ ที่ไม่เหมาะสมกับหลักสิทธิมนุษยชน ก็อาจทำให้เกิดการเลือกปฏิบัติในกระบวนการจ้างงาน การประกอบอาชีพ ขึ้นได้

องค์การการศึกษา วิทยาศาสตร์และวัฒนธรรมแห่งสหประชาชาติ (United Nations Educational, Scientific and Cultural Organization: UNESCO) ได้จัดตั้ง Global AI Ethics and Governance Observatory หรือศูนย์สังเกตการณ์จริยธรรมและการกำกับดูแลด้าน AI โลกขึ้น เพื่อการกำกับดูแลการใช้ AI อย่างถูกต้อง มีจริยธรรมและรับผิดชอบต่อสังคม ซึ่งถือเป็นหนึ่งในความท้าทายที่สำคัญที่สุดในปัจจุบัน โดยต้องอาศัยการเรียนรู้ร่วมกันจากกรณีศึกษา การวิจัย ชุดเครื่องมือและแนวทางปฏิบัติที่ดีจากที่เกิดขึ้นทั่วโลก รวบรวมไว้เป็นแหล่งข้อมูลระดับโลก สำหรับผู้กำหนดนโยบายภาครัฐ หน่วยงานกำกับดูแล นักวิชาการ ภาคเอกชน และสังคมมนุษย์ ในการค้นหาวិธีแก้ปัญหาที่เร่งด่วนที่สุดที่เกิดจากปัญญาประดิษฐ์

โดยวางกรอบค่านิยมหลัก 4 ประการ ไว้เป็นรากฐานให้ระบบ AI ไว้ 1) สิทธิมนุษยชนและศักดิ์ศรีความเป็นมนุษย์ 2) การใช้ชีวิตอย่างสงบสุข สังคมที่ยุติธรรมและเชื่อมโยงกัน 3) การสร้างหลักประกันความหลากหลายและความครอบคลุม 4) สิ่งแวดล้อมและระบบนิเวศเจริญรุ่งเรือง ซึ่งสามารถศึกษาเพิ่มเติมได้จากเอกสาร UNESCO's Recommendations on the Ethics of Artificial Intelligence: key facts



สำหรับในประเทศไทย มีอย่างน้อย 2 หน่วยงานที่ออกแนวปฏิบัติด้านจริยธรรมปัญญาประดิษฐ์ สำหรับประเทศไทย ขึ้น อาทิ

สำนักงานคณะกรรมการดิจิทัลเพื่อเศรษฐกิจและสังคมแห่งชาติ (สดช.) ในกำกับกระทรวงดิจิทัลเพื่อเศรษฐกิจและสังคม ได้จัดทำ “เอกสารแนวปฏิบัติจริยธรรมปัญญาประดิษฐ์” (Thailand AI Ethics Guideline) ขึ้น โดยกำหนดหลักการทางจริยธรรมปัญญาประดิษฐ์ ไว้ 6 หลักการด้วยกัน ประกอบด้วย

THAILAND AI ETHICS GUIDELINE

by MDES



1. ความสามารถในการแข่งขันและการพัฒนาอย่างยั่งยืน (Competitiveness and Sustainability Development)

- AI ควรถูกสร้างและใช้งานเพื่อสร้างประโยชน์และความผาสุกให้แก่มนุษย์ สังคม เศรษฐกิจ และสิ่งแวดล้อมอย่างยั่งยืน

- AI ควรถูกใช้งานเพื่อเพิ่มความสามารถในการแข่งขันและสร้างความเจริญให้กับมนุษย์ สังคม ประเทศ ภูมิภาค และโลกอย่างเป็นธรรม
- AI ควรได้รับการวิจัยและพัฒนาอย่างต่อเนื่อง เพื่อให้มนุษย์เกิดการสร้างสรรค์นวัตกรรมและอุตสาหกรรมใหม่

2. ความสอดคล้องกับกฎหมาย จริยธรรม และมาตรฐานสากล (Law Ethics and International Standards)

- AI ควรได้รับการวิจัย ออกแบบ พัฒนา ให้บริการ และใช้งานสอดคล้องกับกฎหมาย บรรทัดฐาน จริยธรรม คุณธรรมของมนุษย์ และมาตรฐานสากล โดยเคารพต่อความเป็นส่วนตัว เสรี สิทธิเสรีภาพ และสิทธิมนุษยชน
- การออกแบบ AI ควรใช้หลักการมนุษย์เป็นศูนย์กลางและเป็นผู้ตัดสินใจ
- AI ไม่ควรถูกใช้ในการกำหนดชะตาชีวิตของมนุษย์

3. ความโปร่งใสและภาระรับผิดชอบ (Transparency and Accountability)

- AI ควรได้รับการวิจัย ออกแบบ พัฒนา ให้บริการ และใช้งานด้วยความโปร่งใส สามารถอธิบาย และคาดการณ์ได้ รวมถึงสามารถตรวจสอบกิจกรรมต่างๆ ที่เกิดขึ้นย้อนหลังได้
- AI ควรมีความสามารถในการสืบย้อนหลัง เฝ้าระวัง ตรวจสอบความผิดปกติและวินิจฉัยปัญหา ความล้มเหลวได้
- ผู้วิจัย ผู้ออกแบบ ผู้พัฒนา ผู้ให้บริการและผู้ใช้งาน AI ควรมีภาระความรับผิดชอบต่อผลกระทบที่เกิดขึ้นจาก AI ตามภาระหน้าที่ตน

4. ความมั่นคงปลอดภัยและความเป็นส่วนตัว (Security and Privacy)

- AI ควรถูกสร้างเพื่อบริการ แต่ไม่ควรถูกใช้เพื่อหลอกลวง ต่อด้าน และคุกคามมนุษย์
- AI ควรได้รับการออกแบบโดยใช้หลักการป้องกันความเสี่ยง เพื่อป้องกันการโจมตีจากภัยคุกคาม

เพื่อรักษาไว้ซึ่งความมั่นคงปลอดภัยของข้อมูล และระบบ รวมถึงการคุ้มครองข้อมูลส่วนบุคคล จริยธรรม และความปลอดภัยของชีวิตและสิ่งแวดล้อมภายนอกตลอดวัฏจักรชีวิตของระบบ มีความสามารถในการตรวจสอบ รายงานและตอบสนองเพื่อหลีกเลี่ยงหรือลดผลกระทบ

- AI ควรมีการเฝ้าระวังให้มนุษย์แทรกแซงระบบเพื่อควบคุมความเสี่ยงที่อาจมีผลกระทบกับมนุษย์ได้
- หน่วยงานรัฐควรวางแผนกำกับดูแลการพัฒนา และให้ความร่วมมือกับนานาชาติในการหลีกเลี่ยงการแข่งขันสร้างอาวุธอัตโนมัติจากปัญญาประดิษฐ์ที่ร้ายแรง

5. ความเท่าเทียม หลากหลาย ครอบคลุม และเป็นธรรม (Fairness)

- การออกแบบและพัฒนา AI ควรคำนึงถึงความหลากหลาย หลีกเลี่ยงการผูกขาด ลดการแบ่งแยกและเอื้อเอียง เพื่อก่อให้เกิดประโยชน์ต่อผู้คนจำนวนมากเท่าที่จะทำได้โดยเฉพาะกลุ่มคนผู้ด้อยโอกาสในสังคม
- การตัดสินใจที่เกี่ยวข้องกับวิจัย ออกแบบ พัฒนา ให้บริการและใช้งาน AI ควรสามารถพิสูจน์ถึงความ เป็นธรรมได้

6. ความน่าเชื่อถือ (Reliability)

- AI ควรได้รับการสนับสนุนให้มีความน่าเชื่อถือ และความมั่นใจในการใช้งานต่อสาธารณะ
- AI ควรสามารถคาดการณ์ ตัดสินใจ และให้คำแนะนำได้อย่างแม่นยำถูกต้อง สร้างผลลัพธ์ที่สามารถเชื่อถือได้และสร้างใหม่ได้เมื่อต้องการ
- AI ควรมีการควบคุมคุณภาพและตรวจสอบความครบถ้วนสมบูรณ์ของข้อมูล
- AI ควรมีกระบวนการและช่องทางรับผลสะท้อนกลับ (feedback) จากผู้ใช้งานเพื่อให้ผู้ใช้งานสามารถแจ้งความต้องการเพิ่มเติม รับเรื่องร้องเรียน แจ้งปัญหาของระบบที่ตรวจสอบพบ และให้ข้อเสนอแนะได้โดยง่ายและรวดเร็ว

นอกจากนี้ ในเอกสารยังมีกรณีศึกษา กรอบวิธีการปฏิบัติและกิจกรรมตามแนวทางจริยธรรม ตลอดจนหลักเกณฑ์ และรูปแบบการประเมินกระบวนการกำกับดูแลจริยธรรม ที่ออกแบบให้เหมาะสมกับแต่ละผู้เกี่ยวข้องกับการพัฒนา จริยธรรมปัญญาประดิษฐ์อีกด้วย

คณะกรรมการจริยธรรมปัญญาประดิษฐ์ สำนักงานพัฒนาวิทยาศาสตร์และเทคโนโลยีแห่งชาติ (สวทช.) ในกำกับกระทรวงการอุดมศึกษา วิทยาศาสตร์ วิจัยและนวัตกรรม ได้จัดทำ “**แนวปฏิบัติจริยธรรมด้านปัญญาประดิษฐ์**” ขึ้นมาสำหรับใช้เป็นแนวทางการวิจัยและใช้งานปัญญาประดิษฐ์ของประเทศไทย โดยมีหลักการด้านจริยธรรมปัญญาประดิษฐ์ ที่ควรคำนึงถึงมี 7 ข้อ ดังนี้



1. ความเป็นส่วนตัว (privacy) : AI ควรถูกออกแบบให้สามารถป้องกันความเป็นส่วนตัวและเคารพต่อสิทธิเสรีภาพของบุคคล ศักดิ์ศรีของมนุษย์ และต้องเคารพต่อกฎระเบียบ หรือกฎหมายที่เกี่ยวข้องกับการคุ้มครองข้อมูลส่วนบุคคลด้วย เช่น พระราชบัญญัติคุ้มครองข้อมูลส่วนบุคคล พ.ศ. 2562 (PDPA) เป็นต้น

2. ความมั่นคงและปลอดภัย (security and

safety) : AI ควรถูกสร้างให้มีความมั่นคงและปลอดภัย รวมถึงป้องกันภัยอันตรายที่อาจเกิดขึ้นต่อมนุษย์ สิ่งแวดล้อม สังคม เศรษฐกิจ และประเทศ จากการใช้งานปัญญาประดิษฐ์ที่มากเกินไป (overused) และที่ไม่ถูกต้องเหมาะสม (misused) การใช้งานและการตัดสินใจของ AI ควรเกิดจากความตั้งใจของมนุษย์ หรือมีกลไกให้มนุษย์สามารถแทรกแซงการดำเนินการต่างๆ ได้

3. ความไว้วางใจ (reliability) :

ผู้พัฒนา AI ต้องสามารถสร้างความไว้วางใจและเชื่อมั่นในการใช้งานให้แก่สาธารณะชนได้ว่า AI ให้ผลที่แม่นยำถูกต้อง สร้างผลลัพธ์ที่เชื่อถือได้ และสร้างผลแบบเดียวกันใหม่ได้เสมอ (reproducible) รวมถึงมีการควบคุมคุณภาพของข้อมูลที่น่ามาใช้งาน เพื่อป้องกันไม่ให้เกิดการตัดสินใจที่ผิดพลาด ซึ่งอาจส่งผลกระทบต่อความน่าเชื่อถือของระบบ AI ได้

4. ความเป็นธรรม เท่าเทียม และไม่แบ่งแยก

(fairness and non-discrimination) : AI ควรถูกออกแบบและนำไปใช้งานเพื่อส่งเสริมความเป็นธรรม ความเท่าเทียม ความหลากหลาย ความสามัคคี และความยุติธรรมของคนทุกกลุ่มในสังคม โดยไม่เกิดความอคติหรือความเอนเอียงใดๆ รวมถึงให้โอกาสประชาชนทุกคนในสังคม เช่น กลุ่มคนผู้ด้อยโอกาส ผู้พิการและผู้ทุพพลภาพ ให้ได้รับประโยชน์จาก AI ได้อย่างทั่วถึงและเท่าเทียม โดยไม่แบ่งแยกเชื้อชาติ สีผิว และไม่ก่อให้เกิดความเหลื่อมล้ำทางสังคม

5. ความโปร่งใสและอธิบายได้ (transparency

and explainability) : AI ควรได้รับการออกแบบและนำไปใช้ โดยให้มนุษย์สามารถรู้ได้ว่ากำลังใช้ AI อยู่ เข้าใจได้ว่าข้อมูลถูกนำไปใช้อย่างไร เข้าใจได้ถึงกระบวนการการตัดสินใจ คาดการณ์การกระทำต่างๆ ได้ รวมถึงกำกับดูแลและตรวจสอบ AI ได้ โดยควรทำให้มีการแปลผลการดำเนินการของระบบให้เป็นข้อมูลที่สามารถอธิบายและเข้าใจได้โดยมนุษย์ สามารถตรวจ-

สอบย้อนกลับไปยังแหล่งที่มาของชุดข้อมูลที่ได้รับ กระบวนการทำงาน และการตัดสินใจของ AI รวมถึงสถานที่ เวลา และวิธีการนำไปใช้ เพื่อใช้เฝ้าระวัง ตรวจสอบความผิดปกติ วินิจฉัยปัญหา และหาผู้รับผิดชอบได้อย่างมีประสิทธิภาพ อีกทั้งเพื่อช่วยสร้างความเชื่อมั่นให้แก่ผู้ใช้งาน

6. ภาระความรับผิดชอบ (accountability) :

การวิจัย ออกแบบ หรือพัฒนา AI ต้องมีกลไกที่ทำให้เกิดความมั่นใจถึงการรับผิดชอบต่อผลกระทบที่เกิดจาก AI ของผู้ที่มีส่วนร่วมในการวิจัย ออกแบบ พัฒนา และนำไปใช้งาน ซึ่งต้องสามารถตรวจสอบย้อนกลับถึงผู้รับผิดชอบได้โดยชัดเจน รวมถึงมีกลไกแก้ไขปัญหา หรือรับผิดชอบต่อความเสียหายที่อาจเกิดขึ้น ตามภาระหน้าที่ของตนได้อย่างเพียงพอ รวมถึงมีการวางแผนถึงการบริหารจัดการความเสี่ยง หรือผลกระทบในระยะยาวที่อาจเกิดขึ้น

7. มนุษย์เป็นผู้ควบคุมปัญญาประดิษฐ์ เพื่อความ

ยั่งยืนของมนุษยชาติ (human oversight and human agency) :

ในการออกแบบและการนำ AI ไปใช้งานนั้น ควรกำหนดให้คำนึงถึงมนุษย์เป็นหลัก (human centric) และคงไว้ซึ่งความสามารถในการควบคุม AI และสิทธิให้มนุษย์เป็นผู้ตัดสินใจในขั้นตอนการตัดสินใจที่เป็นกระบวนการสำคัญ นอกจากนี้ระบบ AI จะต้องถูกออกแบบและนำมาใช้เพื่อให้เกิดประโยชน์ต่อมนุษยชาติ ไม่ก่อให้เกิดอันตรายต่อมนุษย์ คำนึงถึงความเป็นอยู่ที่ดี และส่งเสริมคุณค่าของมนุษย์ รวมถึงเป็นการนำ AI มาใช้เพื่อให้เกิดการพัฒนาที่ยั่งยืน ต่อทั้งสังคมและสิ่งแวดล้อม ตลอดจนทำให้เกิดการพัฒนาในด้านอื่นๆ ต่อไป

โดยแนวปฏิบัติจริยธรรมด้านปัญญาประดิษฐ์ของ สวทช. ยังมีการให้รายละเอียดไปถึง คุณสมบัติผู้วิจัยปัญญาประดิษฐ์และผู้ร่วมโครงการ การดำเนินการตามหลักการจริยธรรมปัญญาประดิษฐ์ การบริหารจัดการเพื่อลดความเสี่ยงและผลกระทบจากการพัฒนาและการนำปัญญาประดิษฐ์ไปใช้ พร้อมกรณีศึกษาที่เกี่ยวข้องกับจริยธรรมด้านปัญญาประดิษฐ์ไว้ด้วย

การพัฒนาและใช้งาน AI เป็นสิ่งที่มีศักยภาพในการเปลี่ยนแปลงโลกในทางที่ดี แต่ก็มีความเสี่ยงที่อาจเกิดขึ้นได้ หากไม่มีการบริหารจัดการอย่างเหมาะสม จริยธรรมด้าน AI เป็นแนวทางที่ช่วยให้เราสามารถใช้เทคโนโลยีนี้ได้อย่างมีความรับผิดชอบและสร้างสรรค์ เพื่อให้เทคโนโลยีนี้ไม่เป็นเพียงเครื่องมือที่มีประสิทธิภาพ แต่ยังเป็นสิ่งที่สร้างความเป็นธรรม และคุณค่าต่อสังคมในระยะยาวของมนุษยชาติด้วย 🌐

เอกสารอ้างอิง

- สำนักงานคณะกรรมการดิจิทัลเพื่อเศรษฐกิจและสังคมแห่งชาติ. 2565. เอกสารแนวปฏิบัติจริยธรรมปัญญาประดิษฐ์ Thailand AI Ethics Guideline. [ออนไลน์]. เข้าถึงได้จาก: [https://www.onde.go.th/assets/portals/1/files/ThailandAIEthicsGuideline\(Whitepaper\)EditVersion.pdf](https://www.onde.go.th/assets/portals/1/files/ThailandAIEthicsGuideline(Whitepaper)EditVersion.pdf), [เข้าถึงเมื่อ 3 สิงหาคม 2567].
- สำนักงานพัฒนาวิทยาศาสตร์และเทคโนโลยีแห่งชาติ. 2565. แนวปฏิบัติจริยธรรมด้านปัญญาประดิษฐ์. [ออนไลน์]. เข้าถึงได้จาก: <https://waa.inter.nstda.or.th/stks/pub/ori/docs/20220831-aw-book-ai-ethics-guideline.pdf>, [เข้าถึงเมื่อ 3 สิงหาคม 2567].
- Isaac Asimov. 1942. Three laws of robotics. *Britannica*. [online]. Available at: <https://www.britannica.com/topic/Three-Laws-of-Robotics>, [accessed 3 August 2024].
- Satya Nadella. 2016. The Partnership of the Future: Microsoft's CEO explores how humans and A.I. can work together to solve society's greatest challenges. *SLATE*. [online]. Available at: <https://slate.com/technology/2016/06/microsoft-ceo-satya-nadella-humans-and-a-i-can-work-together-to-solve-societys-challenges.html>, [accessed 3 August 2024].
- UNESCO. 2024. Ethics of Artificial Intelligence: The Recommendation. [online]. Available at: <https://www.unesco.org/en/artificial-intelligence/recommendation-ethics>, [accessed 3 August 2024].
- UNESCO. 2024. UNESCO's Recommendation on the Ethics of Artificial Intelligence: key facts. [online]. Available at: <https://unesdoc.unesco.org/ark:/48223/pf0000385082.page=12>, [accessed 3 August 2024].