

עקרונות מדיניות, רגולציה ואתיקה בתחום הבינה המלאכותית

2023



14 דצמבר 2023
ב' בטבת תשפ"ד

דברי שר החדשנות, המדע והטכנולוגיה

מהפכת 'הבינה המלאכותית' מתחוללת במהירות לנגד עינינו. המהפכה מובלת על ידי השוק הפרטי, אך ממשלות בכל העולם נדרשות לגבש מדיניות ברמה הלאומית ולאור סטנדרטים בינלאומיים בהתהוות.

גם ממשלות ישראל נדרשו לסוגיה ואכן בשנתיים האחרונות התקבלו שתי החלטות מהותיות בעניין: החלטות הממשלה מס' 212 מיום 1.8.2021 (להלן – החלטה 212) ומס' 173 מיום 24.2.2023 (להלן – החלטה 173). שתיהן עסקו בקידום החדשנות, בעידוד צמיחת ההייטק ובחיזוק המובילות הטכנולוגית-מדעית של מדינת ישראל.

לאחר השלמת עבודת מטה מקיפה ומקצועית, החלטתי לאמץ את מסקנותיו של המסמך שעניינו מדיניות רגולציה ואתיקה בתחום הבינה המלאכותית בישראל.

ראש הממשלה, בנימין נתניהו, הציב השנה את נושא הבינה המלאכותית כיעד מרכזי וחיוני, במטרה למצב את ישראל כמובילה עולמית בתחום כדי לשקף את הפוטנציאל הרב למען עתיד האנושות. משרד החדשנות, המדע והטכנולוגיה בראשותי, מוביל את התוכנית הלאומית לקידום הבינה המלאכותית.

הצבנו מספר יעדים ובראשם: חיזוק המובילות הישראלית בתחום הבינה המלאכותית לטובת שיפור איכות החיים, והבטחת חוסנה הלאומי וצמיחתה הכלכלית של ישראל. במסגרת זו, התחייבנו שעד סוף שנת 2023 נפרסם לציבור מסמך 'מדיניות לאתיקה ורגולציה בבינה מלאכותית'.

הפוטנציאל המהפכני של הטכנולוגיה משפיע במהירות על שיפור איכות אזרחי מדינת ישראל: ברפואה, מדע, טכנולוגיה, חקלאות, חינוך ביטחון ועוד. אולם, לצד היתרונות הרבים ישנם גם סיכונים רבים - יש לוודא כי האדם נשאר במרכז קבלת ההחלטות לצד ערכת עקרונות אתיים. עקרונות אלו מאפשרים, פיתוח ושימוש אחראי ומאפשר תוך הקפדה על שמירת זכויות יסוד ואינטרסים ציבוריים, כבוד האדם ופרטיותו, שוויון ומניעת אפליה ושקיפות מלאה.

פיתוח ושימוש נכון, אחראי וחכם בבינה מלאכותית מהימנה הוא אמצעי חיוני לעידוד צמיחה, פיתוח בר-קיימא, רווחה חברתית וקידום המובילות הישראלית בתחום החדשנות והטכנולוגיה.

מדינת ישראל מדורגת במיקום גבוה בעולם בטבלת המדינות המובילות את תחום הבינה המלאכותית. עם פרסום מסמך זה אנו משפרים את מעמדנו העולמי באסדרת וארגון התחום, למען העתיד הכלכלי של ישראל, ולקראת דו"חות עתידיים של סוכנויות הדירוג.

כבר היום קיימות בישראל כ-2,200 חברות הייטק, שמפתחות או עושות שימוש משמעותי בטכנולוגיה זו. עלינו לקדם אותן ולאפשר להן לפתח שיטות וכלים חדשניים למען עתיד האנושות כולה, ולסייע במתן הכללים שיאפשרו לפתח חדשנות אחראית.

בימים אלו, בשיאה של מלחמתנו באויבים בכל הגזרות, כאשר עיני העולם נשואות אלינו, יש חשיבות יתרה למיצובה של ישראל כמעצמה מובילה בתחום. בעוד הפיתוח העולמי נמשך ביתר שאת, תפקידנו לוודא כי מדינת ישראל לא תישאר מאחור ולא תפגע בעקבות המלחמה בכל הנוגע לתשתית ולפיתוח החדשנות והטכנולוגיה.

במטרה לצמצם סיכונים ולתת מענה לאתגרים שעולים מהתחום, משרד החדשנות, המדע והטכנולוגיה יחד עם משרד המשפטים גיבש מסמך מקצועי עליו נשענים עקרונות המדיניות בנושא רגולציה ואתיקה עבור פיתוח ושימוש מערכות מבוססות בינה מלאכותית בישראל. המסמך פורסם לפני כשנה להערות הציבור, ולאחר הליך מקיף של איסוף ובחינה של הערות הציבור, עודכן המסמך ועתה הוא מהווה מסמך סופי המציע את מדיניות הרגולציה והאתיקה של ממשלת ישראל בתחום הבינה המלאכותית.

עיקרי המדיניות קובעים כי הרגולציה בתחום הבינה המלאכותית תיעשה באופן ענפי ולא רוחבי, תהא מתקדמת וגמישה ותתאם ככל הניתן למדיניות הרגולציה המתהווה בזירה הבינלאומית.

בנוסף, המסמך מציע להקים מוקד ידע ממשלתי, שנועד לסייע לגופי הממשלה הרלוונטיים בגיבוש הרגולציה ובהטמעתה, לתאם ביניהם ולהיות אחראי להניח בפני מקבלי ההחלטות הצעות מדיניות ממשלתית כוללת בהקשרי רגולציה של בינה מלאכותית. מוקד הידע יאפשר לממשלה לעקוב אחר ההתפתחויות בעולם ובמידת הצורך להמליץ על היבטים מסוימים בהם נדרשת רגולציה. חשיבותו של מוקד הידע התחדדה מאוד בעקבות הליך שמיעת הערות הציבור.

עתה, עם השלמת העבודה על מסמך המדיניות ופרסומו לציבור, בכוונתי לקדם החלטת ממשלה שתאפשר את אימוץ המדיניות על ידי ממשלת ישראל, תנחה את הרגולטורים לפעול בהתאם לה ותאפשר את יישומה של המדיניות האמורה. זאת, לשם חיזוק מעמדה של מדינת ישראל כמובילה טכנולוגית ובכדי לאפשר קידום מדיניות של חדשנות אחראית.

אני מודה מקרב לב למנכ"ל המשרד גדי אריאלי, לעו"ד מאיר לוי, המשנה ליועץ המשפטי לממשלה על העבודה המשותפת. תודות לרמ"ט מנכ"ל משרד החדשנות, המדע והטכנולוגיה הדסה גטשטיין, לד"ר זיו קציר מנהל התוכנית הלאומית לבינה מלאכותית של פורום תל"ם, ליועמ"ש משרד החדשנות המדע והטכנולוגיה עו"ד דני חורין ולעו"ד משה אוסטר ולגב' נועה אלקין, ממשרד המשפטים לעו"ד ד"ר יובל רויטמן ראש אשכול רגולציה, לעו"ד יוסף גדליהו מייעוץ וחקיקה משפט כלכלי, לעו"ד שרית פלבר ולעו"ד סדריק (יהודה) צבע מהמחלקה הבינלאומית במשרד המשפטים, ולכלל השותפים שסייעו בעבודה, בהתנצלות מראש אם נשמט שמו של מישהו שעסק במלאכה.

בכבוד רב,
אופיר אקוניס
שר החדשנות המדע והטכנולוגיה

14 דצמבר 2023
ב' בטבת תשפ"ד

הנדון: מסמך מקצועי – עקרונות מדיניות, רגולציה ואתיקה בתחום הבינה המלאכותית

בשנים האחרונות המונח "בינה מלאכותית" אינו שגור רק בקרב אנשי טכנולוגיה ומדע. השפעתה של טכנולוגיה זו על כל תחומי החיים הביאו לכך שמתקיים בה דיון ציבורי נרחב חוצה סקטורים: תעשייה, אקדמיה, ארגוני חברה אזרחית, רשויות ממשלתיות וכדומה. דיון זה הוא חשוב וחיוני בהתחשב בשינוי היסודי שטכנולוגיה זו מזמנת לנו. המסמך שמובא לפניכם הוא חלק מהמאמץ הממשלתי להביא בפני כלל הגורמים הנזכרים, כמו גם שותפינו בעולם, את עמדתה של ישראל ביחס לשאלות הסבוכות (והמסעירות) שמתעוררות אגב הפיתוח והשימוש בבינה מלאכותית.

כפי שתתרשמו, מדובר במסמך מקיף הדן בסוגיות החברתיות, המשפטיות והטכנולוגיות שהבינה המלאכותית מעוררת. המסמך בוחן סוגיות אלה מנקודת מבט מורכבת המבקשת לאזן בין יתרונות הטכנולוגיה וחשיבות עידוד החדשנות למשק ולציבור הישראלי לבין האתגרים המשמעותיים שהיא מעוררת. מתוך בחינה זו, המסמך מבקש להתוות ולהנחות את אופן הפעולה הממשלתי ביחס לטכנולוגיה במישור הרגולטורי, המשפטי והאתי. ההמלצות במסמך מתבססות על גישת ה"חדשנות האחראית", קרי, עידוד חדשנות תוך מתן דגש על השפעת הטכנולוגיה על האדם והחברה. כך, ההמלצה היא לאמץ, בשלב זה, רגולציה ענפית ולא רוחבית כדי שכל רגולטור יזהה את האתגרים בתחומו ויתן להם מענה מידתי ומתאים; להשתמש בכלי רגולציה גמישים והדרגתיים; לעודד שימוש בטכנולוגיה בסביבות ניסוי כדי ללמוד את השפעותיה ועוד.

עם פרסום המסמך, מדינת ישראל מצטרפת למדינות מפותחות וארגונים בינלאומיים שכבר פרסמו את "האני מאמין" שלהן ביחס לשאלת אופן אסדרתה של בינה מלאכותית. אנו רואים חשיבות גדולה בכך שישראל, על תעשיית החדשנות המשגשגת שבה, משדרת לעולם ולציבור כי היא גם נותנת דגש ומענה לאתגרי הטכנולוגיה במישור האתי והמשפטי.

במהלך העבודה על כתיבת המסמך ובעודי כותב שורות אלה – טכנולוגיית הבינה המלאכותית ממשיכה להתפתח ולהתקדם, ואגב כך ולהציב אתגרים חדשים. בעוד שהמסמך מציב תשתית איתנה לקראת הבאות, עלינו גם לחשוב על המשך העבודה בתחום ועל האופן שבו יש לעדכן את חשיבתנו המשפטית והרגולטורית בהתאם להתפתחויות הטכנולוגיות.

בהזדמנות זו, אבקש להודות לשר המשפטים על הרוח הגבית לפרויקט ונכונותו להמשיך ולעסוק בנושא בתחומי המשרד ומעבר לכך. אני מבקש להודות לעובדי ייעוץ וחקיקה ממחלקתי וממחלקות אחרות שנטלו חלק בכתיבת המסמך, העירו וטייבו את התוצר הסופי.

תודתי נתונה גם למשרד החדשנות, המדע והטכנולוגיה על שיתוף פעולה איתן ומתמשך בחשיבה, בכתיבה ובקידום המסמך ובציפייה להמשך עבודה משותפת בתחום חשוב זה.

מאיר לוין

המשנה ליועצת המשפטית לממשלה

תוכן עניינים

7.....	פתח דבר
10.....	תקציר מנהלים
19.....	1. חלק ראשון: רקע
19	1.1 מהי בינה מלאכותית?
20	1.2 חשיבות מדיניות רגולציה ואתיקה בתחום הבינה המלאכותית
23.....	2. חלק שני: המצב המשפטי והרגולטורי בהיבטי בינה מלאכותית בישראל
23	2.1 התפתחות המשפט והרגולציה להסדרת בינה מלאכותית
25	2.2 העיסוק הממשלתי בבינה מלאכותית
28	2.3 תחולה כללית של הדין והרגולציה בהיבטי בינה מלאכותית
28	2.3.1 דיני זכויות יוצרים
29	2.3.2 דיני פרטיות
30	2.3.2 אפליה אסורה
30	2.3.3 דיני חוזים
31	2.3.4 הגנת הצרכן
32	2.3.5 דיני נזיקין
33	2.3.6 הסדרה ייעודית מגזרית
33	2.3.7 משפט מנהלי
35.....	3. חלק שלישי: פעילות ארגונים בינלאומיים ומדינות מפותחות
35	3.1 המלצות ה-OECD בתחום הבינה המלאכותית
37	3.2 טיוטת החקיקה של האיחוד האירופי
39	3.3 מועצת אירופה וקידום אמנה בנושא בינה מלאכותית
40	3.4 מדיניות הרגולציה בארה"ב
43	3.5 מדיניות הרגולציה בבריטניה
45	3.6 מדיניות הרגולציה בקנדה
47.....	4. חלק רביעי: סוגיות ואתגרים המתעוררים בקשר לבינה המלאכותית
48	4.1 אפליה
49	4.1.1 סיכונים לאפליה אסורה במערכות מבוססות בינה מלאכותית
51	4.1.2 אתגרים בהתמודדות עם אפליה במערכות מבוססות בינה מלאכותית
53	4.2 מעורבות אנושית
54	4.2.1 מהי מעורבות אנושית?
55	4.2.2 יתרונות המעורבות האנושית
56	4.2.3 אתגרים ביחס למעורבות אנושית
59	4.3 הסברתיות
59	4.3.1 על הסברתיות ותופעת "הקופסא שחורה"
61	4.3.2 יתרונות וחסרונות הדרישה להסברתיות
63	4.4 גילוי
64	4.4.1 מה טיבה של דרישת הגילוי?
65	4.4.2 היתרונות והחסרונות להצבת דרישת גילוי
66	4.5 אמינות, עמידות, אבטחה ובטיחות
67	4.5.1 אמינות המערכת

69	עמידות המערכת ואבטחתה	4.5.2
71	בטיחות המערכת	4.5.3
72	אחריות	4.6
74	האתגרים בשימוש בהסדרים הקיימים להטלת אחריות	4.6.1
77	פרטיות	4.7
78	האתגרים לפרטיות בהם מתאפיין השימוש במערכות מבוססות בינה מלאכותית	4.7.1
81	5. חלק חמישי: דרכי התמודדות עם סוגיות ואתגרים אלה	
81	אפליה	5.1
82	מעורבות אנושית	5.2
84	הסברתיות	5.3
85	גילוי	5.4
86	אמינות, עמידות ובטיחות	5.5
88	אחריות	5.6
89	פרטיות	5.7
91	6. חלק שישי: המלצות למדיניות רגולציה ואתיקה לתחום הבינה המלאכותית	
91	אימוץ מדיניות רגולציה ממשלתית לתחום הבינה המלאכותית	6.1
91	החשיבות של מדיניות רגולציה ומעמדה	6.1.1
92	רכיבי מדיניות הרגולציה המוצעים לאימוץ	6.1.2
102	אימוץ עקרונות אתיים לתחום הבינה המלאכותית	6.2
102	החשיבות של אימוץ עקרונות אתיים ומעמדם	6.2.1
103	התבססות על עקרונות ה-OECD	6.2.2
103	העקרונות האתיים המוצעים לאימוץ	6.2.3
109	מיסוד מוקד ידע ותיאום ממשלתי להסדרת בינה מלאכותית	6.3
113	מיפוי השימושים בבינה מלאכותית והאתגרים הנלווים להם בענפים מוסדרים	6.4
114	הקמת פורום רגולטורים ופורום לשיתוף ציבור לתחום הבינה המלאכותית	6.5
114	תיאום המעורבות בפיתוח הרגולציה והתקינה בפורומים בינלאומיים	6.6
115	הנגשת מידע וכלים לשימוש אחראי בבינה מלאכותית, ובכלל זה כלי לניהול סיכונים	6.7
120	הערות הציבור	
122	תודות	

פתח דבר

ביום 1.8.2021 התקבלה בממשלה החלטה מס' 212 שעניינה "תכנית לקידום חדשנות, עידוד צמיחת ענף ההייטק וחיזוק המובילות הטכנולוגית והמדעית" (להלן: **החלטת ממשלה 212**),¹ במסגרתה, בין היתר, הוטל על שר החדשנות, המדע והטכנולוגיה להוביל את גיבוש מדיניות הממשלה בתחום הבינה המלאכותית, לרבות בהיבטי מדיניות רגולציה ואתיקה.

מסמך זה עוסק בעקרונות מדיניות, רגולציה ואתיקה בתחום הבינה המלאכותית, בהמשך להחלטת הממשלה 212. המסמך נערך במשותף על ידי גורמי המקצוע במשרד החדשנות, המדע והטכנולוגיה בראשות מנכ"ל המשרד מר גדי אריאלי, ובמחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים בראשות המשנה ליועצת המשפטית לממשלה (משפט כלכלי) עו"ד מאיר לוי.

ריכוז את העבודה על המסמך, בגרסתו הסופית, עו"ד משה אוסטר מהלשכה המשפטית של משרד החדשנות, המדע והטכנולוגיה, ועו"ד יוסף גדליהו ממחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים; בליווי עו"ד דני חורין, היועץ המשפטי למשרד החדשנות, המדע והטכנולוגיה ולמשרד התרבות והספורט, ועו"ד ד"ר יובל רויטמן, ראש אשכול רגולציה בייעוץ וחקיקה (משפט כלכלי). נטל חלק בכתיבת המסמך עו"ד עמית אשכנזי, לשעבר היועץ המשפטי למערך הסייבר הלאומי.

תודות מיוחדות למר גל קורן, לעו"ד מיכל בן סימון ולגב' נועה אלקין על עבודתם המסורה בכתיבת והגהת המסמך ולד"ר זיו קציר, מנהל התכנית הלאומית לבינה מלאכותית של פורום תל"מ, על הסיוע המקצועי. רשימת תודות מפורטת למי שהיו מעורבים בהכנת המסמך ובעבודה עליו מצורפת בסוף המסמך.

טיוטה של מסמך זה הופצה להערות הציבור ביום 30.10.2022, ובעקבותיה התקיים תהליך שימוע ציבורי נרחב. במסגרת זאת התקבלו תגובות בכתב מאת גורמים שונים מהמגזר הציבורי הרחב, ומאת נציגי אקדמיה, ארגוני חברה אזרחית ותעשייה. בנוסף, טיוטת המסמך עמדה במוקד יום עיון במתכונת "שולחנות עגולים" בארגון משרד החדשנות, המדע והטכנולוגיה וייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים, וכן במוקד יום עיון נוסף שארגן מרכז הנשיא מאיר שמגר למשפט דיגיטלי וחדשנות באוניברסיטת תל אביב בהשתתפות המעורבים בגיבוש המסמך. בהמשך אף נערכו פגישות עם מספר גופים ששלחו הערות בכתב וחפצו בשיח משלים בעל-פה.

בעקבות הדברים שנשמעו, ולאחר ליבון הערות הציבור, נבחנו בשנית מספר עניינים שהוארו במסגרת השיח, ונעשו תיקונים וחידודים בתוכן המסמך ובמדיניות המוצעת. ביחס לחלק אחר מההערות הומלץ על המשך מעקב ובינה בעתיד. כמו כן, חלק מההערות התייחסו לנושאים שמומלץ במסמך לעסוק בהם בעתיד כגון תוכן הרגולציה או הכלי לניהול סיכונים, ולפיכך אף שבהמשך הדרך יהיה בהחלט מקום להתחשב בהערות אלה, הן אינן בהכרח רלוונטיות למסמך זה.

מאז פרסום הטיוטה להערות הציבור חלו שינויים והתפתחויות דרמטיות בטכנולוגיות הבינה המלאכותית. המשמעותית ביותר בהם היא הבשלתם של "מודלי מסד" או "מודלי בסיס" (Foundation Models) ובאמצעותם הבשלתו של תחום הבינה המלאכותית "היוצרת" (Generative AI). חשיבותם של מודלי המסד היא ביכולתם לשרת מגוון גדול של משימות, וזאת מבלי שאומנו

¹ החלטה 212 של הממשלה ה-36 "תכנית לקידום חדשנות, עידוד צמיחת ענף ההייטק וחיזוק המובילות הטכנולוגית והמדעית" (01.08.2021) https://www.gov.il/he/departments/policies/dec212_2021

לכך באופן מפורש. כך מודלי מסד בתחום השפה מסוגלים לבצע שורה ארוכה של משימות ניתוח שפתי, תרגום, תמצות, הסקה, סיכום והשוואה של מקורות מידע שונים. התקדמות טכנולוגית מרשימה זו מהווה צעד משמעותי במעבר שבין "בינה מלאכותית צרה" (כזו אשר נועדה לשרת משימה נתונה), ובין "בינה מלאכותית רחבה" או "כוללת" (Artificial General Intelligence) שתהיה מקבילה ביכולותיה וגמישותה לחשיבה וההסקה האנושית. לצד זאת, מערכות "בינה מלאכותית יוצרת" מסוגלות ליצור תוכן מסוגים שונים, לרבות טקסט, וידיאו, ושמע, באיכות גבוהה המאתגרת את היכולת האנושית להבחין בין תוכן אמיתי לבין כזה אשר נוצר על ידי מכונה. שילוב טכנולוגיות אלו אפשר את יצירתם של ממשקי משתמש שיחתיים המסוגלים לענות על שאלות ולשרת בקשות של משתמש הקצה במגוון עצום של תחומי תוכן. מודלים פורצי דרך (כגון ChatGPT ו-Bard) השפיעו באופן ניכר על חיי היום יום והעלו את החשיבה על בינה מלאכותית למקום גבוה בסדר היום העולמי. התפתחויות אלה הביאו עמן אתגרים חדשים ושאלות מורכבות, ובכלל זה: כיצד להתמודד עם החשש משימוש בזיופים עמוקים (Deep fakes) של וידיאו וקול לביצוע הונאות למשל, או להפצת דיסאינפורמציה ומידע כוזב; כיצד להגיב לשימוש לרעה במודלי שפה, כגון ליצירה של מיילים הנחזים למהימנים במיוחד כחלק ממתקפת סייבר ("פשינג"), או לטובת יצירת תכנים שליליים ומסיתים; ומה המשמעויות של העובדה שחלק ממודלי הבסיס מוסרים מידע שגוי (illusions) או מפרשים מידע באופן שאינו נכון, לעיתים לגבי אנשים.

חלק ניכר מההשלכות של השינויים הטכנולוגיים וההתפתחויות הללו בתחום הבינה המלאכותית עדיין לא ברור, ויהיה צורך בהמשך חשיבה ועדכון המדיניות הממשלתית בהתאם למידע שיצטבר ולרגולציה שתתפתח בעולם בעקבותיהם. מעיון בחומרים מקצועיים ושותפות בצוותי חשיבה שונים בעולם, וכן משיח עם גורמים באקדמיה ובתעשייה, מצטיירת בשלב זה תמונה ראשונית שלפיה ישנה חפיפה משמעותית במרבית הסוגיות האתיות והרגולטוריות הנוגעות למודלי המסד ומערכות בינה מלאכותית פשוטות יותר, כך שחלק מן ההמלצות בדוח זה בוודאי יהיו רלבנטיות גם למודלים אלה. עם זאת, נראה כי מודלי המסד והמערכות הנבנות באמצעותם מציגים גם מספר אתגרים חדשים, או כאלו שהועצמו בצורה משמעותית – ראשית היכולת ליצור תוכן אשר נחזה כאמיתי בהיקפים גדולים ובמחיר נמוך, וכנגזרת מכך היכולת להשפיע על דעת הקהל, לשבש תהליכים דמוקרטיים, ליצור מאמצי הונאה בהיקף נרחב, או להשתלב במתן שירותים הכפופים כיום לאסדרה וכדומה. שנית היכולת לרתום את יכולתם של מודלי הבסיס למטרות בעלות היבטים של פגיעה ביטחונית ושלישית חידוד האתגר בהסדרת כלים טכנולוגיים אשר יכולים לשמש בלא שינוי למגוון גדול מאוד של מטרות.

בטווח המיידי, מוצע כי בחינת אופן הרגולציה הנדרש במודלים אלה, תיעשה במסגרת בחינת הרגולציה הדרושה ביחס ליישום ספציפי אשר כפוף לרגולציה, וזאת תוך שימוש בעקרונות המפורטים במסמך זה. כך למשל, כשייעשה שימוש במודל שפה לטובת שימושים של מסחר אלגוריתמי או מתן אשראי, ניתן יהיה להטיל על הגופים המיישמים, חובות רגולטוריות הנוגעות לשימושים אלו. חלק אחר מפעילות המודלים אינו מצריך בהכרח רגולציה בשלב זה, כגון שימוש ישיר במודלים הללו לשאלות ומחקר, כתיבת שירים וסיכומים ויצירת תמונות ושירים. לצד האמור, אין זה מן הנמנע שיהיה צורך לעסוק גם בבחינת הסדרה של המודלים הבסיסיים עצמם, וכן את הממשקים בינם לבין סוגיות של ביטחון, קניין רוחני, תחרות, אחריות ועוד, וזאת בדומה

להמלצות בדו"ח זה, הנוגעות להמשך מעקב אחר ההתפתחויות הטכנולוגיות בתחומי הבינה המלאכותית ואחר הרגולציה הנדרשת נוכח ההתפתחויות האמורות.

במקביל לשינויים בטכנולוגיה, חלו תמורות גם בתפיסה הציבורית וברגולציה בתחום הבינה המלאכותית. בהקשר הציבורי, בין היתר על רקע ההתפתחות הטכנולוגית המהירה שתוארה לעיל, פורסמו התרעות של מומחים בינלאומיים בתחום הבינה המלאכותית על הסכנות האורבות לאנושות בפיתוחים החדשים.² בהקשר הרגולטורי, הליך החקיקה באיחוד האירופי מוסיף להתקדם ומטעם הפרלמנט האירופי הופצה גרסה מעודכנת להצעת חוק הבינה המלאכותית, כן פורסם נוסח מעודכן של מדיניות הרגולציה הבריטית, בחודש אוקטובר 2023 פורסם צו נשיאותי אמריקאי המפרט את האסטרטגיה הלאומית האמריקאית בתחום הבינה המלאכותית וקובע הוראות שונות המנחות את משרדי הממשלה ואת הרגולטורים בארה"ב ועוד. בגרסה הנוכחית של המסמך נעשה מאמץ להתייחס לחלק מההתפתחויות בתחום זה. עם זאת, בשים לב לטיבן של ההתרעות שפורסמו וההתפתחויות התכופות, יהיה על מוקד הידע הממשלתי (שמוצע להלן להקימו) לעקוב באופן צמוד אחרי ההתפתחויות בהקשרים אלה ובמידת הצורך לנקוט בצעדים הנדרשים על מנת להתמודד עם האתגרים שהבינה המלאכותית מציבה, בתאימות למהלכים שיתבצעו במדינות מובילות בעולם.

יש להדגיש כי מסמך זה אינו עוסק בהסדרת השימושים של רשויות מינהליות (משרדי ממשלה, יחידות סמך, חברות ממשלתיות וכדומה) במערכות מבוססות בינה מלאכותית. אמנם חלק ניכר מהסוגיות, הסיכונים והאתגרים הנסקרים בו עשוי להיות רלוונטי גם ביחס לשימושים מסוג זה, אולם בכל הנוגע להיבטים המשפטיים של שימוש רשויות מינהליות וגופים דו-מהותיים בבינה מלאכותית מתקיימת עבודת מטה של הגורמים המשפטיים הרלוונטיים בממשלה, בהובלת מחלקת ייעוץ וחקיקה (משפט כלכלי) ומחלקת ייעוץ וחקיקה (משפט מנהלי) במשרד המשפטים, תוך עמידה על הדמיון והשוני בין שימוש במערכות מבוססות בינה מלאכותית על ידי רשויות מינהליות לבין השימוש בהן על ידי גופים פרטיים. לכן, ואף שבמקביל נעשית עבודה לעידוד הטמעה של מערכות מבוססות בינה מלאכותית בשימוש ממשלתי (בין היתר בהתאם להחלטת ממשלה 212, ולהחלטת ממשלה 173 מיום 24.2.2023 שעניינה "חיזוק המוביליות הטכנולוגית של מדינת ישראל" (להלן: **החלטת ממשלה 173**),³ הסדרת שימושים אלה והמעטפת המשפטית להם חורגת מנושא מסמך זה והיא מגובשת בשים לב לאמור בו, אך בנפרד ממנו.

² ראו למשל העצומה "Pause Giant AI Experiments: An Open Letter", 22.3.2023, Future of Life Institute, שנחתמה ע"י למעלה מ-30,000 איש ובהם מומחים בתחום הבינה המלאכותית ואנשי טכנולוגיה ידועים <https://futureoflife.org/open-letter/pause-giant-ai-experiments>.

³ החלטה 173 של הממשלה ה-37 "חיזוק המוביליות הטכנולוגית של מדינת ישראל" (24.02.2023) <https://www.gov.il/he/departments/policies/dec173-2023>.

תקציר מנהלים

בשנים האחרונות הפך המונח "בינה מלאכותית" (Artificial Intelligence - AI) למטבע לשון שגורה הבאה לתאר מכונות (מחשבים) הפועלות באופן שיכול להיתפש כחכם, מורכב במיוחד או תבוני. למערכות מבוססות בינה מלאכותית שימושים רבים, כגון נהיגת כלי רכב עצמאיים (אוטונומיים), פענוח תצלומי רנטגן, הערכת סיכוני אשראי, מסחר בניירות ערך, למידה מותאמת אישית ובחינת מועמדים לתעסוקה. כיום מערכות אלה נמצאות בשימוש מואץ במגזר הפרטי והציבורי בעולם ומחוללות מהפכה בענפים רבים במשק. בעתיד הקרוב, מערכות בינה מלאכותית צפויות להשפיע באופן משמעותי, ולעיתים אף מהפכני, על תהליכי עבודה, ייצור ואספקה ולהביא להתייעלות, חדשנות, ושיפור איכות החיים בתחומים רבים.

לצד היתרונות הרבים והפוטנציאל הכלכלי הגבוה,⁴ השימוש בבינה מלאכותית טומן בחובו אתגרים משמעותיים עבור המערכת המשפטית והרגולטורית ברחבי העולם ובישראל. הוא מעורר סיכונים וחששות שונים ומעלה סוגיות משפטיות ורגולטוריות חדשות, בכלל זה ניתן למנות עניינים אלה:

אפליה – השימוש במערכות מבוססות בינה מלאכותית עשוי למתן אפליה הנובעת מהטיות אנושיות וממבנים חברתיים, אולם במקביל לכך הוא מעורר חשש לאפליה בקנה מידה רחב, בין היתר כאשר ישנה הטיה בנתונים המשמשים לאימון המערכות (כגון שהאימון נעשה על בסיס מאגרי מידע לא ייצוגיים או כאלה המשקפים אפליה קיימת), או כאשר מובאים בחשבון משתנים המתואמים עם מאפיינים מפלים (כך למשל, מקום מגורים, הנלמד מכתובת או מיקוד, הוא נתון שלעיתים קרובות מקושר (proxy) לשיוך קבוצתי, למאפיינים אתניים או למעמד סוציו-אקונומי, ובהתאם התחשבות בו עלולה להביא לתוצאות מפלות). בנוסף, אתגרים משפטיים וטכנולוגיים שונים עלולים להקשות באופן ספציפי את ההתמודדות עם אפליה במערכות מבוססות בינה מלאכותית.

מעורבות אנושית – מאפיין מרכזי שמייחד מערכות מבוססות בינה מלאכותית הוא יכולתן לפעול ברמות שונות של אוטונומיות, מה שמעורר שאלות באשר לצורך במעורבות אנושית בפעילות המערכת (human in the loop), כגון בדרך של פיקוח שוטף על החלטותיה או קיום מנגנון ערעור בפני אדם על החלטותיה. מעורבות כזו עשויה לשפר את תהליך קבלת ההחלטות ולצמצם טעויות של המערכת; לחזק את מידת האחריותיות כלפי המערכת; להגביר את הלגיטימציה של ההחלטות ולהפחית את הפגיעה בכבודם של מושאי ההחלטות. מנגד, קיים חשש כי המעורבות האנושית לא תהיה אפקטיבית ואף תסב נזק, בין היתר עקב נטייה מוכרת של אנשים לאמץ המלצות של מערכות כאלה בלי להפעיל שיקול דעת משמעותי (Automation/Machine bias). על רקע זה, נדרש לבחון, בין היתר, מתי צריכה להיות דרישה למעורבות אנושית, וכיצד נכון לעצב את האינטראקציה בין הגורם האנושי לבין המערכת על מנת לנצל את היתרונות היחסיים של כל צד ולמנוע הטיות.

הסברתיות – "הסברתיות" היא היכולת להציג בצורה שניתנת להבנה על ידי בני אדם את אופן פעולת המערכת או את נימוקי ההחלטה שלה. בינה מלאכותית מבוססת על יכולות חישוביות היוצרות מודל תחזיות וקבלת החלטות, שלעיתים לא ניתן "לחלץ" מתוך המערכת (תופעה המתוארת כ-"קופסא שחורה" (black box)). בהיעדר יכולת להבין או להסביר את ההחלטה שהתקבלה על ידי המערכת, מלבד החשש שההחלטה תהיה שגויה, ישנו חשש לפגיעה בכבוד האדם

⁴ ראו למשל את הדו"ח של חברת הייעוץ הבינ"ל מקינזי ושות', The Economic Potential of Generative AI: The Next Productivity Frontier (14.6.2023) <https://www.mckinsey.com/capabilities/mckinsey-digital/our-insights/the-economic-potential-of-generative-ai-the-next-productivity-frontier#introduction>

ובאוטונומיה שלו, בכך שהוא נתון להחלטה שנחזית כשרירותית או עלולה להיות כזו. על רקע זה מתעוררות שאלות כגון באילו מצבים ראוי לדרוש כי יינתן הסבר ביחס לאופן פעילות המערכת, ומה צריכה להיות רמת הפירוט של ההסבר שיינתן בכל תרחיש רלוונטי. על אף יתרונותיה, בהקשרים שונים קיימת מורכבות טכנית או כלכלית בהחלת הדרישה להסברתיות באופן רחב.

גילוי – עם התקדמות הטכנולוגיה והשימוש ההולך והגובר במערכות מבוססות בינה מלאכותית לשם קבלת החלטות, במערכות צ'אט-בוט (chat-bot) לשם תקשורת עם משתמשים או במערכות יוצרות ליצירת תכנים ופרסומים, גוברת האפשרות שאנשים לא יהיו מודעים לכך שלמערכות אלה תפקיד משמעותי "מאחורי הקלעים" בקבלת ההחלטות בעניינם, שהם משוחחים איתן (ולא עם אנשים), או שהם נחשפים לתכנים שנוצרו על ידן. במקרים מסוימים, לגופים המפעילים מערכות אלה יהיה תמריץ להסתיר זאת, לעיתים תוך ניסיון להטעות את משתמש הקצה. נזכיר כי הסתרת העובדה שתוכן או תקשורת מסוימים נוצרים בידי בינה מלאכותית עשויה לעמוד ביסוד התופעה של ידיעות כזב (Fake news), ולהטעות את הציבור בדרכים מגוונות ורבות משמעות – לא רק בתחום צרכני ועסקי אלא גם בתחומים הנוגעים למשטר הדמוקרטי עצמו ולחופש הביטוי, הדעה והמחאה, כגון בתעמולת בחירות או בעיצוב תודעה פוליטית ביחס לשלטון או למחאה נגדו, אף שלא בתקופת בחירות. אשר על כן, הדיון בחובת הגילוי צריך להתחשב גם בהשפעות אפשריות אלה ובפגיעה האפשרית בזכויות חוקתיות ובעקרונות היסוד של המשטר הדמוקרטי.

על רקע זה, מתעוררת השאלה מה מידת הגילוי הנדרשת לגבי מעורבות מערכות מבוססות בינה מלאכותית.

אמינות, עמידות, אבטחה ובטיחות – מערכות מבוססות בינה מלאכותית חשופות לטעויות ותקלות טכניות, או למניפולציות מכוונות. קיים חשש שמערכת תסבול מביצועים ירודים משום שאינה מסוגלת לבצע כהלכה את המשימה שיועדה לה, וכן חשש שגורם חיצוני ישבש את פעילותה על ידי ניצול של נקודת תורפה הטבועה בה. לרוב מפתחי המערכות והמשתמשים בהן יהיו בעלי האינטרס המרכזי להתמודד עם חששות אלה, אך במקרים מסוימים עשויה להיות הצדקה לשקול התערבות רגולטורית על מנת לוודא כי מערכות מבוססות בינה מלאכותית הן אמינות, עמידות, מאובטחות ובטוחות במידה מספקת. זאת, בפרט בנסיבות בהן מתקיימות הצדקות מקובלות להתערבות רגולטורית, כדוגמת החצנות שליליות הכרוכות בסיכונים הנובעים מפעילותן של מערכות אלו.

אחריותיות – המאפיין האוטונומי של מערכות מבוססות בינה מלאכותית והקושי לצפות את פעולתן מראש ולהסביר אותה, עלולים במקרים מסוימים לערער את המבנה המקובל המושתת על הימצאות אדם שכלל נושא באחריות למעשיו במוקד ההתרחשות. ככל שהמעורבות האנושית בהחלטה או בפעולה מצטמצמת ורחוקה יותר, מתעוררות שאלות בנוגע לנשיאה באחריות המוסרית, החברתית והמשפטית (אזרחית או פלילית) בגין טעויות שנגרמו אגב השימוש במערכות מבוססות בינה מלאכותית. בכלל זה, עולות שאלות כגון מי נושא באחריות, מה סוג האחריות שניתן לייחס לו, ומה המשטר המשפטי המתאים לבחינת שאלת האחריות (למשל רשלנות, אחריות קפידה ומוחלטת, או נמל מבטחים). כמו כן, עולות שאלות בנוגע להגברת האחריותיות באמצעות קידום נורמות ארגוניות המבטאות נטילת אחריות, כגון ביצוע תהליך עיתי להערכת סיכונים או מינוי אחראי ארגוני להתמודדות עם סיכונים הקשורים לבינה מלאכותית. קשיים אלו מתעצמים לאור ריבוי הגורמים שעשויים להיות מעורבים בפיתוח ובשימוש במערכות בינה מלאכותית.

פרטיות והגנת מידע – ככלל פיתוח ושימוש במערכות מבוססות בינה מלאכותית מצריך שימוש בנתונים רבים, שחלקם עשויים לכלול מידע אישי, שהשימוש בו מוסדר בדיני הגנת הפרטיות. לפיכך, עשויים להתעורר אתגרים ייחודיים הנוגעים לצורך להבטיח את ההגנה הנדרשת לזכות לפרטיות. כך למשל, לעיתים מבוקש "לאמץ" מערכת במידע שנאסף למטרות אחרות, באופן שעלול לאתגר את העמידה בעיקרון צמידות המטרה, שהוא עקרון יסוד בדיני הגנת הפרטיות. כמו כן, קשיי ההסברתיות עלולים לאתגר את היכולת לקבל הסכמה מדעת מנושא המידע, או לעמוד בדרישת השקיפות ובחובת היידוע. לכך מצטרפים גם קשיים מעשיים לקבל הסכמה מנושאי המידע ביחס למידע שכבר נאסף בעבר. בנוסף, השימוש במערכות מבוססות בינה מלאכותית עלול ליצור מתח בין הצורך בנתוני עתק (Big Data), לבין מחיקת מידע עודף בהתאם לעיקרון צמצום המידע. איסוף ושמירה של מידע אישי עודף מגדילים את הסיכון לפגיעה בפרטיות, בין היתר בשל סיכוני אבטחת מידע. כמו כן, שמירה לאורך זמן של מידע מעלה את השאלה האם יש לאדם יכולת לשנות את המידע שנאסף עליו והאם יש להעניק לו זכות להתנגד להמשך עיבוד המידע לגביו. בנוסף, קיים חשש שייעשה שימוש במערכות אלה לשם זיהוי חוזר של מידע שעבר התממה (Anonymization). על כן, יש להידרש להיבטי פרטיות בפיתוח ושימוש מערכות מבוססות בינה מלאכותית, וגם במסגרת גיבוש מדיניות רגולטורית בתחום זה.

ההסדרים המשפטיים (בין היתר בדיני החוזים, הנזיקין, הגנת הצרכן והגנת הפרטיות) והרגולטוריים הקיימים בישראל (בענפים שונים) עשויים לחול על פיתוח ושימוש במערכות מבוססות בינה מלאכותית, ובהתאם להיות רלוונטיים לטיפול בחלק מהאתגרים והסוגיות המוזכרות לעיל ובעניינים נוספים. עם זאת, הניסיון מלמד כי בין התפתחויות טכנולוגיות באשר הן לבין המערכת המשפטית והרגולטורית נוצרות פעמים רבות נקודות חיכוך, וכפועל יוצא מכך, לא אחת עולה צורך בהתאמה של כללי המשפט והרגולציה הקיימים בעקבות הופעת טכנולוגיות חדשות המשבשות את מצב הדברים הקיים (disruptive innovation). על רקע זה יש להניח כי יהיה צורך בעיסוק מקבלי ההחלטות בכלל, וגורמי הממשלה והרגולטורים בפרט, באתגרים ובסוגיות שנסקרו לעיל ובעניינים נוספים הקשורים לפיתוח ושימוש במערכות מבוססות בינה מלאכותית.

מדיניות הרגולציה והאתיקה בתחום הבינה המלאכותית המוצעת במסמך זה נועדה להתוות את האופן שבו יבחנו מקבלי ההחלטות, בדגש על משרדי הממשלה והרגולטורים, את השאלה אם יש לקבוע כללי התנהגות לפיתוח ושימוש בבינה מלאכותית במגזר הפרטי בכלל או בתחומם, ואת האופן בו נכון לעשות כן. זאת, במטרה לקדם פיתוח ושימוש בטכנולוגיות בינה מלאכותית ואף להפחית חסמים משפטיים ורגולטוריים הניצבים בפניהן, ובד בבד לצמצם פגיעות אפשריות בזכויות יסוד ואינטרסים ציבוריים אגב פיתוח ושימוש בבינה מלאכותית.

המלצות למדיניות רגולציה ואתיקה בתחום הבינה המלאכותית

מספר עבודות קודמות עסקו בהסדרת תחום הבינה המלאכותית בישראל החל משנת 2018, אך עד עתה לא גובשה מדיניות ממשלתית אחידה לבחינת השאלות הרגולטוריות והאתיות בתחום, כאשר חלק מהעבודות אף המליצו על גיבוש מדיניות כזאת. מסמך זה כולל עקרונות מוצעים למדיניות אתיקה ורגולציה בתחום הבינה המלאכותית בישראל, תוך התחשבות בעקרונות מקובלים בעולם. לצורך גיבוש עקרונות אלה נבחנו המאפיינים הייחודיים של הבינה המלאכותית לאור החידושים הטכנולוגיים בעת הנוכחית, וכן האופן בו פועלים ארגונים בינלאומיים ומדינות מפותחות אשר עוסקים בתחום זה. המדיניות המוצעת מבקשת לאזן בין הצורך בוודאות ובהירות, האינטרסים

הציבוריים והזכויות שעל הפרק והחשש מהתערבות שאינה מוצדקת העלולה לפגוע בחדשנות טכנולוגית. כן כולל המסמך המלצות נוספות במטרה ליישם את המדיניות המוצעת.

להלן יפורטו בתמצית עיקרי ההמלצות המופיעות במסמך זה:

1. אימוץ מדיניות רגולציה ממשלתית לתחום הבינה המלאכותית

נקודת המוצא להצעה זו היא שישנה חשיבות רבה לכך שרגולציה המשפיעה על פיתוח ושימוש בבינה מלאכותית תקודם לפי מדיניות רגולציה ממשלתית אחידה וקוהרנטית או תותאם לה. זאת, על מנת להשיג את יעדי המדיניות של הממשלה; לקדם את תחום הבינה המלאכותית; להגן על זכויות יסוד ואינטרסים ציבוריים באופן רוחבי ואפקטיבי; ולצמצם חשש לפגיעה בחדשנות הטכנולוגית. בהקשר זה חשוב לציין שקביעת הסדרים רגולטוריים שיבטיחו שימוש אחראי ואתי בבינה מלאכותית הם לא רק הכרח אשר עומד במתח עם הרצון לעודד חדשנות; אלא גם אמצעי לעידוד צמיחה, פיתוח בר-קיימא, רווחה חברתית וקידום המובילות הישראלית בתחום החדשנות. מעורבות רגולטורית מבטיחה ודאות עבור התעשייה, מגבירה את האמון אצל המשתמשים, ויכולה למנוע או לצמצם פיתוח יישומי בינה מלאכותית שעלולים להזיק לציבור.

בהתאם לכך, מוצע לאמץ מדיניות רגולציה ממשלתית אחידה, ולכוון את גורמי הממשלה, ובפרט את הרגולטורים הרלוונטיים, בפעילותם לקביעת רגולציה ביחס לשימושים בבינה מלאכותית.

בפרט, מוצע לאמץ את רכיבי מדיניות הרגולציה המפורטים להלן, שגובשו בין היתר בשים לב ל"עקרונות מנחים לאסדרה מיטבית" הקבועים בחוק עקרונות האסדרה, התשפ"ב-2021; למסמכי מדיניות רגולציה שפרסמו המדינות המובילות בעולם בפיתוח רגולציה בתחום הבינה המלאכותית; להמלצות שניתנו ולעבודות קודמות שנעשו בישראל במסגרת הפעילות הציבורית בתחום; ולדיוני הצוות הבין-משרדי שעסק ברגולציה ואתיקה במהלך גיבוש התכנית הלאומית לבינה מלאכותית.

יצוין כי במסגרת גיבוש רכיבי מדיניות הרגולציה ניתן משקל של ממש לכך שמדובר בתחום חדשני ומורכב, שההתפתחויות הטכנולוגיות במסגרתו הן מהירות ועשויות להקדים במקרים רבים את גיבוש הרגולציה. עוד נלקחה בחשבון השונות הרבה שיש בין השימושים המגוונים בבינה מלאכותית, והעובדה שהמדובר בתחום שיש בו שימושים בעלי רגישות רבה או חשש לסיכון ולפגיעה בזכויות יסוד ואינטרסים ציבוריים, בצד שימושים שאינם מעוררים את אותה הרגישות או החשש. לבסוף נלקח בחשבון הצורך ללכת "עקב בצד אגודל" בתחום חדשני זה, מתוך רצון להוסיף ולעודד את המשק הישראלי כגורם מוביל מבחינה בינלאומית בפיתוח ושימוש בבינה מלאכותית, כמו גם מתוך הרצון לפעול באופן התואם לנעשה במדינות מובילות בתחום זה בעולם.

רכיבי מדיניות הרגולציה המוצעים עומדים כל אחד בפני עצמו, אך מתחברים יחד לכדי מדיניות שלמה. מטבע הדברים, לא כל הרכיבים ניתנים ליישום ביחס לכל רגולציה בתחום הבינה המלאכותית, וההתחשבות בהם ראויה להיות תמיד על פי ההקשר והנסיבות ובשים לב לשיטות העבודה המקובלות, אך מוצע כי ככלל יהיה צורך להתחשב בהם ולשאוף ליישם.

רכיבי מדיניות הרגולציה המוצעים הם:

ראשית, מוצע כי לעת הזו ככל שתקודם רגולציה שמטרתה להסדיר פיתוח ושימוש בבינה מלאכותית, הדבר ייעשה ברמת הענף שבו היא נדרשת, על בסיס צורך קונקרטי, ובהובלת הרגולטור האמון עליו. זאת, תוך שמירה על מדיניות ממשלתית אחידה באמצעות תיאום, אך לא באמצעות

רגולציה רוחבית (למשל, לא באמצעות חקיקה ייעודית המשתרעת על כל תחום הבינה המלאכותית). יחד עם זאת, נבקש לחדד כי אין באמור בכדי לגרוע מהצורך לשוב ולשקול מעת לעת את הצורך בקביעתה של רגולציה רוחבית בתחום, כמו גם להסדרת אתגרים רוחביים שמתעוררים בענפים שונים ומתאימים להסדרה כזו (בין היתר, בשים לב להתפתחויות בזירה הבינלאומית בהקשרים אלה).

שנית, מוצע כי רגולציה ביחס לפיתוח ושימוש בבינה מלאכותית, ככל שתאומץ, תביא בחשבון את הרגולציה הנוהגת במדינות מובילות בתחום ושאימצה על ידי ארגונים בינלאומיים, ותותאם במידת האפשר לרגולציה רלוונטית שנקבעה במדינות ובארגונים כאמור. זאת, בין היתר, כדי להימנע מהצבת חסמים רגולטוריים ייחודיים בפני כניסת מערכות מבוססות בינה מלאכותית לישראל.

שלישית, מוצע כי רגולציה המסדירה פיתוח ושימוש בבינה מלאכותית, תהיה מותאמת ככלל לסיכונים הנשקפים מסוג הטכנולוגיה ומהשימוש הספציפי שאותו היא נועדה להסדיר, ותהיה תוצר של ניהול סיכונים שביצע הרגולטור ביחס אליו. כך, שכלל, הרגולציה לא תחול באופן אחיד על טכנולוגיות ושימושים שרמת הסיכונים והחששות לגביהם שונה באופן ניכר.

רביעית, מוצע כי רגולציה שנועדה להסדיר פיתוח ושימוש בבינה מלאכותית תקודם ותתפתח בשלבים ובאופן הדרגתי בהתאם להתפתחויות הטכנולוגיות. במסגרת זו, אף מוצע כי ייעשה שימוש בנסיינות רגולטורית, ובכלל זה בפיילוטס וב"ארגוני חול" רגולטוריים, אשר יאפשרו הכנסה בטוחה של מערכות מבוססות בינה מלאכותית והפקה של התועלות החברתיות והכלכליות מכך.

חמישית, על מנת לעודד פיתוח ושימוש בבינה מלאכותית ולהפיק את התועלות החברתיות והכלכליות הנובעות מכך, מוצע כי במקרים המתאימים ייעשה שימוש ברגולציה מאפשרת. בכלל זה, יבחן שימוש בכלי רגולציה "רכים", גמישים ומתקדמים, כגון עקרונות אתיים לא מחייבים, תקינה, המלצות רגולטור לאימוץ וולונטרי ורגולציה עצמית (מפוקחת או בלתי מפוקחת).

שישית, מוצע כי רגולציה המסדירה פיתוח ושימוש בבינה מלאכותית תגובש תוך שיתוף בעלי הידע והמומחיות ובעלי העניין, ובכלל זה נציגי התעשייה (כולל "חברות הזנק"), האקדמיה, ארגוני חברה אזרחית והציבור בכללותו. זאת, במידה הנדרשת בנסיבות העניין, ועל מנת שתהיה מבוססת על תשתית מקצועית וטכנולוגית איכותית ותבטא איזון בין הזכויות והאינטרסים השונים.

2. אימוץ עקרונות אתיים לתחום הבינה המלאכותית

בהתאם להחלטת ממשלה 212, ובדומה למהלכים נוספים המקודמים במדינות ובארגונים שונים בעולם, מוצע לאמץ עקרונות אתיים משותפים עבור תחום הבינה המלאכותית בישראל, על בסיס עקרונות ה-OECD (Organization for Economic Co-operation and Development), הארגון לשיתוף פעולה ולפיתוח כלכלי, כדי לסייע לרגולטורים ולארגונים בתחום זה.

העקרונות מאפשרים שפה אתית משותפת ביחס לאופן ההתמודדות עם חששות וסיכונים הכרוכים במערכות מבוססות בינה מלאכותית, לרבות הסוגיות והאתגרים שצוינו לעיל. כן הם מאפשרים לגשר בין ערכים, יישומים טכנולוגיים, והיבטים משפטיים, תוך ממשק לתפיסה הבינלאומית המתהווה של "בינה מלאכותית מהימנה". אימוץ ופרסום העקרונות במסגרת המדיניות מאפשר גם שקיפות ושיח ציבורי לגבי העקרונות כנקודת מוצא למדיניות רגולציה לתחום הבינה המלאכותית. העקרונות יוכלו אף לסייע בקידום מדיניות רגולציה קוהרנטית בתחום הבינה המלאכותית, בכך

שהם מכונים את הרגולטורים להיבטים המרכזיים שיש להידרש אליהם, ומגבירים את הוודאות בקרב המגזר הציבורי והפרטי בעניין המדיניות הממשלתית לגבי בינה מלאכותית.

מוצע כי העקרונות לא יהיו בעלי מעמד משפטי, ובהתאם לכך הם לא מחייבים גורמים פרטיים לפעול או שלא לפעול בדרך כלשהי, לא מחליפים את המסגרת המשפטית החלה ולא מהווים כלי לפרשנות משפטית. עם זאת, עקרונות אלה משקפים היבטים בהם ראוי להתחשב במסגרת פיתוח ושימוש בבינה מלאכותית ובעת קביעת רגולציה בתחום זה.

העקרונות המוצעים מבוססים על עקרונות ה-OECD בהתאמות שונות, וזאת בשים לב לתפיסה הכוללת שלפיה יש לפעול ככלל בהתאם למדיניות המקובלת במדינות המפותחות בעולם.

העקרונות האתיים המוצעים הם:

א. **בינה מלאכותית לקידום צמיחה, פיתוח בר קיימא ומובילות ישראלית בחדשנות** – שימוש אחראי בבינה מלאכותית מהימנה הוא אמצעי לעידוד צמיחה, פיתוח בר קיימא, רווחה חברתית וקידום המובילות הישראלית בתחום החדשנות.

ב. **האדם במרכז: כיבוד זכויות יסוד ואינטרסים ציבוריים** – פיתוח ושימוש בבינה מלאכותית יש לעשות תוך כיבוד שלטון החוק, זכויות יסוד ואינטרסים ציבוריים ובפרט תוך שמירה על כבוד האדם ועל הזכות לפרטיות.

ג. **שוויון ומניעת אפליה פסולה** – בפיתוח ושימוש בבינה מלאכותית יש להביא בחשבון את הצורך בשוויון וגיוון, את הפוטנציאל הגלום בבינה מלאכותית לקידום שוויון מצד אחד, ואת החשש להטיה במערכות בינה מלאכותית והסיכון לאפליה פסולה כנגד יחידים או קבוצות מצד שני.

ד. **שקיפות והסבריות** – בפיתוח ושימוש בבינה מלאכותית יש להביא בחשבון, בשים לב למגוון השיקולים הנוגעים לעניין, את הצורך ליידע מי שבא במגע עם בינה מלאכותית או מושפע מפעילותה על כך, וכן את הצורך במתן הסבר להחלטתה או לאופן שבו פעלה. זאת, בין היתר בהתחשב במידת השפעתה, השלכותיה על מי שמושפע והאפשרויות הטכנולוגיות הזמינות.

ה. **אמינות, עמידות, אבטחה ובטיחות** – בפיתוח ושימוש בבינה מלאכותית יש להביא בחשבון את הצורך בכך שמערכות בינה מלאכותית תהיינה אמינות, מאובטחות ובטיחותיות לאורך כל "מחזור חייהן", כך שבתנאי שימוש רגילים, שימוש צפוי או שימוש שגוי או תנאים מסוכנים אחרים, הן יפעלו כראוי ולא יהוו סיכון בטיחותי בלתי סביר. לשם כך יש לנקוט אמצעים סבירים, בהתאם לתפיסות מקצועיות מקובלות של ניהול סיכונים, על מנת לצמצם סיכוני בטיחות ואבטחת מידע בכל "מחזור החיים" של מערכות בינה מלאכותית.

ו. **אחריותות** – מפתחי בינה מלאכותית, מפעיליה או המשתמשים בה יגלו אחריות לתפקודה התקין, ולקיום העקרונות האתיים האחרים בפעילותם, בין היתר בשים לב לתפיסות מקובלות של ניהול סיכונים ולאפשרויות הטכנולוגיות הזמינות.

3. מיסוד מוקד ידע ותיאום ממשלתי להסדרת בינה מלאכותית

בהמשך להמלצה בעניין זה בטיטת המסמך שפורסמה, בהחלטת ממשלה 173 נקבע כי על משרד החדשנות, המדע והטכנולוגיה להקים מוקד ידע ותיאום ממשלתי בתחומי הבינה המלאכותית אשר

יעסוק בנושאי רגולציה, מדיניות מידע ונתונים, אתיקה, שיתוף פעולה בינלאומי אזרחי והטמעה במגזר הציבורי האזרחי; וזאת בשיתוף מחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים.

מוצע לפעול בהקדם להקמת ומיסוד מוקד הידע האמור, לנוכח ההתפתחויות בתחום הבינה המלאכותית ופוטנציאל ההשפעה הרב של מערכות מבוססות בינה מלאכותית, ובשים לב לחסרונות הכרוכים בהחלטה שלא לקדם בשלב זה חקיקת מסגרת רגולטורית רוחבית (בפרט החשש מפערים, חפיפות וסתירות ברגולציה של הרגולטורים הענפיים שישדירו בינה מלאכותית בתחומם).

בהקשרים בהם עוסק מסמך זה – רגולציה ואתיקה – מוצע כי מוקד הידע ישמש כגוף פנים-ממשלתי מקצועי ובעל מומחיות, אשר ימלא, בין היתר, את התפקידים הבאים:

א. ייעוץ לרגולטורים הענפיים בבחינת הצורך ברגולציה בתחום הבינה המלאכותית בענף עליו הם מופקדים, ובאפיון וקידום רגולציה כזו ככל שהיא נחוצה;

ב. קידום תכנון, תיאום ושיתוף-פעולה בין רגולטורים בתחום הבינה המלאכותית, על מנת לצמצם חפיפות וסתירות בין רגולציה שנקבעת על ידי רגולטורים שונים, ולהגביר ודאות ואחידות;

ג. הובלת מהלכים מתואמים ורוחביים למימוש מדיניות הממשלה בתחום הבינה המלאכותית, ובפרט המדיניות המוצעת במסמך זה;

ד. ייעוץ לממשלה בענייני רגולציה ומדיניות רגולציה בתחום הבינה המלאכותית, ובכלל זאת בעניינים הנוגעים לעיצוב, עדכון ופיתוח המדיניות המוצעת במסמך זה, וכן מעקב אחר יישומה;

ה. הובלת הייצוג והמעורבות של ישראל בפיתוח רגולציה ותקינה בפורומים בינלאומיים;

ו. הנגשה ופיתוח של מידע וכלים לשימוש אחראי בבינה מלאכותית עבור הרגולטורים והציבור.

ז. הקמת פורומים רלוונטיים לקיום שיח ושיתוף ידע עם גורמי תעשייה, אקדמיה, מגזר שלישי וממשל, כמפורט בהמשך ההמלצות.

לנוכח תפקידיו המשמעותיים, מוקד הידע נדרש להיות מורכב מעובדי מדינה בעלי מומחיות, בין היתר, בתחומי מדיניות ממשלתית, רגולציה, קשרי חוץ וקשרי ציבור, טכנולוגיה ומשפטים. כמו כן, עליו להחזיק ביכולת המעשית, בהיבטים של תקנון ותקצוב, לתת מענה אפקטיבי לאתגרים השונים והמורכבים בתחומי הרגולציה של בינה מלאכותית. זאת, על מנת שיוכל להיות בעל ראייה רוחבית ולהחזיק באחריות לעידוד פיתוח ושימוש במערכות מבוססות בינה מלאכותית תוך שמירה על זכויות ואינטרסים, והשגת היעדים הממשלתיים. חשיבות מיוחדת יש להיבט זה נוכח ההתפתחות המהירה בתחומי הבינה המלאכותית, הן במישור הטכנולוגיה והשימושים והן בזירה הבין לאומית, ונוכח הצורך לתת מענה מהיר להתפתחויות האמורות, לצרכי השוק ולרגולטורים השונים.

על מנת לכוון את מוקד הידע בהתאם למכלול הזכויות והאינטרסים הנוגעים לעניין, להקנות לו ראיית רוחב על הנעשה בממשלה ובמשק הישראלי, לחזק את מומחיותו ויכולותיו, ולהעצים את השפעתו – מוצע להקים למוקד הידע ועדת היגוי, שבה יכהנו הנציגים הבאים, או מי מטעמם: מנכ"ל משרד החדשנות, המדע והטכנולוגיה (שיכהן כיו"ר הוועדה); מנהל מוקד הידע; המשנה ליועץ המשפטי לממשלה (משפט כלכלי) במשרד המשפטים; הממונה על התקציבים במשרד האוצר; ראש רשות הרגולציה; ראש רשות הגנת הפרטיות; ראש מערך הדיגיטל הלאומי; ומנכ"ל רשות החדשנות.

4. מיפוי השימושים בבניה מלאכותית והאתגרים הנלווים להם בענפים מוסדרים

בשנים האחרונות חל גידול משמעותי בהיקף השימושים במערכות מבוססות בינה מלאכותית והן מוטמעות כמעט בכל ענף במשק, כשההתפתחויות ביחס לבינה המלאכותית היוצרת מאיצות תהליך זה. בהתאם, ועל מנת שהחלטה על רגולציה נדרשת או נכונה תתקבל על בסיס מצע עובדתי מספק, מוצע כי גורמי הממשלה והרגולטורים הרלוונטיים יפעלו בהקדם לזיהוי ומיפוי של השימושים הקונקרטיים במערכות מבוססות בינה מלאכותית הנעשים על ידי השחקנים בענף המוסדר שעליו הם אמונים; האתגרים, החששות והסיכונים הכרוכים בכך; והמענים האפשריים.

מוקד הידע הממשלתי יסייע לגורמי הממשלה והרגולטורים במלאכה זו, ויקדם וירכז אותה.

5. הקמת פורום רגולטורים ופורום לשיתוף ציבור לתחום הבינה המלאכותית

מוצע להקים פורום מקצועי פנים-ממשלתי הכולל נציגי רגולטורים ומומחים לטכנולוגיה, מדיניות ומשפט, על מנת לקדם את הרגולציה הענפית בתחום הבינה המלאכותית באופן מתואם וקוהרנטי, ותוך שיתוף פעולה ולמידה הדדית. הפורום נועד לאפשר למידה משותפת ודיונים בסוגיות משותפות, בדגש על תחומים עתירי "חיכוך" בין רגולטורים שונים. זאת, במטרה לקדם את התיאום הממשלתי וכן לבחון במשותף את הצורך בעדכון מדיניות הרגולציה לתחום הבינה המלאכותית.

כן מוצע להקים פורום לשיתוף ציבור הכולל נציגי תעשייה, אקדמיה, וארגוני חברה אזרחית, וכן בעלי עניין חוץ-ממשלתיים רלוונטיים נוספים, במטרה לאתר את התחומים העיקריים שבהם יש צורך בהתאמה של הרגולציה או מדיניות הרגולציה בתחום הבינה המלאכותית, כדי להסיר חסמים לפעילות כלכלית או חברתית רציפה ויעילה או לגבש מענים מותאמים לצרכים.

6. מעורבות אקטיבית בפיתוח הרגולציה והתקינה בפורומים בינלאומיים

מוצע כי גורמי הממשלה הפועלים במסגרת פורומים וארגונים בינלאומיים, או עומדים בקשרי עבודה עם מדינות מפותחות אחרות בתחום הבינה המלאכותית, יפעלו יחד לקידום רגולציה ומדיניות רגולציה ואתיקה מאוזנת ותואמת למדיניות רגולציה זו וליעדים הממשלתיים.

זאת, בשים לב להחלטות הממשלה 212 ו-173, שהטילו על שר החדשנות, המדע והטכנולוגיה להוביל את שיתופי הפעולה הבינלאומיים האזרחיים בתחום הבינה המלאכותית, לייצג את ישראל בפורומים בינלאומיים אזרחיים, לרבות פורום השרים של ארגון ה-OECD, ולפעול לקידום שיתופי פעולה בינלאומיים אזרחיים בתחומי הבינה המלאכותית בשיתוף עם העוסקים בכך ברשות הלאומית לחדשנות טכנולוגית (להלן: **רשות החדשנות**), במשרד החוץ, במחלקת ייעוץ וחקיקה (משפט בינלאומי ומשפט כלכלי) ובגופים רלוונטיים נוספים. המלצה זו נובעת מהבנה כי לתקינה הבינלאומית שתגובש, כמו גם לתוצרי עבודה נוספים של מדינות וארגונים בינלאומיים העוסקים בסוגיות אלה, תהיה השפעה גלובלית משמעותית, כולל על הנעשה בישראל ועל יכולתו של המשק הישראלי להיות מוביל עולמי בתחום הבינה המלאכותית.

7. הנגשת מידע וכלים לשימוש אחראי בבניה מלאכותית, ובכלל זה כלי לניהול סיכונים

נוכח החשיבות של גיבוש מדיניות ממשלתית אחידה, בין היתר להגברת הוודאות הרגולטורית והמשפטית בתחום זה, מוצע להנגיש מידע וכלים לשימוש אחראי בבניה מלאכותית, ובכלל זה לשקול אימוץ של כלי אחיד לניהול סיכונים ביחס לפיתוח ושימוש בבניה מלאכותית שייצור שפה



משותפת בין גורמי הממשלה והרגולטורים ובינם לבין גורמים פרטיים. השפה האחידה תסייע למגזר הפרטי להעריך את הסיכונים הנלווים לשימוש מסוים בבינה מלאכותית, וכן לרגולטורים לבחון את הסיכונים הכרוכים בשימוש בה בתחום שעליו הם אמונים ואת הצורך בהתערבותם.

1. חלק ראשון: רקע

1.1. מהי בינה מלאכותית?

המושג "בינה מלאכותית" מתאר מכונות (מחשבים) הפועלות באופן שיכול להיתפש כחכם, מורכב במיוחד או תבוני. ריבוי הניסיונות ליצירת הגדרה פורמאלית עבור הבינה המלאכותית מרמז כי מדובר במשימה שאינה פשוטה כלל.⁵ לצורך הדיון במסמך זה, תחום הבינה המלאכותית הוא שם כולל להתפתחות בתחום טכנולוגיות המידע, התקשורת ומדעי הנתונים,⁶ המאפשרת קבלת החלטות, ייצור תחזיות, או ביצוע פעולות על ידי מחשב ברמת עצמאות גבוהה, באופן המדמה או מסוגל להחליף בינה אנושית.

למערכות בינה מלאכותית יתרונות בלתי מבוטלים. כך, בינה המלאכותית מסוגלת לקבל החלטות או לגבש המלצות באופן עצמאי ומהיר, ובכך להחליף שיקול דעת אנושי בשורה של מצבים. כמו כן, מערכות בינה מלאכותית מסוגלות לייצר תובנות ותחזיות על בסיס נתונים רבים באופן שפעמים רבות נחשב לבלתי אפשרי בידי אדם.

לבינה המלאכותית שימושים רבים, ובהם נהיגת כלי רכב עצמאיים (אוטונומיים), פענוח הדמיות רפואיות וצילומי רקמות, פיקוד על רובוטים הפועלים בשירות האדם, הערכת סיכוני אסרטי, זיהוי מזיקים בחקלאות, טיוב שרשראות אספקה, זיהוי מוקדם של תקלות בקווי ייצור, יצירת תכניות הוראה מותאמות אישית, מסחר בניירות ערך ובחירת מועמדים לתעסוקה. בהתאם, מערכות בינה מלאכותית מעורבות ביחסי עבודה; ביחסים בין צרכנים לעוסקים; בין עסקים לעסקים אחרים; בין בעלי מקצוע ללקוחות; בין גופים במגזר הציבורי; ובין המגזר הציבורי לכלל הציבור.

במקביל לקושי בהגדרת תחום הבינה המלאכותית בהיבט המקצועי והטכנולוגי, הגדרת "בינה מלאכותית" לצרכי גולציה, באופן שמצד אחד יכלול את כלל מערכות המחשוב עליהן רצוי שהרגולציה תחול, ומצד שני יחריג מערכות פשוטות ויומיומיות שאינן מעוררות כל סיכון, היא משימה מאתגרת.

בתוך שלל הניסיונות בעולם ליצירת הגדרה הולמת, בעת האחרונה מסתמנת מגמה של אימוץ ההגדרה שמציע ה-OECD, המתמקדת בהיבט הפונקציונלי של מערכת העושה שימוש בבינה מלאכותית. הגדרה זו, אשר גובשה על ידי ועדת מומחים בשנת 2018,⁷ נבחנת מחדש כיום בעקבות התפתחויות קיימות וצפויות בתחום הבינה המלאכותית, וייתכן שתעודכן בעתיד הקרוב. מכל מקום, על פי ההגדרה הנוכחית, מערכת בינה מלאכותית היא:

⁵ את הניסיונות השונים להגדרת הבינה המלאכותית ניתן לחלק לשתי קבוצות עיקריות – הגדרות טכנולוגיות והגדרות פונקציונליות. ההגדרות הטכנולוגיות מתרכזות על פי רוב ביכולות הבסיסיות של כלי הבינה המלאכותית, ובשיטות ליישומן, וניתן לחלקן לשתי שיטות מרכזיות: האחת, מזוהה עם תחום מדעי הנתונים (data science) ומתארת את הבינה המלאכותית בתור כלל הפעולות הטכנולוגיות למיצוי מידע והפקת תובנות מתוך מאגרי נתונים. השנייה, מזוהה עם תחום למידת מכונה (machine learning) ולפיה בינה מלאכותית היא היכולת של מכונה "להתאמן" על ביצוע פעולה ולטייב את ביצועה, בהסתמך על נתונים, דוגמאות וניסיון מצטבר. שתי השיטות אינן מתייחסות למושגים כגון "תבונות" או "חשיבה", אלא למהות הטכנולוגית של איסוף ועיבוד נתונים במטרה להגשים משימה אלגוריתמית נתונה. לעומת זאת, הגדרות פונקציונליות מתייחסות או מגדירות בינה מלאכותית באמצעות השוואתה לתהליכי חשיבה או הסקה אנושיים. הדגש בסוג זה של הגדרה הוא על יכולתן של מכונות לפעול באופן אשר יכול להיתפש כתבוני, או אשר מחקה את דפוסי הפעולה האנושיים.

⁶ לעניין הקשר בין בינה מלאכותית ומדעי הנתונים ראו את מסקנות דוח ועדת ההיגוי המייעצת לות"ת לנושא מדעי הנתונים: המועצה להשכלה גבוהה, הוועדה לתכנון ולתקצוב, דוח ועדת ההיגוי המייעצת לות"ת לנושא מדעי הנתונים, דצמבר 2020, זמין כאן.

⁷ OECD, ARTIFICIAL INTELLIGENCE IN SOCIETY, Chapter 1 (2019) (להלן: AI in Society). <https://www.oecd-ilibrary.org/>

"מערכת ממוכנת שיכולה, בהתאם למטרות מוגדרות בידי אדם, להפיק תחזיות, המלצות, או החלטות המשפיעות על סביבות פיזיות או וירטואליות, ומתוכננת לפעול ברמת עצמאות משתנה".⁸

הגדרה דומה מאוד נכללה בטיוטה העדכנית של חוק הבינה המלאכותית האירופי,⁹ ובטיוטת החקיקה הקנדית.¹⁰

מערכות בינה מלאכותית רבות מבוססות על למידת מכונה, כלומר תהליכים שבהם המערכת "לומדת" מהמידע שהיא אוספת או מקבלת ומפיקה תובנות הנדרשות לפעולתה. "לימוד" המכונה נעשה בין היתר דרך הצגה של דוגמאות רבות לתוצאות המבוקשות, ניסוי וטעייה.¹¹ למידת מכונה מאפשרת יצירת "מודל" המייצג תובנות מהמידע, וקבלת החלטות או המלצות על פיו.¹²

בינה מלאכותית מבוססת למידת מכונה נשענת בהקשרים רבים על היכולת של המכונה לטייב את התחזיות שהיא מספקת באמצעות זיהוי דפוסים סטטיסטיים חבויים המתקיימים בתוך נתוני האימון – אופן לימוד שכלל אינו דומה לדרך בה בני אדם לומדים ומסיקים מסקנות. משמעות גישה זו של אימון מודל על בסיס נתונים, היא כי פעילות מערכת התוכנה מוגדרת במידה רבה על ידי תהליך האימון והנתונים אשר שימשו במהלכו, ולא על ידי המתכנת. משמעותה המעשית של עובדה זו היא כי ייתכן ויהיה קושי לשחזר או לעקוב אחרי חוקי הפעולה של מערכת המחשב, ולעיתים ההתחקות אחר הליך קבלת ההחלטות של המערכת כלל אינה אפשרית – תופעה אשר מכונה לעיתים "קופסה שחורה" (black box).

1.2. חשיבות מדיניות רגולציה ואתיקה בתחום הבינה המלאכותית

מדיניות הרגולציה והאתיקה בתחום הבינה המלאכותית, המוצעת במסמך זה, נועדה להתוות את האופן שבו תבחן הממשלה את השאלה האם יש לקבוע כללי התנהגות לפיתוח בינה מלאכותית או

⁸ RECOMMENDATION OF THE COUNCIL ON ARTIFICIAL INTELLIGENCE, OECD/LEGAL/0449 (Adopted on 8 May 22, 2019), OECD LEGAL INSTRUMENTS <https://legalinstruments.oecd.org/en/instruments/oecd-legal-0449> (להלן: **המלצות ה-OECD**). בשפת המקור:

"AI system: An AI system is a machine-based system that can, for a given set of human-defined objectives, make predictions, recommendations, or decisions influencing real or virtual environments. AI systems are designed to operate with varying levels of autonomy".

⁹ ההגדרה בגרסת הפרלמנט האירופי לטיוטת החקיקה: "artificial intelligence system" (AI system) means a machine-based system that is designed to operate with varying levels of autonomy and that can, for explicit or implicit objectives, generate outputs such as predictions, recommendations, or decisions, that influence physical or virtual environments" Artificial Intelligence Act, Amendments adopted by the European Parliament on 14 June 2023 on the proposal for a regulation of the European Parliament and of the Council on laying down harmonized rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts (COM(2021)0206 – C9-0146/2021 – 2021/0106(COD)) https://www.europarl.europa.eu/doceo/document/TA-9-2023-0236_EN.pdf

¹⁰ ההגדרה בטיוטת החקיקה הקנדית: "artificial intelligence system" means a technological system that, autonomously or partly autonomously, processes data related to human activities through the use of a genetic algorithm, a neural network, machine learning or another technique in order to generate content or make decisions, recommendations or predictions. (*système d'intelligence artificielle*) <https://www.parl.ca/DocumentViewer/en/44-1/1000/Document/1000>; HOUSE OF COMMONS OF CANADA, BILL C-27, An Act to enact the Consumer Privacy Protection Act, the Personal Information and Data Protection Tribunal Act and the Artificial Intelligence and Data Act and to make consequential and related amendments to other Acts, Part 3, Article 2.

¹¹ אופן הלימוד מושפע בין היתר מהיבטים אלה:
א. למידה מפוקחת או לא מפוקחת.
ב. למידה אקטיבית מול למידה פסיבית.
ג. מידת ה-"סיוע" המתקבל מהגורם המלמד.
ד. קיום שלב מקדים של למידה לפני הפעלת המערכת בתנאי אמת, לבין מערכת שנפרסת ללא שלב מקדים של למידה.
¹² החוקרים משתמשים שונות תחת תחום למידת המכונה. שיטות אלה כוללות רגרסיה לינארית ולוגית, עצי החלטה, principle component analysis וגם רשתות נוירונים עמוקות (deep neural networks).

שימוש בה במגזר הפרטי, וכן את האופן בו ייקבעו כללים כאמור במקרים המתאימים. זאת מתוך הנחה שמדיניות רגולציה ואתיקה נכונה יכולה לעודד רגולציה מיטיבה, אשר בכוחה לייצר וודאות, להסיר חסמים ולפתור כשלי שוק שעשויים להקשות על מיצוי הפוטנציאל של בינה מלאכותית; ובד בבד ליצור כללים וסטנדרטים בתחום הבינה המלאכותית באופן שימקסם את התועלות הנובעות ממערכות בינה מלאכותית ויצמצם פגיעות אפשריות שלהן בזכויות יסוד ובאינטרסים ציבוריים.

הצורך בגיבוש ואימוץ מדיניות רגולציה ואתיקה לתחום הבינה המלאכותית חשוב במיוחד ממספר טעמים, כמפורט להלן:

ראשית, מאפייני הבינה המלאכותית עשויים לעורר שאלות בעניין התערבות רגולטורית, ומדיניות רגולציה ואתיקה תסייע בבחינת הצורך בהתערבות רגולטורית ואופי ההתערבות.

בהתאם לתפיסות מקובלות של רגולציה, יש לבחון התערבות באופן הפעילות של המגזר הפרטי באמצעות כללים רגולטוריים (דהיינו, כללים מחייבים וניתנים לאכיפה שנועדו להסדיר פעילות כלכלית או חברתית המבוצעת על ידי גורמים פרטיים),¹³ בעיקר במקום שבו קיימים כשלי שוק או כאשר יש חשש לפגיעה בזכויות יסוד או באינטרסים ציבוריים, שאינם מטופלים כראוי על ידי מנגנוני השוק.¹⁴ גם במקרים אלה, רגולציה תיעשה רק כאשר העלות הכרוכה בהתערבות הרגולטורית צפויה להיות נמוכה מזו הנובעת מכשל השוק או מהפגיעה באינטרס המוגן, ולאחר בחינת חלופות להתערבות הרגולטורית. קביעת רגולציה מהווה אפוא מלאכת איזונים עדינה שנועדה לנווט את פעילות השוק מבלי לייצר נטל רגולטורי עודף.

הצורך בהחירות במלאכת איזונים זו מועצם פי כמה כאשר מושא הרגולציה הוא טכנולוגיה חדשנית, כיוון שבמקרים אלו קשה להעריך מראש הן את עצמת הפגיעה באינטרס המוגן והן את העלות הרגולטורית. יתר על כן, קיים חשש לפיו הנטל הרגולטורי כשלעצמו יהיה גרסיבי, שכן הוא יפגע בתמריץ לייצר טכנולוגיות חדשות או ינווט את פיתוחן למקומות "בטוחים מבחינה רגולטורית" אך פחות רצויים. אתגר זה מלווה את המדינות השונות והרגולטורים השונים הפועלים בהן, בעת שהם בוחנים שאלות הנוגעות לקביעת רגולציה בתחום הבינה המלאכותית.

על רקע זה, המדיניות המוצעת יכולה לסייע בבחינת הצורך בהתערבות רגולטורית ובבחינת אופייה. זאת, בין היתר בשים לב לאיזון הנדרש בין הרצון לקדם חדשנות טכנולוגית לבין הצורך להגן על זכויות יסוד ואינטרסים ציבוריים, ולפי תפיסת ניהול הסיכונים הבוחנת את היחס בין עוצמת הסיכון והסיכון להתרחשותו, לבין סוג האמצעים הנדרשים לניהולו.

שנית, כללים רגולטוריים עלולים לשמש כחסם בפני פיתוח או שימוש בבינה מלאכותית, ומדיניות רגולציה ואתיקה תועיל בהכוונה לצמצום מאוזן ונכון של חסמים אלה.

רגולציה מחילה מגבלות או איסורים על פעילות של גורמים פרטיים, ובהתאם לכך לעיתים היא עלולה למנוע או להגביל פיתוח בינה מלאכותית או שימוש בה. כך למשל, רגולציה עשויה לאפשר פעילות מסוימת על ידי אדם בלבד (כגון נהיגה עם "ידיים על ההגה" או זיהוי "פנים-אל-פנים" בפתיחת חשבון בנק) או על ידי איש מקצוע בלבד (כגון אבחון רפואי או ייעוץ משפטי), ובכך לחסום או להגביל מערכת מבוססת בינה מלאכותית המסוגלת לבצע פעילות דומה.

¹³ הגדרת "אסדרה" בס' 3 לחוק עקרונות האסדרה, התשפ"ב-2021.
¹⁴ ס' 1 לחוק עקרונות האסדרה.

בהתאם לכך, המדיניות המוצעת עשויה להוביל ולתמוך בקידום מהלכים להפחתת חסמים, הנובעים מרגולציה קיימת או חדשה שאינה מותאמת לטכנולוגיה מבוססת בינה מלאכותית, וכפועל יוצא מונעת או מגבילה פיתוח ושימוש בבינה מלאכותית שלא לצורך. זאת, בין היתר, באמצעות עיצוב נכון של הכללים הרגולטוריים או על ידי שימוש בכלי רגולציה מתקדמים כגון הוראות שעה, סעיפי "שקיעה" (sunset clauses) ו"ארגזי חול" רגולטוריים (regulatory sandbox).

שלישית, מדיניות רגולציה ואתיקה צפויה להגביר את הוודאות העסקית, ולחזק את אמון הציבור בטכנולוגיות בינה מלאכותית, ובכך לעודד פיתוח ושימוש בבינה מלאכותית.

ככל טכנולוגיה חדשה, בחלק מהשימושים בבינה מלאכותית צפויה אי-ודאות לגבי אופן ההשפעה שלהם על זכויות יסוד ואינטרסים ציבוריים, כמו גם על בעלי עניין בהקשרים שונים. לחששות אלה עלולה להיות השפעה כפולה: האחת, שגורמים עסקיים הפועלים לפיתוח והפצה של מערכות מבוססות בינה מלאכותית יימנעו מכך בשל החשש מרגולציה עתידית שתהיה חוסמת או מגבילה; השנייה, שהציבור הרחב שמיועד להשתמש במערכות אלה או להיות מושפע מהן, יירתע מכך.

לפיכך, מעבר להגנה על זכויות יסוד ואינטרסים ציבוריים, למדיניות המוצעת חשיבות גם בהקניית ודאות לגבי אופן בחינת קידום כללים והסדרים שיחולו על בינה מלאכותית, במידת הצורך, וזאת, בהלימה לנעשה במדינות המפותחות שהם שווקי היעד של הטכנולוגיה. לקיומה של מדיניות בהירה, המעידה על בחינת האינטרסים השונים בידי הגורמים הציבוריים המופקדים עליהם תפקיד חשוב בקידום הוודאות לעוסקים בתחום ולהגברת אמון הציבור בשימושים חדשניים בטכנולוגיה.

כפי שיפורט להלן, המדיניות במסמך זה מוצעת לנקודת הזמן הנוכחית. בהתאם, יש צורך להוסיף ולבחון אותה מעת לעת בהתאם להתפתחויות הטכנולוגיות, החברתיות, המשפטיות והרגולטוריות בעולם. כך, ביתר שאת, בשים לב להתפתחויות העת האחרונה, בפרט בכל הנוגע לבינה מלאכותית יוצרת (Generative Artificial Intelligence) והשיח אודותיה, כמפורט לעיל.

2. חלק שני: המצב המשפטי והרגולטורי בהיבטי בינה מלאכותית בישראל

2.1. התפתחות המשפט והרגולציה להסדרת בינה מלאכותית

לצד היתרונות למשק ולחברה הנובעים מההתפתחות הטכנולוגית בכלל, ומהתפתחות טכנולוגיות הבינה המלאכותית בפרט, התפתחויות אלה מייצרות גם אתגרים משמעותיים עבור המערכת המשפטית והרגולטורית, ברחבי העולם ובישראל.

הניסיון מלמד כי בין התפתחויות טכנולוגיות באשר הן, לבין המערכת המשפטית והרגולטורית נוצרות פעמים רבות נקודות חיכוך. זאת, בעיקר משום שלא אחת הטכנולוגיות החדשות משנות את מצב הדברים שהיה קיים טרם הופעתן או את אופן התנהגותם של יחידים והחברה, ואגב כך מעוררות סוגיות חברתיות, משפטיות ורגולטוריות (יש המתמייחסים למצב זה כ"חדשנות משבשת" (disruptive innovation) ו"שיבוש רגולטורי" (regulatory disruption)).¹⁵

כתוצאה מכך, לא אחת עולה צורך בהתאמה של כללי המשפט והרגולציה הקיימים עקב הופעת טכנולוגיות חדשות, ובעשורים האחרונים מקבלי ההחלטות עוסקים במלאכה זו.¹⁶ ההתאמה נעשית, ככלל, באמצעות פרשנות, תיקון הכללים, או קביעת כללים חדשים. מלאכה זו אינה קלה, וכרוכה בהתמודדות עם קשיים מהותיים, כגון הקצב המהיר של ההתפתחות הטכנולוגית, הצורך במידע ובמומחיות לשם הסדרתה, והמאפיינים הבינלאומיים הנלווים לה. עמד על הקושי הכרוך במלאכה זו, בין היתר, השופט סולברג בפרשת איגוד האינטרנט הישראלי:

"בידוע, כי המשפט מדדה בעצלתיים אחר חידושי העולם, וכי החקיקה אינה מדביקה את קצב התקדמות המדע וישומיו. מפרי-החוק מסתגלים לקידמה מהר יותר מאוכפיו. זו אקסיומה. לראשונים אין עכבות; לאחרונים יש. שנים רבות חלפו מעת המצאת המחשב ועד לחקיקת חוק המחשבים, התשנ"ה-1995. לא חלפו אז דור או שניים במונחי מחשב, והחוק כבר נמצא מיושן, כי המחוקק לא צפה – ולא יכל לצפות – את חידושי הטכנולוגיה. אך לא רק עולם המשפט עומד נבוך. גם מדע הפסיכולוגיה נתקל בתופעות חדשות של התמכרות ופגיעות בנפש, ומנסה ללמוד דרכי התמודדות עדכניות 'תוך כדי תנועה'; כך גם הסוציולוגיה, ושאר ענפים מתחום מדעי החברה, הטבע והרוח. אין תימה אפוא שגם עולם המשפט איננו ערוך עדיין עם מלוא הכלים העומדים לרשותו.

¹⁵ Nathan Cortez, *Regulating Disruptive Innovation*, 29 BERKELEY TECH. L. 175 (2014).

¹⁶ כך לדוגמה, בישראל הוקמו בשנים האחרונות מספר ועדות שעסקו בהתאמת המשפט והרגולציה להתפתחויות טכנולוגיות, עליהן נמנות, בין היתר: ועדת המשנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה של בינה מלאכותית (2019); צוותי העבודה הבין-משרדיים לגיבוש התכנית הלאומית לבינה מלאכותית, בהובלת משרד החדשנות, הטכנולוגיה והמדע בהתאם להחלטת ממשלה 212 (2021); הוועדה להתאמת המשפט לאתגרי החדשנות והאצת הטכנולוגיה (2021); הוועדה לבחינת אסדרת פעילות הרשתות החברתיות (2021); הוועדה לגיבוש אמצעים להגנה על הציבור ונושאי משרה בשירות הציבור מפני פעילות ופרסומים פוגעניים, כמו גם בריונות ברשת האינטרנט (2020); הוועדה הציבורית לבחינת חוק הבחירות (דרכי תעמולה), התשי"ט-1959; הוועדה לבחינת אסדרה של הנפקת מטבעות קריפטוגרפיים מבוזרים לציבור (2019); והצוות המקצועי לבחינת "כלכלת הפלטפורמה" – צורות העסקה ייחודיות שנוצרו במשק, בהובלת זרוע העבודה (2022).

"[...] אין לכחד: הסדרה חוקית של הנעשה במרחב הווירטואלי היא מורכבת ומסובכת. קשה להגדיר את הדין הרצוי, וגם לא בנקל ניתן להחיל את הדין המצוי. לא בכדי יש שהגיעו למסקנה כי מדובר בתחום שראוי להסדרה חוקית; יש הסבורים כי הפסיקה היא המסגרת הראויה לעשיית ההתאמות הנדרשות לעידן האינטרנט; ואלה ואלה מתלבטים גם לגבי מידת השיתוף של קהילת משתמשי האינטרנט בעיצוב הכללים במרחב הווירטואלי וביישומם".¹⁷

הנחיית היועץ המשפטי לממשלה מס' 1.2500/18 עוסקת באופן קידום ופיתוח המשפט במסגרת היחסים שבין טכנולוגיה, ערכים ומשפט. ההנחיה עוסקת ב"ביצוע המעבר מן העולם הפיסי אל העולם הדיגיטלי, תוך ניהול סיכונים וקיום התכליות הנדרשות", הן בקשר "לשירותים ממשלתיים וציבוריים" והן במסגרת "הסדרים החלים על השוק הפרטי", המוגדרים כ"הסדרים דיגיטליים". ההנחיה מתייחסת לגישה הפרשנית בעת עיצוב הסדר דיגיטלי, הכוללת בין היתר בדיקה לאיתור התכליות של ההסדר המשפטי שהופך לדיגיטלי, הסיכונים במעבר להליך דיגיטלי, ואופן ההתמודדות איתם באמצעים טכניים, נהלים ומשפטיים. על פניו, נראה כי לגישה זו רלוונטיות גם להחלת הדין הקיים על פיתוח ושימוש במערכות מבוססות בינה מלאכותית.

טכנולוגיות הבינה המלאכותית מצויות כבר כיום בשימוש בידי גופים פרטיים וציבוריים בישראל ובעולם. בדומה לאופן שבו התפתחו המשפט והרגולציה ביחס לטכנולוגיות אחרות, ניתן להצביע על מאפיינים מרכזיים הרלוונטיים גם לאופן קידום ופיתוח המשפט והרגולציה ביחס לטכנולוגיות בינה מלאכותית. המאפיין העיקרי והחשוב ביותר הוא שמערכת הדינים החלה על כל פעילות באופן כללי, הן בתחום המשפט הפרטי והן בתחום המשפט הציבורי, חלה גם על בינה מלאכותית.

כפי שיתואר להלן, השימוש במערכות מבוססות בינה מלאכותית מעורר סוגיות ואתגרים שונים, ובכלל זה אפליה; שקיפות; צורך במתן הסבר; אבטחת מידע והגנת סייבר; ופרטיות. ניתן להניח, כי חלק מסוגיות ואתגרים אלו יצריכו חשיבה מחודשת על הדין הקיים ואולי אף קביעת כללים חדשים במקרים המתאימים. מנגד, נראה שחלק מסוים מהסוגיות והאתגרים האלה צפויים לקבל מענה במסגרת הדין הקיים, בין אם במישרין ובין אם באמצעות פרשנות מתאימה.¹⁹

¹⁷ ע"מ 3782/12 **מפקד מחוז תל אביב-יפו במשטרת ישראל נ' איגוד האינטרנט הישראלי**, סו(2) 159, פס' 23-25 לפסק הדין של השופט סולברג (2013).

¹⁸ הנחיית היועץ המשפטי לממשלה 1.2500 "כללים מנחים לגיבוש הסדרים דיגיטליים" (10.10.2019) (להלן: **הנחיות היועץ המשפטי לממשלה 1.2500**).

¹⁹ הסקירה בדבר אופן הפיתוח המשפטי של הדין באופן כללי לגבי בינה מלאכותית, ובתחומי משפט קונקרטיים מחייבת בחינה מפורטת נוספת. באופן כללי נראה כי הגישה העולה מהנחיית היועץ המשפטי לממשלה 1.2500 רלבנטית במובן זה שבחלק מהשימושים החברתיים של בינה מלאכותית ניתן יהיה להתאים את ההסדר הנורמטיבי באמצעות פרשנות הדין הקיים בכלים הקיימים, אולם לעיתים שיקולים הקשורים בשינוי חלוקת אחריות או שינוי נטלי ראיה יובילו לצורך בהסדר ספציפי. ראו הנחיית היועץ המשפטי לממשלה 1.2500, בעמוד 12. בספרות הועלו שאלות הקשורות באחריות, סיבתיות, ונוק הקשור בתוכנות ובבינה מלאכותית. המענה על שאלות אלה יבחן בהמשך.

2.2. העיסוק הממשלתי בבינה מלאכותית

החלטת ממשלה 212, כמו גם העבודה שנעשתה לקראת עריכת מסמך זה, התבססו, בין היתר, על עבודות קודמות שנעשו במהלך השנים האחרונות בנוגע לתחום הבינה המלאכותית.

בשנת 2018 ריכז המטה לביטחון לאומי במשרד ראש הממשלה (להלן: **המל"ל**) את הבחינה האסטרטגית בכל הנוגע לטכנולוגיות הבינה המלאכותית. במסגרת זאת ריכז המל"ל, בין היתר, את הקשר עם "המיזם הלאומי למערכת נבונות בטוחות" בראשות פרופ' איציק בן ישראל ופרופ' אביתר מתניה,²⁰ כאשר תחום האתיקה והמשפט נדון בצוות משנה של המיזם, בראשות פרופ' קרין נהון ובהשתתפות מומחים מהאקדמיה, מהממשלה ומהתעשייה.²¹ הוועדה קראה לבנות תכניות ידע והכשרה בתחום, לאמץ מערכת עקרונות אתיים בחברות וארגונים, לזהות עקרונות אתיים החשובים למקבלי החלטות ולהטיל אחריות על כלל העוסקים בבינה מלאכותית לפעול באופן חוקי ואתי. בנוסף, הוועדה הציעה כלי עבור מקבלי החלטות, לבחינת האתגרים האתיים במערכות בינה מלאכותית. בתחום הרגולציה, הוועדה מיפתה את סוגי הרגולציה הקיימים והמליצה, בין היתר, להקים מנגנון תיאום פנים-ממשלתי על-משרדי במטרה ליצור מדיניות רגולטורית אחידה, ברורה וקוהרנטית; לבחון את התפיסה של סביבה רגולטורית מבוקרת; לסייע לרשויות המאסדרות את משאבי המידע בהסרת חסמים; ולהטיל על רשות התחרות לגבש דרכי התמודדות שמטרתן שמירה על תחרות הוגנת, על צרכנים ועל נגישות הטכנולוגיה.²² בנוסף המל"ל סייע לרכז את עמדות ישראל לצורך ליווי הדיונים ב-OECD על עקרונות לתחום הבינה המלאכותית בשנת 2019.

בהמשך לכך, בשנת 2020 הפורום לתשתיות לאומיות למחקר ופיתוח (פורום אשר בראשו יושב נציג האקדמיה הלאומית למדעים וחברים בו משרד החדשנות, המדע והטכנולוגיה, משרד הביטחון, משרד האוצר, הות"ת ורשות החדשנות; להלן: **פורום תל"ס**) הקים ועדת בדיקה לבחינת הצורך בהתערבות ממשלתית לשם האצת התפתחות תחום הבינה המלאכותית ומדע הנתונים, בראשות ד"ר ארנה ברי.²³ באותה שנה פרסמה הוועדה את המלצותיה, ובמסגרת זו, בין היתר, המליצה על יצירת סביבה רגולטורית מאפשרת ומעודדת לבינה מלאכותית.

בשנת 2021 הפיצו מחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים ורשות החדשנות "קול קורא" משותף, שהזמין פניות ציבור לגבי חסמים רגולטוריים ורגולציה אפשרית לתחום הבינה המלאכותית בישראל, בדגש על ניסוי והטמעה של טכנולוגיות בתחום זה,²⁴ וזאת על רקע עיסוקן בבינה מלאכותית ובנסינות רגולטורית. במסגרת זו, התקיימו גם שולחנות עגולים בהם התקבלה תמונת מצב מסוימת מהתעשייה על חסמים, צרכים וכדומה.

²⁰ המיזם הלאומי למערכות נבונות בטוחות להעצמת הביטחון הלאומי והחוסן המדעי-טכנולוגי: אסטרטגיה לאומית לישראל: דו"ח מיוחד לראש הממשלה (חלק א': תמצית והמלצות, יצחק בן ישראל, אביתר מתניה וליהיא פרידמן עורכים (2020) (להלן: **המיזם הלאומי למערכות נבונות**) <https://icrc.tau.ac.il>; בהתאם לאמור בדין וחשבון של המיזם, ראשי המיזם התניעו את המיזם ביולי 2018, וגייסו מאות מומחים מהממשלה מערכת הביטחון האקדמיה והתעשייה, אשר פעלו בהתנדבות - 14 צוותים, וצוות מתכלל.
²¹ קרין נהון (יו"ר) ואח', **דין וחשבון של ועדת משנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה של בינה מלאכותית** (2019) (להלן: **דוח ועדת המשנה**) <http://ekarine.org>.

²² פורום תל"ס ועדת בינה מלאכותית ומדע הנתונים (2020) (להלן: **פורום תל"ס**) <https://innovationisrael.org.il>.
²³ ראו פנייה לציבור לקבלת מידע בנוגע לרגולציה וחסמים רגולטוריים ביחס לתחום הבינה המלאכותית, המחלקה הכלכלית במחלקת ייעוץ וחקיקה במשרד המשפטים הרשות הלאומית לחדשנות טכנולוגית (15.06.2021) <https://innovationisrael.org.il>.

בנוסף, בשנת 2021 הפיצה מחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים "קול קורא" לקבלת סקירה בנושא יישומי בינה מלאכותית בתחום השירותים הפיננסיים. בהמשך לכך, בשנת 2022 הוגש למחלקה ופורסם לציבור דו"ח בנושא "בינה מלאכותית במגזר הפיננסי: שימושים נפוצים, אתגרים וסקירה השוואתית של התמודדות רגולטורית", המסכם מחקר שביצעו חוקרים בכירים מאוניברסיטת תל-אביב (להלן: **דו"ח בינה מלאכותית בשירותים פיננסיים**).²⁵ הדו"ח כולל סקירה מקיפה של השימושים הנפוצים בבינה מלאכותית במגזר הפיננסי (ובכלל זה שימוש בבינה מלאכותית לחיתום אשראי וביטוח, לפעולות מסחר בניירות ערך, לייעוץ השקעות וניהול כספים, לשירות לקוחות ולציאט בוטים), ובחינה של ההזדמנויות הנובעות מהשימוש בטכנולוגיה זו לצד האתגרים המשפטיים והרגולטוריים שהיא מייצרת. בנוסף, הדו"ח מתייחס לאתגרים הכלליים שמעוררת הבינה המלאכותית בתחומים נוספים, וכולל הצגה מפורטת של יוזמות ההסדרה בעולם. כמו כן, הדו"ח כולל את המלצות החוקרים באשר לאופן שבו יש להמשיך את העיסוק בתחום זה.

בהמשך לדו"ח בינה מלאכותית בשירותים פיננסיים, הוקם הצוות הבין-משרדי לבחינת השימוש בבינה מלאכותית ולמידת מכונה במגזר הפיננסי. הצוות כולל נציגים מטעם רשות ניירות ערך, בנק ישראל, רשות שוק ההון הביטוח והחיסכון, רשות התחרות, משרד האוצר וייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים, וביום 24.4.2023 פרסם "קול קורא" להערות הציבור בעניינים אלה.²⁶

בנוסף, קודמו בשנים האחרונות פעולות סקטוראליות הקשורות להסדרת היבטים שונים של שימוש בבינה מלאכותית בהקשרים ספציפיים (כגון כלי רכב עצמאיים (אוטונומיים), שימושים בתחום הפיננסי כמפורט לעיל, וסוגיות של פרטיות. להרחבה על אודות אלה ראו להלן בחלקים "דיני פרטיות ו"הסדרה ייעודית מגזרית").²⁷

בהמשך לעבודות הנסקרות לעיל, במסגרת החלטת הממשלה 212 הוחלט על השקעה משמעותית בתחום הבינה המלאכותית, והוטל על שר החדשנות, המדע והטכנולוגיה להוביל את מדיניות הממשלה בנושא. במסגרת זו נקבע כי על משרד החדשנות, המדע והטכנולוגיה יחד עם גורמים נוספים, ובכלל זאת משרד רה"מ, משרד הביטחון, משרד האוצר, רשות החדשנות והוועדה לתכנון ולתקצוב הפועלת במועצה להשכלה גבוהה (להלן: **ות"ת**), לגבש תכנית לאומית לבינה מלאכותית.

יצוין כי החלטת הממשלה 212 מתווספת לפעילויות שקיימו גופים שונים בתחום אחריותם. כך למשל, נכון להזכיר את פעילות הות"ת לפיתוח והגדלת מספר החוקרים והסגל האקדמי בתחום מדעי הנתונים, וכן פעילות רשות החדשנות לתמיכה בתעשיית ההייטק לצורך קידום תחום הבינה המלאכותית – ובכלל זה באמצעות תמיכה ישירה (כגון תמיכה בפרויקטים טכנולוגיים ובהון אנושי)

²⁵ ראו משרד המשפטים, ארי אחיעז, אסף חמדני, דן עמירם וקובי קסטיאל **בינה מלאכותית במגזר הפיננסי: שימושים נפוצים, אתגרים וסקירה השוואתית של התמודדות רגולטורית** (2022) (להלן: **אחיעז, חמדני, עמירם וקסטיאל**) https://www.gov.il/he/departments/news/ai_report.

²⁶ "קול קורא מטעם הצוות הבין-משרדי לבחינת השימוש בבינה מלאכותית ולמידת מכונה במגזר הפיננסי" (2023), זמין לעיון כאן: <https://www.new.isa.gov.il/>.

²⁷ בשנת 2020, פרסמה הרשות להגנת הפרטיות במשרד המשפטים סקירה בנושא "דירוג חברתי בראי הזכות לפרטיות", ובה מציגה הרשות את הסיכונים הכרוכים בשימוש במידע רב ובנתוני עתק (big data) ובאלגוריתמים, לצורך "דירוג חברתי" של אנשים להשגת מטרות שלטוניות. בשנת 2022 פרסמה הרשות להגנת הפרטיות מסמך בנושא "חובת יידוע במסגרת איסוף ושימוש במידע אישי", ובמסגרתו נאמר כי חובת הגילוי לפי ס' 11 לחוק הגנת הפרטיות, התשמ"א-1981, החלה בעת איסוף מידע מאדם, חלה גם בהקשר של "מערכות לאיסוף מידע וקבלת החלטות מבוססות אלגוריתם או בינה מלאכותית". הרשות להגנת הפרטיות מציינת כי "הוראות סעיף 11 לחוק הגנת הפרטיות מטילות חובת יידוע בשלב איסוף המידע גם על גורמים האוספים מידע אישי באמצעות מערכות לקבלת החלטות מבוססות אלגוריתם (לרבות מערכות מבוססות בינה מלאכותית (AI)) וגם על גורמים האוספים מידע אישי לשם שימוש בו במסגרת מערכות שכולן" הרשות להגנת הפרטיות. הרשות להגנת הפרטיות **חובת יידוע במסגרת איסוף ושימוש במידע אישי** (2022) (להלן: **עמדת הרשות להגנת הפרטיות בעניין חובת יידוע במסגרת איסוף ושימוש במידע אישי**) <https://www.gov.il>.

ופעילות לקידום מדיניות שמטרתה לעודד את החדשנות (כגון אפיון תשתיות טכנולוגיות וקידום גולציה תומכת).²⁸ ברשות החדשנות מרוכזת תכנית תל"מ לבינה מלאכותית, שמטרתה תכלול וניהול התשתיות הלאומיות לבינה מלאכותית אשר על הקמתן סוכם בפורום תל"מ.

לאחר פרסום טיוטת מסמך זה להערות הציבור שבה הממשלה ועסקה בנושא במסגרת החלטה 173 מיום 24.2.2023. במסגרת זו החליטה הממשלה לאשר תכנית בבינה מלאכותית ומדעי הנתונים שתהווה הרחבה של תכנית תל"מ. עיקרי התכנית שאושרה על ידי הממשלה עוסקים בהאצת המחקר הבסיסי והיישומי בתחום הבינה המלאכותית; ביצירת נכסים לשימוש רחב של התעשייה, קידום מו"פ תשתיתי והפעלת מיזמים מתואמים בין מגזריים לניסוי טכנולוגיות מבוססות בינה מלאכותית; ובהטמעת יישומי בינה מלאכותית במגזר הציבורי האזרחי לטובת יעול עבודת המגזר הציבורי. כן הנחתה הממשלה, בהמשך להמלצות בטיטת מסמך זה את משרד החדשנות, המדע והטכנולוגיה להקים מוקד ידע ותיאום ממשלתי בתחומי הבינה המלאכותית, וזאת בשיתוף עם מחלקת ייעוץ וחקיקה (משפט כלכלי), והורתה לשר החדשנות, המדע והטכנולוגיה לפעול לקידום שיתופי פעולה בינלאומיים אזרחיים בתחומי הבינה המלאכותית, בשיתוף עם העוסקים בכך ברשות החדשנות, במשרד החוץ ובמחלקת ייעוץ וחקיקה (משפט בין לאומי ומשפט כלכלי).

²⁸ שם.

2.3. תחולה כללית של הדין והרגולציה בהיבטי בינה מלאכותית

חלק זה מוקדש לסקירה של הדין הכללי בישראל, מתחומים שונים, אשר חל, ככלל, גם על מצבים המערבים בינה מלאכותית. נציין דוגמאות שונות לתחולה אפשרית של הדין הקיים, באופן שמסדיר פיתוח מערכות מבוססות בינה מלאכותית או שימוש בהן. לשם הגיוון, דוגמאות אלה לקוחות מתוך ענפי משפט שונים (כגון דיני הגנת הצרכן, זכויות אדם, זכויות יוצרים, חוזים, נזיקין ופרטיות), וכן יעסקו בדינים מיוחדים שמטרתם הגנה על אוכלוסייה מסוימת (כגון צרכנים או קטינים).

מטרת חלק זה היא להמחיש באמצעות דוגמאות כיצד הדין והרגולציה הקיימים יכולים להיות רלוונטיים לטיפול בהיבטי בינה מלאכותית, אולם הדברים נכתבים לשם ההמחשה בלבד, ואין בהם כדי להשליך על שאלות משפטיות או פרשניות שטרם הוכרעו על ידי הגורם המוסמך. ממילא, הניתוח המפורט להלן הוא חלקי ואינו מתיימר למצות את הדיון בסוגיות השונות המועלות.

2.3.1. דיני זכויות יוצרים

איכותה של מערכת בינה מלאכותית תלויה ראש וראשית בהיקף הנתונים המוזנים לתוכה בשלב למידת המכונה, באיכותם ובגיוונם.²⁹ כאשר התכנים המשמשים לאימון המכונה אינם מוגנים בזכויות יוצרים, או שהם בבעלות מי שמאמן את המכונה, אין כל מניעה - ככל שבדיני זכויות יוצרים עסקינן - לעשות בהם שימוש. ואולם, הצורך בהיקף נרחב של תכנים איכותיים מחייב כמעט תמיד שימוש גם בתכנים המוגנים בזכויות יוצרים של גורמים נוספים. אי הבהירות באשר לנסיבות בהן ניתן לעשות שימוש בתכנים מוגנים במצבים אלה עלול להוות חסם משפטי לביצוע למידת מכונה אפקטיבית ולצמיחתו של שוק הבינה המלאכותית. אי ודאות זו עלולה גם להוות אבן נגף בדרכם של יוצרים לאכוף את זכויותיהם בזירה זו. בנובמבר 2022, פרסמה מחלקת ייעוץ וחקיקה (אזרחית – אשכול קניין רוחני) חוות דעת מקיפה בתחום שנועדה להסיר חסם זה.³⁰

חוות הדעת נועדה להבהיר את היקף הזכות של מיזמים מבוססי בינה מלאכותית לעשות שימוש בתכנים המוגנים בזכויות יוצרים לצורך למידת מכונה. על פי גילוי הדעת, מלבד מקרים חריגים, שימוש בתכנים מוגנים בזכויות יוצרים לצורך אימון מכונה חוסה תחת הסדרי השימושים המותרים בדיני זכויות היוצרים ולכן אינו מהווה הפרה של זכויות יוצרים.³¹ ראשית, כמפורט בגילוי הדעת, שימוש בתכנים מוגנים לשם אימון מכונה יהווה לרוב שימוש הוגן.³² שנית, בחלק ממיזמי למידת המכונה, יחול סעיף 22 לחוק זכות יוצרים, העוסק בשימוש אגבי ביצירה.³³ שלישית, ככל שהנתונים המשמשים ללמידה נמחקים בתום התהליך, יחול גם סעיף 26 לחוק זכות יוצרים בדבר יצירה ארעית.³⁴ במצבים מסויימים, הדין הקיים אינו מתיר שימוש בתכנים מוגנים בזכויות

²⁹ נתונים כאמור יכולים לכלול סוגים מגוונים של תוכן, כגון תמונות, קבצי טקסט, קבצי קול, קבצי וידאו וכדומה, בהתאם למשימה שהמערכת לומדת לבצע. ראו למשל, Google, *Training and Test Sets: Splitting Data*, <https://developers.google.com/machine-learning/crash-course/training-and-test-sets/splitting-data>.
³⁰ משרד המשפטים, *חוות דעת: שימושים בתכנים מוגנים בא לצורך למידת מכונה* (2022) <https://www.gov.il/BlobFolder/legalinfo/machine-learning/he/machine-learning.pdf> (להלן: *חוות דעת זכויות יוצרים*).

³¹ ראו ס' 19-30 ו-32 לחוק זכות יוצרים, התשס"ח-2007; ס' 33 לפקודת זכות יוצרים, 1911. ראו גם למשל ע"א 8393/96 *מפעל הפיס נ' The Roy Export Establishment Company*, פ"ד נד(1) 577, 596 (2000) ("בבואנו להגן על היצירה המקורית יש לתת את הדעת גם על כך שהגנה מוגברת יתר על המידה עלולה לבלום התפתחות תרבותית וחברתית, הסומכת על הישגי העבר").

³² ס' 19 לחוק זכות יוצרים; ראו גם פרק ב.1. לחוות דעת זכויות יוצרים, לעיל ה"ש 30.

³³ ס' 22 לחוק זכות יוצרים; ראו גם פרק ב.2. לחוות דעת זכויות יוצרים, לעיל ה"ש 30.

³⁴ ס' 26 לחוק זכות יוצרים; ראו גם פרק ב.3. לחוות דעת זכויות יוצרים, לעיל ה"ש 30.

יוצרים בהקשר של למידת מכונה, כגון כאשר המערכת מתאמנת על יצירות של יוצר יחיד ומייצרת כלים שמתחרים בו בשווקים בהם הוא פועל.³⁵

הבהרה של חוקיות השימוש בתכנים מוגנים לצורך למידת מכונה היא חשובה לנוכח מרכזיותן של מערכות בינה מלאכותית לכלכלה העולמית ולצורך שמירה על יתרונה התחרותי של ישראל כיצרנית של מערכות בינה מלאכותית.³⁶ הסרת החסם של עמימות משפטית באמצעות חוות הדעת נועדה לתמרץ פעילות רצויה בשוק למידת המכונה ולמקסם את התחרותיות של חברות ישראליות הן בתעשיית התוכן והן בתחום הבינה המלאכותית בשווקים גלובליים.

2.3.2. דיני פרטיות

סעיף 7 לחוק יסוד: כבוד האדם וחירותו קובע כי הזכות לפרטיות היא זכות יסוד. חוק הגנת הפרטיות, התשמ"א-1981 קובע עוולות ועבירות הקשורות בפגיעה בפרטיותו של אדם, והוראות הקשורות באיסוף מידע אישי לצורך עיבוד ממוחשב במאגר מידע. עקב כך, גם בכל הקשור לפעולות עיבוד מידע הנעשות כחלק מתהליכי "נתוני עתק" או למידת מכונה, צפויים לחול דיני הפרטיות.

הדין בישראל מאפשר פגיעה בפרטיות רק אם הפגיעה נעשית לפי חוק או אם ניתנה לכך הסכמה מדעת. כמו כן, הדין קובע הסדרים שונים לשימוש במידע, ובהם כי יש למסור לנושא מידע אודות השימוש לשמו נאסף מידע הנוגע אליו. כמו כן, יש לעשות שימוש במידע שנאסף רק למטרה אשר לשמה הוא נמסר (כלומר, לא ניתן לאסוף מידע בהסכמה מדעת למטרה אחת ולאחר מכן לעשות בו שימוש למטרה אחרת בלא קבלת הסכמה נוספת או היתר חוקי אחר).³⁷

הרשות להגנת הפרטיות במשרד המשפטים פרסמה גילוי דעת, שבו היא מתייחסת, בין היתר, לפרשנותה בעניין חובת היידוע במסגרת איסוף ושימוש במידע אישי כאשר נעשה שימוש בבינה מלאכותית.³⁸ במסגרת זו מציינת הרשות, כי כשאיסוף המידע נעשה באמצעות מערכות אוטומטיות, יש לוודא שיוצגו לנושא המידע כל הפרטים הנדרשים בהתאם להוראות החוק. עוד מציינת הרשות, כי כשעיבוד המידע נערך באמצעות מערכות אוטומטיות, על הליך היידוע לפרט על אופן פעולת המערכות, ככל שהדבר רלוונטי לגיבוש ההסכמה וככל שפירוט זה אפשרי מבחינה משפטית, טכנולוגית, ומסחרית. עוד ממליצה הרשות, כי יוסבר לנושא המידע על אודות פרטי המידע בהם עשויות המערכות להשתמש במסגרת השימוש במידע הנוגע אליו, והמקור של פרטי מידע אלו.

לצד ההוראות שצוינו לעיל, חוק הגנת הפרטיות ותקנות הגנת הפרטיות (אבטחת מידע), התשע"ז-2017 קובעים הוראות בעניין אבטחת מידע שבמאגר מידע. בנוסף, בהתאם להוראות החוק קיימת לאדם זכות לעיין במידע שעליו שבמאגר המידע וכן לתקן אותו או, בנסיבות מסוימות, למחקו. כאמור לעיל, על פני הדברים, הוראות אלה ואחרות צפויות לחול אף ביחס לשימוש בבינה מלאכותית.

³⁵ שם.

³⁶ שם; ראו גם, *Sizing the prize: PwC's Global Artificial Intelligence Study: Exploiting the AI Revolution*, Neil Savage, *The race to the top*; PwC Global <https://www.pwc.com/gx/en/issues/data-and-analytics/among-the-worlds-leaders-in-artificial-intelligence>, NATURE INDEX, (Sept. 9, 2020), <https://www.nature.com/articles/d41586-020-03409-8>; לירן ענתבי *בינה מלאכותית וביטחון לאומי בישראל* 85 (2020 INSS); דאטאנישן "שנת שיא: סטארטאפים של בינה מלאכותית גייסו 25.2 מיליארד דולר ב-2018" <https://www.themarker.com/technation/datanation/1.7017976>, (21.3.19).

³⁷ ראו עת"מ (מינהליים ת"א) 24867-02-11 אי.די.איי חברה לביטוח בע"מ נ' רשם מאגרי המידע, הרשות למשפט טכנולוגיה ומידע במשרד המשפטים (2012).

³⁸ עמדת הרשות להגנת הפרטיות בעניין חובת יידוע במסגרת איסוף ושימוש במידע אישי, לעיל ה"ש 27.

2.3.2. אפליה אסורה

על גופים שעליהם לא חל המשפט המנהלי חלים דינים שונים שמטילים חובה לנהוג בשוויון, כגון חוק שוויון הזדמנויות בעבודה, התשמ"ח-1988 (אשר אוסר על הפליה בין עובדים ובין דורשי עבודה), וחוק איסור הפליה במוצרים, בשירותים ובכניסה למקומות בידור ולמקומות ציבוריים, התשס"א-2000. חוק איסור הפליה מחיל חובות שוויון ואיסור אפליה על פעולות מסוימות במגזר הפרטי ועל גופים פרטיים (בנוסף לתחולתו על המדינה), ובין היתר אוסר על פרסום הכולל אפליה כאמור, ואף מעביר את הנטל בתביעה אזרחית בין השאר בהתאם לתוצאה מפלה. בהיעדר הסדר חוקי ספציפי, במקרים שבהם שימוש בבינה מלאכותית יגרום לפעולה שתוצאתה אפליה אסורה לפי החוק – על פני הדברים, החוק והוראותיו האמורות עשויים לחול על מי שאחראי לפעילות הבינה המלאכותית. כך למשל, עשויה להיות להן רלוונטיות במקרה של קבלת החלטה מפלה על ידי מערכת מבוססת בינה מלאכותית, או במקום שמוצע פרסום בהתבסס על פילוח המבוצע על ידי בינה מלאכותית. עם זאת, ראוי להדגיש כי תחולת החוק היא רק על הפעילות המוגדרת בו.³⁹ הפרת חוק זה היא עבירה פלילית וכן עוולה בנזיקין (מלבד עוולת הרשלנות שעשויה להיות רלוונטית), כך שהוא מאפשר גם תביעה נזיקית.

2.3.3. דיני חוזים

חוק החוזים (חלק כללי), התשל"ג-1973 (להלן: **חוק החוזים**), לא מונע ביצוע פעולות משפטיות באמצעים ממוחשבים, משום שאין בו ככלל דרישות של צורה וטכניקה. כיום נערכות עסקאות על סמך התקשרויות מול אתרי אינטרנט, בלי שעצם ביצוע הפעולה מעורר קושי משפטי. בהתאם לסעיף 23 לחוק החוזים ככלל אין דרישה צורנית לעריכת חוזה, ולכן ככלל אין מגבלה על שימוש במחשבים ובתקשורת אלקטרונית לעריכת חוזים.⁴⁰ בהקשר זה קבעה הפסיקה, בזיקה לפסיקה האמריקאית, הוראות לעניין היצע וקיבול ביחס לחוזים ממוחשבים.⁴¹

עמדה זו מבוססת גם על מסקנות ועדת המסחר האלקטרוני שפעלה במשרד המשפטים, על בסיס התפיסה בדבר הגמישות של הדין הקיים ויכולתו לחול גם על מצבים מורכבים מבחינה טכנולוגית שיתעוררו בעתיד.⁴²

עם זאת, קיומו של אלגוריתם המעורב ביצירת חוזה שונה משימוש במדיה אלקטרונית בלבד, ומעלה שאלה שנדונה בוועדת המסחר האלקטרוני בדבר היכולת של "סוכן אלקטרוני", כלומר

³⁹ ראו ס' 3 לחוק איסור הפליה במוצרים, בשירותים ובכניסה למקומות בידור ולמקומות ציבוריים, תשס"א-2000: "מי שעיסוקו בהספקת מוצר או שירות ציבורי או בהפעלת מקום ציבורי, לא יפלה בהספקת המוצר או השירות הציבורי, במתן הכניסה למקום הציבורי או במתן שירות במקום הציבורי, מחמת גזע, דת או קבוצה דתית, לאום, ארץ מוצא, מין, נטיה מינית, השקפה, השתייכות מפלגתית, גיל, מעמד אישי, הורות או לבישת מדי כוחות הביטחון וההצלה או ענידת סמליהם"; ס' 2 לחוק איסור הפליה:

"**שירות ציבורי**" – שירותי תחבורה, תקשורת, אנרגיה, חינוך, תרבות, בידור, תיירות ושירותים פיננסיים, המיועדים לשימוש הציבורי;

"**שירותים פיננסיים**" – שירותי בנקאות, מתן אשראי וביטוח;

"**שירותי תחבורה**" – אוטובוסים, רכבות, תובלה אווירית, אניות, שירותי הסעה והשכרת רכב";
לביקורת על גישה זו ולהצעת גישה מרחיבה לתחולת החוק ראו רונן אברהם "בנפרד ועדיין שווה? על דרכי ההתמודדות עם מקרי הפליה הנופלים בנפרד מתחולת חוק איסור הפליה" **המשפט** כח 13 (2022) (להלן: **אברהם**).

⁴⁰ ראו משרד המשפטים **הוועדה לבדיקת משפטי הכרוכות במסחר אלקטרוני דו"ח חלקי** 16 (2004) (להלן: **ועדת מסחר אלקטרוני**); וכן ראו תמר קלהורה "הצעת חוק מסחר אלקטרוני: מטרות ועקרונות" **חוקים** 5 ב (2010) (להלן: **קלהורה**).

⁴¹ ראו למשל תא (מרכז) 18763-04-15 **ויוה מדיה בע"מ נ' Google Ireland Ltd**, פס' 12-18 (2019).

⁴² ראו ועדת מסחר אלקטרוני לעיל ה"ש 40.

תוכנה הפועלת לפי הנחיות מי שערך אותה, לייצר התקשרויות באופן עצמאי. על אף שוועדת מסחר אלקטרוני דנה בנושא, בסופו של דבר הכלל שהציעה לא נכלל בהצעת חוק מסחר אלקטרוני.⁴³

על כן, ובהיעדר דין ספציפי העוסק בכך, לו תתעורר שאלה בדבר תחולת חוק החוזים על השימוש בבינה מלאכותית, יאפשר הדין הקיים מענים שונים. בין היתר, יהיה מקום לבחון שאלות בדבר ייחוס לצד המפעיל בינה מלאכותית להצעה להתקשר בחוזה, או למסירת הודעת קיבול להצעה להתקשר בחוזה, כוונה לבצע פעולה זו, כך שיישא בתוצאות הפעולה המשפטית.⁴⁴ עמדה זו נובעת מהאופן שבו מפותחת בינה מלאכותית ומשולבת בפעילות הארגון, המקדימה בדרך כלל לשימוש בבינה מלאכותית, והמעידה על רצונו וכוונתו של מפעיל הבינה המלאכותית לבצע פעולות משפטיות באמצעותה.⁴⁵

שאלה נוספת נוגעת לקיומה של חובת גילוי בנוגע לשימוש בבינה מלאכותית על ידי מי שעושה בה שימוש בדבר עצם השימוש (ראו להלן פרק "גילוי"), וזאת מכוח החובה לנהוג בתום לב במשא ומתן לפי סעיף 12 לחוק החוזים,⁴⁶ או בשל ההוראה בעניין הטעיה לפי סעיף 15 לחוק. זאת, על רקע תכונות הבינה המלאכותית ויכולותיה, והאפשרות כי השימוש בבינה מלאכותית יקנה לצד המתקשר יתרון מהותי בניהול משא ומתן לקראת החוזה.

תחום נוסף הוא החוזים האחידים. השימוש בבינה מלאכותית חושף את הפעילות לסיכונים שונים, ובין היתר סיכונים אבטחת מידע. במידה שמנסח חוזה אחיד יעביר את הסיכון הכרוך בשימוש בבינה מלאכותית לצד השני, יהיה מקום לבחון שאלות הנוגעות לקיומה של תניה מקפחת בחוזה אחיד.

2.3.4. הגנת הצרכן

חוק הגנת הצרכן, התשמ"א-1981 קובע הוראות שמטרתן הגנה על צרכנים. בהתאם לכך, מקום שבינה מלאכותית מעורבת בתהליך יצירת עסקה צרכנית, עשויות לחול על מי שמפעיל אותה החובות המוגברות הקשורות בגילוי בהתאם לחוק זה. בין היתר, יתכן שיהיה צורך לבחון אם אותם פערי הידע בין הצרכן לבין העוסק המפעיל את הבינה המלאכותית מתקיימים או אף מתגברים עקב כך שכנראה המידע שמצוי באופן בלעדי אצל העוסק כעת יהיה גדול עוד יותר.

בסעיף 7א לחוק הגנת הצרכן נקבע איסור על פרסום או שיווק פוגעניים לקטינים.⁴⁷ בהתאם לסעיף נקבעו גם תקנות הגנת הצרכן (פרסומת ודרכי שיווק המכוונים לקטינים), התשנ"א-1991, שמטרתן מניעת ניצול חוסר ניסיונם ופערי הכוחות הקוגניטיביים שבין עוסקים לקטינים. הוראות אלה עשויות להיות רלוונטיות אף הן ככל שנעשה שימוש בבינה מלאכותית בתקשורת עם קטינים. הרשות להגנת הצרכן ולסחר הוגן הודיעה כי היא רואה באתר אינטרנט המציע תכנים לקטינים כמי

⁴³ קלהורה, לעיל ה"ש 40.

⁴⁴ יובהר כי מוקד הדיון בהקשר זה אינו דרישות הצורה לפי ס' 23 לחוק החוזים, כי אם לעניין הביטוי "גמירת דעת" הנדרש כחלק מקיום "הצעה" לפי ס' 2 לחוק ולעניין "קיבול" לפי ס' 5 הדורש כי הקיבול יהיה "על דעתו" של צד לחוזה. "2. הצעה

פנייתו של אדם לחברו היא בגדר הצעה, אם היא מעידה על גמירת דעתו של המציע להתקשר עם הניצע בחוזה והיא מסויימת כדי אפשרות לכרות את החוזה בקיבול ההצעה; הפניה יכול שתהיה לציבור.

...

5. קיבול
הקיבול יהיה בהודעת הניצע שנמסרה למציע והמעידה על גמירת דעתו של הניצע להתקשר עם המציע בחוזה לפי ההצעה."

⁴⁵ כמפורט לעיל, בעמ' 16-19.

⁴⁶ ראה: גבריאלה שלו ואפי צמח דיני חוזים 12-13 (מהדורה רביעית 2019).

⁴⁷ בהתאם לס' 7א(א) לחוק: "לא יפרסם אדם פרסומת ולא ינקוט דרך שיווק אחרת אם הפרסומת או דרך השיווק כאמור עלולות להטעות קטין, לנצל את גילו, תמימותו או חוסר ניסיונו או לעודד פעילות שיש בה כדי לפגוע בגופו או בבריאותו הגופנית או הנפשית."

שאחראי לפרסומות המוצגות לצד התוכן, ללא קשר למידת השליטה או הידיעה בפועל של האתר על אודות פרסומות אלה, ובכלל זה לאפשרות שהפרסומות נובעות ממידע שנאסף ממכשיר הקצה שבו הן מוצגות.⁴⁸

2.3.5. דיני נזיקין

שימוש במערכות מבוססות בינה מלאכותית הגורם לנזק עשוי לחשוף את המזיק לאחריות בנזיקין בגין הנזק, בתנאי שניתן יהיה להוכיח את יסודות האחריות בנזיקין.⁴⁹ במסגרת "קול קורא" שפרסמה מחלקת ייעוץ וחקיקה (משפט אזרחי) במשרד המשפטים⁵⁰ הוצגו באופן כללי שאלות מרכזיות הקשורות בדיני הנזיקין והשימוש בכלי רכב עצמאיים, אשר יש להן רלוונטיות לבחינת קביעת חבות נזיקית.⁵¹ דו"ח בינה מלאכותית בשירותים פיננסיים העלה שאלות דומות לגבי אופן יישום משטרי האחריות המקובלים בנזיקין לגבי שימוש במערכות מבוססות בינה מלאכותית.⁵²

הדיון באחריות נזיקית בנסיבות קונקרטיות הוא תלוי נסיבות והקשר, וגם נגזר מהמסגרת הנזיקית החלה, כגון עוולת הרשלנות או הפרת חובה חקוקה, או העוולות הפרטניות. בלי לקבוע מסמרות, נראה כי תיתכן אפשרות לתחולת דיני הנזיקין על סיטואציות שבהן נגרם נזק או התנהגות עוולתית הקשורות בשימוש בבינה מלאכותית. דיני הנזיקין, במיוחד עוולת המסגרת של עוולת הרשלנות, גמישים דיים כדי להתמודד עם מצבים עובדתיים חדשים ועם שינויי הזמן, לרבות שינויים טכנולוגיים, משום שהם שואלים שאלות עקרוניות וכלליות שניתן לבחון אותן במצבים חדשים מגוונים. במסגרת יישום יסודות עוולת הרשלנות יעלו, בין היתר, שאלות אלה: האם יש חובת זהירות? האם חובת הזהירות הופרה? האם נגרם נזק עקב כך? האם הנזק הוא במתחם הסיכון שיצרה ההתנהגות? האם הוא לא רחוק מדי? שאלה מרכזיות היא האם מי שבונה מערכת של בינה מלאכותית מסוימת יכול וצריך לצפות שייגרם נזק מסוים כתוצאה מהפעלתה או השימוש בה. המסקנה הנורמטיבית תלויה הן בהיבטים עובדתיים טכנולוגיים והן בשיקולי מדיניות.

⁴⁸ ראו עוד את עמדת הרשות להגנת הצרכן וסחר הוגן בעניין ת"צ 40801-09-20, המטילה אחריות על אתרי תוכן למניעת פרסום אסור לקטינים. הרשות קבעה עמדה זו למרות הטענה כי אופן הצגת הפרסומות אינו בשליטה ישירה ומודעות של מנהלי אתרי התוכן, וכי הם עשויים להיות מושפעים ממידע על אודות המשתמש המתקבל ממכשיר הקצה <https://www.law.co.il/news/>.

⁴⁹ "הטלת האחריות בנזיקין מותנית ככלל בהתקיימותם של שני יסודות אחריות הכרחיים: היסוד האחד הוא התנהגות עוולתית, דהיינו מעשה או מחדל או פעילות של המזיק או מצב שבו הוא נמצא, שהם תנאי והצדקה להטלת אחריות נזיקית עליו. היסוד השני של האחריות הוא קיומה של זיקה מספקת בין ההתנהגות העוולתית לבין הנזק, זיקה הנדרשת לצורך ייחוס הנזק למזיק. צוין כי שילוב שני יסודות אלה לרוב הכרחי אף לא תמיד מספיק, וכי יש שנדרש יסוד אחריות נוסף, בעיקר זה הובחן אם תוצאות הטלת האחריות הן תוצאות רצויות. עוולת הרשלנות אכן מושתתת על שלושת יסודות אלה". ישראל גלעד **דיני נזיקין – גבולות האחריות** כרך א 418 פס' 6.2 (2012).

⁵⁰ ראו משרד המשפטים, **אסדרת השימוש ברכבים אוטונומיים - פנייה לקבלת עמדת הציבור בתחום הנזיקין והביטוח** (2021) (להלן: **אסדרת השימוש ברכבים אוטונומיים - פנייה לקבלת עמדת הציבור**) <https://www.gov.il>.

⁵¹ השאלות המרכזיות המועלות בקול הקורא בתחום אחריות נזיקית לרכב עצמאי (אוטונומי) כמצע לעיצוב ההסדר הראוי: "חסרונה של אישיות משפטית" משום שהרכב נשלט בידי מחשב, "הוספת גורם אחריות חדש" (משום שיש שילוב של גורמי בקרה), "העדר צפיות והעדר קשר סיבתי בין שלב התכנון והייצור לבין התנהגות הרכב האוטונומי" (עקב היות הרכב לומד), "מורכבות וחדשנות המידע הנדרש לצורך הוכחת רשלנות בתביעת נזק רכוש", "too much information" (איסוף מידע רב מהווה יעד לתקיפה ומקשה על הבנת התמונה), "אי וודאות בנוגע לתוחלת הנזק הצפוי", "תמחור ביטוח בהעדר מידע", "אחריות תורמת של משתמשי הדרך והמשתמשים ברכב", "העדר מנגנוני תקינה ורישוי". על רקע האמור מועלים לדיון גם פתרונות ובהם: הצורך במשטר של אחריות מוחלטת והגדרת מאפייניה ומי חב לפיה, היחס בין הסדרי האחריות לבין ביטוח, היחס בין רכבים רגילים לרכבים אוטונומיים, סוגיית המידע הנאסף.

⁵² ראו אחיעז, חמדני, עמירם וקסטיאל, לעיל ה"ש 25, בעמ' 102-103, המעלה שאלות לגבי הטלת אחריות על תוכנה, חלוקת האחריות שבין המוסד הפיננסי, כותב התוכנה, והלקוח עצמו, היבטי הצפיות, וכן החשש שהטלת אחריות כפידיה תביא לחסמי כניסה לשימוש בטכנולוגיה זו.

נראה עוד, כי דיני הנזיקין מאפשרים באופן עקרוני לבחון את חלוקת האחריות בין מפתח ומשתמש בבינה מלאכותית בהתאם למנגנון חבות ביחד ולחוד. בהתאם למנגנון זה, האחריות המיוחסת לגורם אחד אינה מוציאה לרוב את האחריות של גורם אחר במישור היחסים שלהם מול הנזוק. חלוקת האחריות בין המזיקים מבוססת על "מבחן האשם המוסרי" שמקצה את חלק הארי של האחריות למי שאשמו המוסרי רב יותר.

2.3.6. הסדרה ייעודית מגזרית

לצד הוראות המשפט הפרטי, והוראות החלות מכוח אסדרת דיני פרטיות במידע, חלות גם הוראות בתחומי אסדרה ספציפיים במשק. בהתאם לאופן שבו התפתחה האסדרה בתחום טכנולוגיית המידע, נראה כי מאסדרים יוכלו, במקרים המתאימים, בכפוף לסמכותם על פי הדין, לקבוע הוראות וכללי התנהגות פרטניים לגבי הבינה המלאכותית. לצד זאת, ייתכן כי כללים מסוימים יהיו טעונים עיגון בחקיקה ראשית, וזאת בהמשך לאופן הצגת הדברים בהנחיית היועץ 1.2500.⁵³

כך, לדוגמה, בשנים האחרונות קידמה הממשלה את תיקון מספר 130 לפקודת התעבורה,⁵⁴ שמטרתו להוות תשתית חוקית לניסוי בשימוש ציבורי בכלי רכב עצמאיים במרחב הציבורי, לגבש תשתית ידע בנושא, ולבסס את אמון הציבור בכלי רכב עצמאיים.⁵⁵ תיקון זה לפקודת התעבורה קובע דרישות רגולטוריות ממי שמפעיל רכב עצמאי, ובכלל זה בתחום הבטיחות והגנת הסייבר, בטרם הפעלת הרכב, וכן בכל מחזור החיים של השימוש בו במרחב הציבורי.⁵⁶ התיקון מבוסס על קביעת העקרונות הנדרשים בחקיקה ראשית, ופירוט בתוספת לחוק או בתקנות מכוח הפקודה, והשלמת ההוראות במסגרת הוראות פרטניות לבעל ההיתר.⁵⁷

כמו כן, נעשית כאמור עבודה לבחינת ההיבטים הקשורים ביישומי בינה מלאכותית בתחום הפיננסי, וזאת בהמשך לדו"ח בינה מלאכותית בשירותים פיננסיים. במסגרת הדו"ח, ממליצים החוקרים לבחון ולהסדיר את השימושים בבינה מלאכותית בתחום הפיננסי באופן ספציפי וייעודי, ולעסוק ברגולציה הפיננסית כמקשה אחת ולהסדירה באופן תואם, אך לא במשותף עם ענפי משק אחרים. דו"ח זה, לרבות ההמלצה האמורה, נבחנים על ידי הצוות הבין-משרדי לבחינת השימוש בבינה מלאכותית ולמידת מכונה במגזר הפיננסי, שעוסק בנושא.

2.3.7. משפט מנהלי

כללי המשפט המינהלי מתווים מבחינה משפטית את דרך הפעולה של גופים מינהליים – משרדי ממשלה, יחידות סמך, חברות ממשלתיות וכדומה. שילוב הבינה המלאכותית בפעילותם של גופים ציבוריים אלה – אם ככלי תומך החלטה או כחלק מביצוע שלבי קבלת ההחלטה המנהלית בידי אלגוריתם – מחייב יישום של עקרונות המשפט המנהלי.

⁵³ הנחיית היועץ המשפטי לממשלה 1.2500, לעיל ה"ש 18, בעמ' 8-9.

⁵⁴ ראו חוק לתיקון פקודת התעבורה (מס' 130), התשפ"ב-2022, ס"ח 2967.

⁵⁵ בהתאם לס' 16 לתיקון לפקודת התעבורה: "מטרתו של סימן זה לקבוע הסדרים שיאפשרו הפעלה של רכב עצמאי בלא נהג, בדרך, למטרת ביצוע ניסוי, תוך שמירה על בטיחות הנוסעים ברכב ועוברי הדרך ושימוש במגוון טכנולוגיות, כדי להביא לגיבוש תשתית ידע לגבי בטיחות הרכב העצמאי, יכולתו להשתלב באופן בטוח בין עוברי הדרך ולתת שירות לנוסעים והשפעתו על התנועה בדרך, לאפשר הנגשה של הידע האמור לציבור, ולבסס את אמון הציבור בו".

⁵⁶ ראו ס' 16 (חובת קבלת היתר הפעלה) ו-16 (תנאים למתן היתר הפעלה) לתיקון לפקודת התעבורה. יוער כי בהתאם לס' 16 (שמירת דינים) לתיקון לפקודה נקבע כי: "אין בהוראות לפי סימן זה כדי לגרוע מחובתו של בעל היתר לעמוד בהוראות כל דין, ובכלל זה חוק פיצויים לנפגעי תאונות דרכים, התשל"ה-1975, ובהוראות הדין לעניין חובת ביטוח, ובכלל זה הוראות פקודת הביטוח". בנושא זה ראו **אסדרת השימוש ברכבים אוטונומיים - פנייה לקבלת עמדת הציבור**, לעיל ה"ש 50.

⁵⁷ ראו ס' 16 ח, יגב) לתיקון לפקודת התעבורה.

בתוך כך עשויות לעלות שאלות לגבי הסמכות המינהלית, ההליך המינהלי ודרך הפעלת שיקול הדעת המינהלי, כמו גם בנוגע להליך הביקורת השיפוטית עצמו ולאכסנייה המוסדית הרלוונטית לבחינת החלטות אלגוריתמיות.

כאמור, בנושא שימושי הממשלה בכלי בינה מלאכותית ותחולת המשפט המנהלי בהקשר זה מתקיימת עבודת מטה של הגורמים המשפטיים הרלבנטיים בממשלה, בהובלת מחלקת ייעוץ וחקיקה במשרד המשפטים, ובמטרה לסייע ליועצים המשפטים במשרדים בקידום שימוש בבינה מלאכותית באופן העונה על דרישות המשפט המנהלי. משכך, הוא לא נידון במסמך זה.

לסיכום, הדין הקיים מאפשר התמודדות מסוימת עם חלק מהאתגרים המתעוררים עקב שימוש במערכות מבוססות בינה מלאכותית, ובכלל זה אפליה, שקיפות, הצורך בהנמקה, אבטחת מידע, סייבר ופרטיות.

3. חלק שלישי: פעילות ארגונים בינלאומיים ומדינות מפותחות

מדינות רבות וארגונים בינלאומיים מנהלים שיח בעל אופי משפטי, רגולטורי ואתי ביחס לבניה מלאכותית, ומסמכי מדיניות רבים עוסקים בעקרונות שיש להביא בחשבון בעת פיתוח ושימוש בבניה מלאכותית.

מסגרות מדיניות אלה כוללות לרוב התייחסות לתועלות המשמעותיות הנובעות מהבניה המלאכותית, ולצידן התייחסות לאתגרים המתעוררים בקשר לבניה המלאכותית, וסקירה של אמצעים להתמודדות עם אתגרים אלה. קביעת מסגרות המדיניות נעשית ככלל, על מנת להגן על זכויות יסוד ואינטרסים ציבוריים, לחזק את הוודאות המסחרית והמשפטית, ולהגביר את אמון הציבור בבניה המלאכותית ולעודד שימוש בה. מכנה משותף מרכזי למסגרות המדיניות הוא התייחסות לעקרונות כגון אחריותיות, בטיחות, שקיפות, מניעת אפליה, וכיבוד זכויות אדם. בנוסף, מסגרות המדיניות כוללות אמצעים ליישום העקרונות וכן שיטות לאיתור שימוש בבניה מלאכותית המייצרת סיכון גבוה לעקרונות וערכים אלה, ושיטות לניהול הסיכונים על מנת לאפשר פיתוח ושימוש בבניה מלאכותית באופן מהימן.

לצד מכנים משותפים אלה, ניתן לראות שוני במאפייני המדיניות, בשני היבטים מרכזיים. מאפיין מרכזי אחד הוא **רמת הפיתוח והפירוט של העקרונות והאמצעים לקידום** – במדינות מסוימות העקרונות והאמצעים ליישומם מתוארים באופן כללי מאוד (כגון סינגפור). לעומת זאת, בטיטות החקיקה של האיחוד האירופי על בניה מלאכותית לדוגמה, מוצעות דרישות מפורטות כבר בשלב זה של התפתחות הטכנולוגיה. מאפיין מרכזי שני הוא **המעמד הנורמטיבי המוצע לקידום המדיניות** – כלומר האם מוצע לקדם אותה בחקיקה ייעודית רוחבית, או שמא מוצע לקדם אותה על בסיס התשתית המשפטית הקיימת תוך קידום הסדרים נקודתיים בהתאם לצורך. כך לדוגמה, האיחוד האירופי פועל לקידום חקיקה רוחבית ומקיפה, הכוללת גם הוראות מפורטות לגבי הדרישות ממי שמפתח ומשתמש בבניה מלאכותית. מנגד, בבריטניה מוצעת מדיניות ממשלתית המבוססת על קביעת מספר עקרונות כלליים, אשר נועדו להנחות את הרשויות הרגולטוריות בלבד. אופן קידום המדיניות נובע, בין היתר, מהגישה הרגולטורית הכללית הנוהגת, וממאפייני השווקים הטכנולוגיים.

עם זאת, גישות שונות אלה עלולות להביא ליצירת חסמי סחר ושיתוף פעולה. על רקע זה, מתנהלים מגעים בין ארה"ב והאיחוד האירופי, במסגרתם גובשה "מפת דרכים משותפת בנוגע לכלי הערכה ומדידה למערכות בניה מלאכותית אמינות ולניהול סיכונים", אשר נועדה בין היתר לממש את המחויבות המשותפת של ארה"ב והאיחוד האירופי לעקרונות ה-OECD.⁵⁸ מאמצים דומים מובלים על ידי ארגונים בינ"ל נוספים, כגון ה-G7 (ארגון המדינות המתועשות).⁵⁹

3.1. המלצות ה-OECD בתחום הבניה המלאכותית

בשנת 2019 אישר ה-OECD שורת המלצות בתחום הבניה המלאכותית (Recommendation of the Council on Artificial Intelligence Principles for responsible stewardship of).⁶⁰ ההמלצות כוללות שני חלקים: "עקרונות לניהול ושימוש אחראי בבניה מלאכותית מהימנה" (Principles for responsible stewardship of).

⁵⁸ TTC Joint Roadmap for Trustworthy AI and Risk Management, 1.12.2022 <https://digital-strategy.ec.europa.eu/en/library/ttc-joint-roadmap-trustworthy-ai-and-risk-management>

⁵⁹ G7 Hiroshima Leaders' Communiqué, 20.5.2023, <https://www.mofa.go.jp/files/100506878.pdf>

⁶⁰ ראו המלצות ה-OECD, לעיל ה"ש 8.

"trustworthy AI", ו-"מדיניות לאומית ושיתוף פעולה בינלאומי לבינה מלאכותית מהימנה" (National Policies and international cooperation for trustworthy AI). זהו המסמך הראשון שגובש על ידי ארגון בינלאומי מוביל, אשר מעגן עקרונות כלליים לבינה מלאכותית. לצד פיתוח השיח הנורמטיבי בתחום זה, מטרת ההמלצות היא גם לעודד תאימות (interoperability) בין מדינות ה-OECD ואחרות במדיניות ובהסדרים המיישמים את העקרונות. זאת, כדי לקדם את הפעילות הכלכלית הבינלאומית בתחום הטכנולוגיה והמידע, כמנוף לקידום יעדי הפיתוח והצמיחה של מדינות ה-OECD.

ההמלצות אומצו על ידי כל מדינות ה-OECD וכן על ידי מדינות נוספות.⁶¹ בשנת 2019 הן שימשו כמקור השראה למסמך העקרונות של ה-G20,⁶² ארגון המאגד את 20 הכלכלות הגדולות בעולם שבו חברות 19 מדינות והאיחוד האירופי, ובהן מדינות נוספות שאינן חלק מ-OECD.⁶³ אף שהעקרונות אינם מחייבים, המלצות ה-OECD הן בדרך כלל בעלות השפעה משמעותית על מי שאימצו אותן,⁶⁴ וכן בקביעת סטנדרטים בינלאומיים ובהכוונת מדיניות.⁶⁵ בקרב מדינות ה-OECD יש גישות שונות לאופן קידום טכנולוגיה ולתפקיד החקיקה והרגולציה בתחום הטכנולוגיה. עקב כך, להמלצות שהתקבלו בהסכמה בין המדינות החברות, ובפרט בין מדינות האיחוד האירופי לארה"ב ולמדינות כמו יפן שגישתן כאמור שונה, נודעת חשיבות רבה כרף בינלאומי המשקף מדיניות מוסכמת. במסגרת שיתוף הפעולה המתואר לעיל בין האיחוד האירופי לארה"ב, מוזכר נושא התיאום במדיניות בינה מלאכותית, והמלצות ה-OECD ופעילותו מצוינות כנקודת מוצא משותפת.⁶⁶

ישראל חברה ב-OECD ולקחה חלק בדיונים על ניסוח ההמלצות. בהתאם, החלטת ממשלה 212 הורתה להתחשב בין היתר בעקרונות ה-OECD במסגרת גיבוש המדיניות הלאומית בתחום הבינה המלאכותית. להמלצות ה-OECD חשיבות גם להמלצות שיפורטו בהמשך מסמך זה, הן בתחום ה"עקרונות אתיים לבינה מלאכותית", והן בהצבעה על כיוונים אפשריים לקידומן בפועל באמצעות הסדרים משפטיים ורגולטוריים.

הפרק הראשון להמלצות ה-OECD כולל עקרונות לניהול אחראי של בינה מלאכותית מהימנה, שמיועדים לחול על ארגונים ציבוריים ופרטיים המפתחים בינה מלאכותית או משתמשים בה. העקרונות נועדו לקדם תפיסה של בינה מלאכותית שבה "האדם במרכז" (Human-Centered AI). זאת, לנוכח עליית חשיבותה של טכנולוגיה זו והשפעת התחזיות, ההמלצות או ההחלטות שלה על

⁶¹ ראו בקישור לעיל: מדינות אלה כוללות את: אוסטרליה, אוסטרליה, בלגיה, קנדה, צ'ילה, צ'כיה, דנמרק, אסטוניה, פינלנד, צרפת, גרמניה, יוון, הונגריה, איסלנד, ישראל, איטליה, יפן, קוראה, לטביה, ליטא, לוקסמבורג, מקסיקו, הולנד, ניו זילנד, נורבגיה, פולין, פורטוגל, סלובקיה, סלובניה, ספרד, שוודיה, שווייץ, טורקיה, הממלכה המאוחדת, ארצות הברית. וכן מדינות נוספות: קולומביה, קוסטה ריקה, ארגנטינה, ברזיל, מצרים, מלטה, פרו, רומניה, סינגפור, ואוקראינה.

⁶² ראו את המידע המופיע לצד המלצות ה-OECD לעיל ה"ש 8, וכן ראו את הצהרת הסיכום של המפגש ביוני 2019: <https://www.oecd.org/science/forty-two-countries-adopt-new-oecd-principles-on-artificial-intelligence.htm> (להלן: **OECD סיכום מפגש יוני 2019**) וכן: <https://www.global-solutions-initiative.org>.

⁶³ בין מדינות אלה סין, רוסיה, ערב הסעודית, הודו, אינדונזיה, דרום אפריקה, ארגנטינה, ברזיל.

⁶⁴ על מעמדן של המלצות ה-OECD ראו <https://www.oecd.org/>.

⁶⁵ **OECD סיכום מפגש יוני 2019**, לעיל ה"ש 62; ההודעה מטעם ה-OECD מציינת למשל את ההשפעה של עקרונות ה-OECD לגבי הגנה על הפרטיות במידע, שהתקבלו ב-1980, והיו בסיס לקידום מדיניות וחקיקה בתחום הפרטיות, ומשמשות מכנה משותף אחיד ברמה הבינלאומית.

⁶⁶ section 19 EU-US Joint Statement of the Trade and Technology Council, https://ec.europa.eu/commission/presscorner/detail/en/IP_22_3034.

בני אדם.⁶⁷ על רקע מאפייני טכנולוגיית הבינה המלאכותית, ועדת המומחים של ה-OECD המליצה על נושאים שצריכים להיות בעדיפות לצורך קידום המטרה של "בינה מלאכותית שבה האדם במרכז".⁶⁸ נושאים אלה באו לידי ביטוי בעקרונות שנכללו בהמלצות, והם: (1) תרומה של הבינה המלאכותית לפיתוח בר קיימא ורווחה חברתית; (2) כיבוד ערכים השמים את האדם במרכז וערך ההגיונות; (3) השימוש במערכות בינה מלאכותית ואופן פעולתן צריך להיות שקוף; (4) מערכות בינה מלאכותית צריכות להיות עמידות ובטוחות; (5) נדרשת אחריות לתוצאות של הבינה המלאכותית.⁶⁹ מרכיב משותף לעקרונות בפרק הראשון של מסמך ההמלצות הוא שיש להחיל אותם בהתאם להקשר ולרמת הסיכון הנובעת מהטכנולוגיה, ובשים לב לתועלות החברתיות מבינה מלאכותית.

בפרק השני להמלצות ה-OECD מוצעים כלי מדיניות שנועדו לקדם את העקרונות ובאופן כללי "בינה מלאכותית מהימנה" (Trustworthy AI) במסגרת מדינות ה-OECD, כגון קידום מדיניות לאומית לתחום הבינה המלאכותית וכן קידום הון אנושי בתחום והשקעה במחקר ופיתוח.

מאז גיבוש ההמלצות בשנת 2019, ה-OECD משקיע משאבים ותשומת לב בקידום יישום העקרונות באמצעות שיח רצוף על אופן מימוש העקרונות והכלים המעשיים הנדרשים לממשלות וארגונים. פעילות זו תומכת בקידום הבנה והתמודדות עם בינה מלאכותית, וכן שואפת לייצר מכנה משותף בינלאומי גם בהיבטים מעשיים של מימוש העקרונות. במסגרת זאת, הוקמה ופועלת ב-OECD קבוצת עבודה ייעודית לבינה מלאכותית (Working Party on Artificial Intelligence Governance (AIGO)) במסגרת הוועדה למדיניות כלכלה דיגיטלית (Committee on Digital Economy Policy (CDEP)). כמו כן, ה-OECD הקים פורטל הכולל מידע רב בנוגע למדיניות בתחום הבינה המלאכותית במגוון היבטים.⁷⁰ תוצר מרכזי אחד של פעילות זו הוא הצעה לשיטה לסיווג ומיון מערכות מבוססות בינה מלאכותית.⁷¹ שיטה זו חשובה לצורך איתור ההזדמנויות והסיכונים, והתאמה של המענה בתחום ניהול הסיכונים לסוג המערכת ולסיכונים הכרוכים בה. פעילות מרכזית נוספת היא קידום מכנה משותף באשר לאופן ניהול סיכונים הבינה המלאכותית, בשים לב לפעילות רבה של ארגוני התקינה השונים בעולם בתחום זה.

3.2. טיוטת החקיקה של האיחוד האירופי⁷²

האיחוד האירופי מבוסס על מספר אמנות רב צדדיות בין מדינות אירופה, שעל בסיסן ניתן לקדם בין היתר חקיקה במדינות החברות בתחומים הקבועים באמנות.⁷³ נציבות האיחוד האירופי, הרשות המבצעת של האיחוד, מקדמת במסגרת מדיניות האיחוד, בתחום הדיגיטלי בכלל ובתחום הבינה המלאכותית בפרט, חקיקה רוחבית שנועדה לחול על רוב השימושים האזרחיים בבינה מלאכותית.

⁶⁷ דברי ההסבר והרקע לעקרונות מופיעים בפרסום שיצא במקביל מה-OECD, שנקרא "Artificial Intelligence in Society". בהתאם לכך, האמור בטקסט מבוסס על ההסברים המופיעים בפרסום זה. AI in Society, לעיל ה"ש 7, פרק 4 בעמ' 82 https://www.oecd-ilibrary.org/science-and-technology/artificial-intelligence-in-society_eedfee77-en.

⁶⁸ שם, בעמ' 83.

⁶⁹ שם.

⁷⁰ OECD AI Policy Observatory <https://oecd.ai/en/>.

⁷¹ OECD, OECD Framework for the Classification of AI Systems: A Tool for Effective AI Policies <https://oecd.ai/en/classification>.

⁷² חלק זה בסקירה מבוסס על גם אחיעז, חמדני, עמירם וקסטיאל, לעיל ה"ש 25, בעמ' 28-36.

⁷³ European Union, European Union – Principles, countries, history https://european-union.europa.eu/principles-countries-history_en.

בעקבות שימוע ציבורי ומחקר, זיהתה הנציבות שש בעיות מרכזיות העלולות לנבוע מפיתוח ושימוש בבינה מלאכותית, המצדיקות התערבות באמצעות חקיקה. בעיות אלה הן: (1) סיכון מוגבר לבטיחות ואבטחה; (2) סיכון מוגבר לפגיעה בזכויות אדם וערכי האיחוד; (3) היעדר כלים ומשאבים נדרשים בידי רשויות הרגולציה לפקח על פיתוח בינה מלאכותית ועל כך שהשימוש בה נעשה בהתאם לכללים; (4) אי-ודאות ומורכבות ביחס לכללים החלים על בינה מלאכותית שמרתיעים עסקים מפיתוח או שימוש בבינה מלאכותית; (5) החשש שחוסר אמון בבינה מלאכותית יפגע בפיתוח בינה מלאכותית באירופה ויפחית את התחרותיות של הכלכלה האירופית; (6) כלי מדיניות מבוזרים מייצרים חסמים לסחר ומסכנים את הריבונות הדיגיטלית של האיחוד.⁷⁴

בהמשך לכך, בשנת 2021 וכחלק מהמדיניות האירופית הכוללת לקידום בינה מלאכותית,⁷⁵ הנציבות פרסמה טיוטת חקיקה בתחום הבינה המלאכותית אשר נועדה להתמודד עם בעיות אלו. החקיקה המוצעת על ידי נציבות האיחוד האירופי ונדונה בפרלמנט האירופי קובעת כללי התנהגות והסדרה מפורטים שמטרתם מימוש העקרונות הכלליים של בינה מלאכותית מהימנה (Trustworthy AI), ובכך היא מהווה את מודל המדיניות החקיקתי המפורט ביותר מסוגו בעולם. זאת, משום שמודל מדיניות זה כולל **הסדרה רוחבית בחקיקה** של טכנולוגיות הבינה המלאכותית, באופן שהוא מגדיר את תפיסת ניהול הסיכונים הכוללת, את התבחינים לדרישות הרגולטריות, ואת הדרישות הרגולטריות.

שיטת ההסדרה המוצעת בטיטת החקיקה האמורה מבוססת על תפיסת ניהול סיכונים, שבה יש התאמה בין עומק הרגולציה לבין רמת הסיכון הנשקפת לזכויות יסוד ולבטיחות. טיוטת החקיקה שהציעה הנציבות מבחינה בין פיתוח או שימוש במערכות בינה מלאכותית אשר הסיכון הגלום בהן הוא "בלתי מקובל" ולכן אסור באופן גורף; מערכות שלגביהן "הסיכון גבוה" ולכן מצריך בחינה מראש וכפוף לשורה של דרישות מפורטות; ומערכות שלגביהן "הסיכון מוגבל" ולכן מחייב מגבלות דרישות בהיקף מצומצם יחסית.

מערכות שיוגדרו בסיכון גבוה יידרשו להירשם במרשם כללי אירופי, ולעבור בדיקה מקדימה. במידה שמדובר בבינה מלאכותית המשולבת במוצר, הבדיקה תיעשה במסגרת בדיקת המוצר. יצרני מערכות בתחום שאינו דורש בחינה מוקדמת לפי חוק כיום, יידרשו לבצע בדיקה עצמית המראה שהם עומדים בדרישות. כמו כן, מערכות אלה יצטרכו לעמוד במגוון דרישות, ובכלל זה עריכת ניהול סיכונים; בדיקת הטכנולוגיה; עמידות; משילות מידע; שקיפות; מעורבות אנושית; והגנת סייבר.

לאחרונה אישר הפרלמנט האירופי את גרסתו להצעת חוק הבינה המלאכותית.⁷⁶ השינויים המרכזיים בגרסה זו כוללים התאמה של הגדרת "בינה מלאכותית" להגדרת ה-OECD, כפי שזכר

European Parliament, European Parliamentary Research Services, Auditing the Quality of Datasets ⁷⁴ (EPRS Study 2022 Used in Algorithmic Decision-Making Systems p.3 (July 2022)).
[https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI\(2021\)698792](https://www.europarl.europa.eu/thinktank/en/document/EPRS_BRI(2021)698792)

⁷⁵ מדיניות זו כוללת ארבעה יעדי מדיניות מרכזיים: ליצור תנאים לפיתוח ושימוש בבינה מלאכותית באיחוד האירופי, להפוך את האיחוד למקום הנכון לעשייה – ממעבדה לשוק, לוודא כי טכנולוגיות בינה מלאכותית משרתות את האוכלוסייה (חקיקת הבינה המלאכותית הינה חלק מיעד מדיניות זה), ובניית הובלה אסטרטגית בסקטורים. מקור: מצגת של נציבות האיחוד האירופי.

⁷⁶ "Artificial intelligence system" (AI system) means a machine-based system that is designed to operate with varying levels of autonomy and that can, for explicit or implicit objectives, generate outputs such as predictions, recommendations, or decisions, that influence physical or virtual environments"; Artificial Intelligence Act, Amendments adopted by the European Parliament on 14 June 2023 on the proposal for a regulation of the European Parliament and of the Council on laying down harmonized rules on artificial intelligence (Artificial Intelligence Act) and amending certain Union legislative acts (COM(2021)0206 –

בפרק שעסק בהגדרת בינה מלאכותית; התייחסות לבינה מלאכותית יוצרת ולמודלי בסיסי כקטגוריה חדשה, אשר איננה נחשבת אוטומטית למערכת בסיכון גבוה אך כפופה לחלק מהכללים החלים על מערכות מסוג זה; איסור גורף על שימוש במערכות ביומטריות הפועלות במרחב הציבורי בזמן אמת; ותחולה כללית של הרגולציה על מערכות אכיפת החוק.

נכון להיום, טיוטת החקיקה האירופית מצויה במשא ומתן בין הנציבות האירופית, מועצת השרים, והפרלמנט האירופי. על פי פרסומי הפרלמנט, הליך החקיקה באיחוד האירופי צפוי להסתיים עד לתום שנת 2023 ולהיכנס לתוקף כשנה וחצי-שנתיים לאחר מכן.

3.3. מועצת אירופה וקידום אמנה בנושא בינה מלאכותית

מועצת אירופה (Council of Europe) היא ארגון של מדינות אירופה שהוקם לאחר מלחמת העולם השנייה כדי לקדם את התיאום ושיתוף הפעולה בעניין זכויות האדם, דמוקרטיה ושלטון החוק באירופה.⁷⁷ כיום הארגון מונה 46 מדינות.⁷⁸ אמנה מרכזית שגיבש הארגון היא האמנה האירופית לזכויות אדם וחירויות יסוד (The European Convention on Human Rights and Fundamental Freedoms).

בשנת 2019 מינתה מועצת השרים של מועצת אירופה ועדת אד-הוק בנושא בינה מלאכותית (CAHAI).⁷⁹ הועדה פעלה עד שנת 2021, והגישה דו"ח מפורט ובו המלצות על נושאים באמנה חדשה של מועצת אירופה בתחום הבינה המלאכותית,⁸⁰ שהועברו להכנה לאמנה על ידי ועדת ההמשך (CAI).⁸¹ זוהי היוזמה המובילה כיום לאמנה בינלאומית בתחום הבינה המלאכותית.

הוועדה מציינת את חשיבות ניסוח האמנה באופן שיאפשר הצטרפות של מדינות מחוץ לאירופה, זאת, על מנת להגביר את השפעת האמנה ולייצר כללי שוק שווים לגורמים הרלוונטיים ובהם התעשייה והאקדמיה.⁸² בהמשך לכך, הוועדה ממליצה להתחשב בעיסוק בנושא בעקרונות ה-OECD, באיחוד האירופי, וב-UNESCO (שלא יידון במסמך זה). ישראל היא מדינה משקיפה בתהליך גיבוש האמנה המתקיים כיום.⁸³

C9-0146/2021 – 2021/0106(COD)) https://www.europarl.europa.eu/doceo/document/TA-9-2023_0236_EN.pdf

⁷⁷ Council of Europe, About the Council of Europe – Overview <https://edoc.coe.int/en/an-overview>
⁷⁸ הארגון כולל את המדינות המייסדות (מערב אירופה) וכן מדינות שהצטרפו בשנות ה-90 של המאה ה-20 לאחר התפרקות הגוש הסובייטי, ראו שם.

⁷⁹ Council of Europe, CAHAI – Ad hoc Committee on Artificial Intelligence <https://www.coe.int/en/web/artificial-intelligence/cahai>

⁸⁰ Council of Europe, Ad Hoc Committee on Artificial Intelligence (CAHAI), Possible elements of a legal framework on artificial intelligence based on the Council of Europe's standards on human rights, democracy and the rule of law, Strasbourg (2021) <https://rm.coe.int/cahai-2021-09rev-elements/1680a6d90d>

⁸¹ ההמלצות אינן חלות בתחום הביטחון הלאומי שאינו במנדט של מועצת אירופה, שם, פס' 6.

⁸² שם, פס' 7.

⁸³ במסגרת עבודת הוועדה היא ביצעה סקירה של הנעשה בעולם. בסקירה זו נכלל מאמר מאת פרופסור יצחק בן ישראל, פרופסור אביתר מתניה וד"ר ליהיא פרידמן, המיזם הלאומי למערכות נבונות, לעיל ה"ש 20, הסוקרים בהרחבה את המלצות ועדת המשנה בראשות פרופסור קרין נהון, לגבי הגישה שהוצעה בידי ועדת המשנה לעקרונות אתיים ולאסדרה;

Counsel of Europe, CAHAI, Towards Regulation of AI Systems, Global perspectives on the development of a legal framework on Artificial Intelligence (AI) systems based on the Council of Europe's standards on human rights, democracy and the rule of law <https://edoc.coe.int/fr/intelligence-artificielle/9656-towards-regulation-of-ai-systems.html>

על פי הטייטה הראשונה שגיבשה מזכירות הארגון כנקודת מוצא למו"מ,⁸⁴ הסעיפים המרכזיים באמנה הם: גישה מבוססת סיכונים, חובה לבצע תהליך בחינת סיכונים כאשר המערכת עלולה לפגוע בזכויות אדם, דמוקרטיה או שלטון החוק, חובות מוגברות לממשלה (חוקיות, שיתוף ציבור, שקיפות וכדומה), והקמת רשות לאומית.

3.4. מדיניות הרגולציה בארה"ב

בשנת 2020 פורסם חוזר מטעם משרד הניהול והתקציב (Office of Management and Budget) בבית הלבן,⁸⁵ בעקבות צו נשיאותי 13859 בנושא זה (להלן: **חוזר OMB**). חוזר זה שיקף מדיניות שעיקרה בהסרת חסמים לפיתוח בינה מלאכותית ושימוש בה, תוך הגנה על טכנולוגיה אמריקאית, על הביטחון הלאומי והכלכלי האמריקאי, על פרטיות, זכויות אזרח, וערכים נוספים כגון חרות, שלטון החוק וכיבוד קניין רוחני. לצד זאת, משרד המדע שבבית הלבן החל בסדרת דיונים ציבוריים לבחינת הצורך ב"מגילת זכויות אדם דיגיטלית" לתחום הבינה המלאכותית.⁸⁶

החוזר מגדיר את העקרונות הבאים לגבי מדיניות רגולציה בהקשרי בינה מלאכותית:

- א. קידום אמון הציבור בבינה מלאכותית – על הרגולציה לשרת את אמון הציבור בטכנולוגיה באמצעות קידום אפליקציות מהימנות, עמידות וראויות לאמון.
- ב. שיתוף הציבור.
- ג. מהימנות מדעית ואיכות המידע – התפיסה הרגולטורית צריכה להיות מבוססת על מידע ותהליכים טכנולוגיים ומדעיים. הערכת סיכונים וניהול סיכונים – הפעילות של הרשויות בתחום הבינה מלאכותית הן בהקשר הרגולטורי והן בהקשר לא רגולטורי צריכות להיות מבוססות על יישום עקבי של הערכת סיכונים וניהול סיכונים.
- ד. שקלול תועלות ועלויות – בקידום רגולציה יש לבחון את מכלול התועלות והעלויות הנובעות משילוב בינה מלאכותית בפעילות, תוך התחשבות בעלותה החברתית, בתועלת הגלומה בה ובהשפעותיה החלוקתיות. יש להשוות בין המצב הקיים לבין המצב הנובע משילוב הבינה המלאכותית, ובהיעדר מקור להשוואה, יש לשקול את הסיכונים והעלויות של אי-שימוש בה.
- ה. גמישות – יש לשקול לקבוע כללים גמישים וניטרליים מבחינה טכנולוגית (כלומר, לא כללים מפורטים המתאימים רק לסוג ספציפי של מערכת), ולהימנע מכללים משפטיים נוקשים.

⁸⁴ Committee on Artificial Intelligence (CAI) Revised Zero Draft [Framework] Convention on Artificial Intelligence, Human Rights, Democracy and the Rule of Law <https://rm.coe.int/cai-2023-01-revised-zero-draft-framework-convention-public/1680aa193f>

⁸⁵ Executive Office of the President of the United States of America, Office of Management and Budget, Memorandum for the Heads of Executive Departments and Agencies, Guidance for Regulation of Artificial Intelligence Applications, M-21-06, 17.11.2020.

The White House, CYMI: WIRED (Opinion): Americans Need a Bill of Rights for an AI-Powered World, 22.10.2021 <https://www.whitehouse.gov/ostp/>

וראו עוד:
The White House, Readout of White House Listening Session on Tech Platform Accountability, 08.09.2022, "Principles for Enhancing Competition and Tech Platform Accountability With the event, the Biden-Harris Administration announced the following core principles for reform: 1. **Promote competition in the technology sector...**5. **Increase transparency about platform's algorithms and content moderation decisions.**...6. **Stop discriminatory algorithmic decision-making**" <https://www.whitehouse.gov/briefing-room/>

ו. הגינות ואי-אפליה – רשויות, במסגרת פעילותן הרגולטורית, צריכות להבהיר כיצד הן מתייחסות להשפעה של מערכות מבוססות בינה מלאכותית על הסיכון להתרחשות אפליה. בהקשר זה, יש לבחון את המצב לפני שילוב בינה מלאכותית ולאחריה.

ז. גילוי ושקיפות – גישת החוזר היא כי לגילוי ושקיפות בנוגע לשימוש בבינה מלאכותית תרומה לא רק להליך הרגולטורי אלא גם להגברת אמון הציבור. היקף הגילוי הנדרש מבוסס על הקשר, הערכת הנזקים האפשריים (כולל שימוש לרעה במידע הנמסר), היקף הנזקים, האפשרויות הטכנולוגיות ("technical state of the art"), והתועלות הפוטנציאליות משימוש.

ח. בטיחות ואבטחה.

ט. תיאום בין רשויות.

י. שימוש בכלים לא רגולטוריים – רשויות יכולות להימנע מקביעת רגולציה ייעודית לבינה מלאכותית, ותחת זאת לעשות שימוש בהנחיות קונקרטיות להסרת אי-ודאות, שימוש בניסויים או קידום סטנדרטים וולונטריים.

יא. הסרת חסמים לשימוש בבינה מלאכותית – בין האמצעים המוזכרים בחוזר בהקשר זה ניתן למנות פעילות להנגשה של מידע הנדרש עבור בינה מלאכותית, מסירת מידע לציבור על הפעילות של הרשות בתחום כדי לעודד את אמון הציבור, השתתפות בפורומים הקובעים סטנדרטים וכללי התנהגות וולונטריים ופעילות לשם כך בפורומים בינלאומיים.

בחודש אוקטובר 2022 פרסם הבית הלבן תכנית שכותרתה "AI Bill of rights – making automated systems work for the American people".⁸⁷ התכנית איננה מסמך בעל תוקף משפטי מחייב, אלא עקרונות שנועדו להנחות את התכנון והשימוש במערכות אוטונומיות, במטרה להגן על הציבור: מערכות בטוחות ויעילות, הגנה מפני אפליה, פרטיות ובעלות על מידע, גילוי והסברתיות, ומעורבות אנושית.

בחודש יולי 2023 פרסם הבית הלבן כי במטרה למתן את הסיכונים הנשקפים ממערכות בינה מלאכותית הגיע הממשל האמריקאי להסכמה עם מספר חברות מובילות בתחום הטכנולוגיה (ובהן אמזון, גוגל, מיקרוסופט, מטא ו-OPEN-AI) על מחויבויות וולונטריות שחברות אלו לוקחות על עצמן לטובת הגברת הבטיחות, הביטחון, והאמון בכל הנוגע לטכנולוגיות בינה מלאכותית.⁸⁸

בסוף חודש אוקטובר 2023 פרסם הבית הלבן צו נשיאותי בנושא בינה מלאכותית, המשקף ומעגן עקרונות מנחים לפיתוח ושימוש בטוח בבינה מלאכותית (הצו פורסם ביום 30.10.23). הצו מתייחס למגוון נושאים לרבות חדשנות, עידוד תחרות, קניין רוחני, חיזוק ההון האנושי באקדמיה ובתעשייה, עידוד מחקר, תשתיות, פיקוח על מודלי מסד (foundation models), שוק התעסוקה העתידי וזכויות עובדים, ביטחון לאומי וסייבר, זכויות סוציאליות, מערכת המשפט ואכיפת חוק, זכויות אדם, הגנת הצרכן ופרטיות. הצו קובע כללים ביחס לשימושים ממשלתיים של בינה מלאכותית, ודורש מכל משרד ממשלתי פדראלי למנות Chief AI Officer. חשוב להדגיש שהצו עצמו, אף שעוסק בסוגיות משפטיות ואסדרתיות, אינו מהווה מקור חקיקתי מחייב, אלא מתווה

⁸⁷ The White House, Blueprint for an AI Bill of Rights <https://www.whitehouse.gov/ostp/ai-bill-of-rights>.

⁸⁸ The White House, FACT SHEET: Biden-Harris Administration Secures Voluntary Commitments from Leading Artificial Intelligence Companies to Manage the Risks Posed by AI <https://www.whitehouse.gov/briefing-room>.

יעדים לתכניות עבודה רב שנתיות, רחבות וכוללניות עבור גופים פדראליים שונים, ומנחה אותם ביחס למטרות, משימות, שיטות הפעלה והנחיות בתחום.

הצו מתבסס על עקרונות בינה מלאכותית מוכרים (אמינות, אי-הפליה, פרטיות, שקיפות וכו'). בניגוד לגישה האירופאית של חקיקה רוחבית, הצו נוקט בגישה של אסדרה ענפית בכך שהוא מנחה את הרגולטורים השונים להתייחס למגוון הסוגיות ולתת את הדעת לצורך באסדרה סקטוריאלית. יצויין, כי בדומה לצו הנשיאותי הקודם, גם הפעם ה-OMB פרסם טיוטת מזכר להערות הציבור ובה הנחיות נוספות (וקונקרטיות יותר), בין השאר, ביחס לניהול סיכונים בכל הנוגע לשימושי בינה מלאכותית, כללי אחריות וכו'".ב.

3.5. מדיניות הרגולציה בבריטניה

ממשלת בריטניה פרסמה מסמך אסטרטגיה לאומית בתחום הבינה המלאכותית, אשר עודכן לאחרונה בשלהי שנת 2022.⁸⁹ נקודת המוצא למסמך היא חשיבותו של שוק הבינה המלאכותית לתעשייה ולכלכלה בכלל, והשפעתו האפשרית על כל מרחבי החיים. מסמך האסטרטגיה ומציב בפני ממשלת בריטניה שלוש מטרות:

- א. לתת מענה לצרכים ארוכי הטווח של האקו-סיסטם בתחום הבינה המלאכותית;
- ב. לתמוך במעבר לכלכלה מאפשרת בינה מלאכותית;
- ג. להבטיח שבריטניה תפעל נכון ביחס למשטר הלאומי והבינלאומי בתחום הבינה המלאכותית באופן שיעודד חדשנות והשקעות ויגן על הציבור והערכים הבסיסיים בהם הוא מחזיק.

המסמך מציע שלא לקדם בשלב זה רגולציה רוחבית בתחום הבינה המלאכותית וחרף זאת לפעול במסגרת של רגולציה ענפית, וזאת בין היתר בשל הקושי להתחשב במאפיינים המורכבים של שימושים פרטניים והשלכותיהם, והעובדה שהרגולטורים הענפיים מסוגלים לספק מענה פרטני מדויק ומהיר יותר לסיכונים הרלוונטיים.

מנגד, המסמך מדגיש את החשיבות בראיה לאומית רוחבית וקוהרנטית.

לפיכך, ממשלת בריטניה תשקול פעולות שונות כגון: הסרה של עומסים רגולטוריים, והצגת עקרונות או כללים חוצי מגזר אשר ישלימו את פעילות הרגולטורים השונים ויבטיחו עקביות גבוהה יותר.

במרץ 2023 פרסמה ממשלת בריטניה להערות הציבור טיוטת מסמך מדיניות בשם "גישה תומכת חדשנות לאסדרת בינה מלאכותית".⁹⁰ במסגרת המסמך מתוארת המחויבות ל-5 עקרונות מנחים שנועדו להבטיח פיתוח ושימוש אחראי בבינה המלאכותית, בכל המגזרים: (1) בטיחות, בטחון ואמינות; (2) שקיפות והסברתיות הולמות; (3) הוגנות; (4) אחריותיות; ו- (5) תחרותיות. כמו כן, מתוארת הגישה אשר נסקרה לעיל, לפיה בשלב הראשון אין מקום לקבוע חקיקה רוחבית ונכון ליישם את העקרונות הללו על ידי הרגולטורים הענפיים תוך שימוש במומחיותם לתחום עליו הם מופקדים. עם זאת, בזמן המתאים ייתכן שתוטל על הרגולטורים חובה לקבוע רגולציה שתיישם עקרונות אלו בתחומם.

בנוסף, המסמך מציין כי זוהו מספר תחומים בהם מסגרת רוחבית תוכל לסייע למיתון הסיכונים הנובעים משימוש בכלי בינה מלאכותית מבלי לפגוע בקידום החדשנות:

1. ניטור והערכת היעילות של הרגולטורים בכל ענף, לרבות יכולתם להטמיע את העקרונות המנחים ולקדם את החדשנות.
2. הערכת הסיכונים הכלכליים הנובעים משימוש בבינה מלאכותית.
3. תמיכה במיזמים וארגוני חול רגולטוריים שיסייעו ליזמים לבחון טכנולוגיות חדשות בשווקים.

⁸⁹ Guidance, National AI Strategy - HTML Version, Updated 18 December 2022 <https://www.gov.uk/government/publications>
⁹⁰ Policy paper, A Pro-innovation Approach to AI Regulation, March 2023 <https://assets.publishing.service.gov.uk/>

4. קידום החינוך והמודעות בקרב בעלי עסקים והציבור הרחב.
5. קידום התאימות למסגרות הרגולציה הבינ"ל.

כמו כן מוצע במסמך להיעזר בכלי רגולציה מתקדמת כגון כלי הערכת סיכונים, מדריכים שאינם מחייבים, וסטנדרטים טכניים. במדיניות זו ממשלת בריטניה סבורה שתביא לפיתוח הסביבה הרגולטורית המובילה בעולם באופן שיקדם את החדשנות והכלכלה.

3.6. מדיניות הרגולציה בקנדה

בשנת 2017 פרסמה קנדה לראשונה את "אסטרטגיית הבינה המלאכותית הפאן-קנדית";⁹¹ במטרה להפוך את קנדה למובילה עולמית בתחום הבינה המלאכותית על ידי השקעה במחקר, חינוך ובחדשנות, ובכפוף לעקרונות לשימוש אחראי בבינה מלאכותית, כגון שקיפות, אחריות ואי-אפליה.

חלק מהשפעות הבינה המלאכותית בקנדה מטופל תחת המסגרות הרגולטוריות קיימות בתחומי הגנת הצרכן, הפרטיות ועוד. במסגרת המאמצים בהם נוקטת קנדה על מנת להבטיח שהחוקים המסדירים את פעילות השוק יעמדו בקצב התפתחות הטכנולוגיה הרגולטוריים השונים פועלים לטיפול בחלק מההשפעות הבינה מלאכותית בתחומם. כך למשל, משרד הבריאות הקנדי פרסם עקרונות מנחים לפיתוח מכשירים רפואיים המבוססים על למידת מכונה,⁹² והמשרד המפקח על המוסדות הפיננסיים עובד על עדכון ההנחיות המתייחסות לניהול סיכונים כדי להתייחס לשימוש בטכנולוגיות חדשות הכוללות בינה מלאכותית.⁹³

עם זאת, ממשלת קנדה עמלה גם על מסגרת רגולטורית רוחבית להסדרת מערכות בינה מלאכותית. ביוני 2022 הגישה ממשלת קנדה את הצעת חוק הבינה המלאכותית והנתונים, Artificial Intelligence and Data Act (AIDA),⁹⁴ כחלק מתהליך חקיקה רוחבי המתייחס גם לפרטיות ולהגנת הצרכן.⁹⁵ ה-AIDA נועד להבטיח פיתוח אחראי של בינה מלאכותית במדינה, ולקדם חברות קנדיות בתחום פיתוח בינה מלאכותית בעולם. הצעת החוק נוקטת בגישה מבוססת סיכון, ומתוכננת להתאים לנורמות הבינלאומיות המתפתחות בתחום הבינה המלאכותית ובהם עקרונות ה-OECD, החקיקה באיחוד האירופי, ומסגרת ניהול הסיכונים של המכון הלאומי האמריקאי לתקנים וטכנולוגיה (NIST).

הצעת החוק מתמקדת במערכות בינה מלאכותית בעלות השפעה גבוהה על הציבור, בהתאם לאמות מידה שיוגדרו על ידי הרגולטורים, אשר ביחס אליהן יוטלו מגבלות שיבטיחו את בטיחותן ויפחיתו סיכונים כגון הטיה ואפליה. כחלק מאמות המידה לזיהוי מערכות בינה מלאכותית בעלות השפעה גבוהה על הציבור ייבחנו הסיכון לפגיעה בבריאות ובבטיחות, הסיכון להשפעה שלילית על זכויות אדם, חומרת הנזקים הפוטנציאליים, תדירות השימוש, האפשרות המעשית או המשפטית למנוע שימוש במערכת והשאלה האם הנושא מוסדר בחוק קיים. במדריך שפורסם מובהר כי הממשל הפדרלי מתכוון להקדיש תשומת לב מיוחדת למערכות שנועדו לקבל החלטות, המלצות או תחזיות הקשורות לגישה לשירותים, כגון אשראי או תעסוקה; ומערכות המשתמשות בנתונים ביומטריים כדי ליצור תחזיות לגבי אנשים, כגון זיהוי אדם מרחוק, או ביצוע תחזיות לגבי המאפיינים, הפסיכולוגיה או התנהגויות אנושיות.

הצעת החוק כוללת שורת דרישות ביחס לפיתוח מערכות בינה מלאכותית שייחשבו לבעלות השפעה גבוהה, שעיקרן: התממה (אנונימיות של מידע); הערכת סיכונים; פעילות לזיהוי ומיתון החשש מנזקים והטיות; ניטור היעילות של פעילות מיתון הנזק והתאמתה לנדרש; שמירת רישומים

⁹¹ <https://ised-isde.canada.ca/site/ai-strategy/en>

⁹² Health Canada, "Good Machine Learning Practice for Medical Device Development: Guiding Principles." Canada.ca, Oct. 27, 2021 <https://www.canada.ca/en/health-canada/services/d>

⁹³ Office of the Superintendent of Financial Institutions, "Proposed Revisions to Guideline E-23 on Model Risk Management." Canada.ca, May 20, 2022 https://www.osfi-bsif.gc.ca/Eng/fi-if/in-ai/Pages/E-23_let.aspx

⁹⁴ The Artificial Intelligence and Data Act (AIDA) <https://ised-isde.canada.ca/site/ised/>
⁹⁵ The Digital Charter Implementation Act, 2022 <https://www.parl.ca/DocumentViewer/en/44-1/bill/C-27/first-reading>

כלליים הנוגעים למילוי חובותיו לפי הצעת החוק; פרסום תיאור בנוגע לאופן פעילות המערכת, סוג התוכן שהיא יוצרת או ההחלטות וההמלצות שהיא מספקת, והצעדים שנקטו למיתון הסיכונים; וחובת הודעה לממשלת קנדה על נזק מהותי שנגרם או צפוי להיגרם כתוצאה מהשימוש במערכת.

במקביל, על מנת להבטיח שהמדיניות החדשה תהיה תואמת להתפתחויות הטכנולוגיות, נקבע שהסמכות לאכוף את החוק תהיה נתונה לשר החדשנות, המדע והתעשייה; ושיוקם בקנדה מרכז חדש בעל מומחיות בתחום הבינה המלאכותית, בראשות נציב ייעודי לתחום זה (AI and Data Commissioner), שיהיה מופקד על כלל ההיבטים הרגולטוריים והמינהליים הכרוכים בחוק. תפקיד המרכז בשלב הראשוני יהיה בתחומי ההסברה והסיוע, אולם בהמשך יחזיק גם בסמכויות אכיפה. זאת, בין היתר, באמצעות הטלת עיצומים כספיים וכן קביעת עבירות פליליות ייעודיות.

4. חלק רביעי: סוגיות ואתגרים המתעוררים בקשר לבניה המלאכותית

לצד היתרונות המשמעותיים הגלומים בו, השימוש בבניה מלאכותית מעורר סיכונים וחששות שונים, לרבות פגיעה בזכויות אדם, כפי שנוזכר לעיל. סיכונים אלה תלויים בהקשר בו נעשה השימוש בבניה מלאכותית, ובין היתר באופן היישום הטכנולוגי וסוג השימוש.⁹⁶ ההתמודדות עם סיכונים וחששות אלה עשויה לעורר סוגיות משפטיות ורגולטוריות מורכבות, אשר יש בהן כדי ליצור אתגרים משמעותיים עבור המערכת המשפטית והרגולטורית ברחבי העולם ובישראל. על רקע זה, הפרק סוקר באופן כללי סוגיות, סיכונים ואתגרים נפוצים המתעוררים בעקבות פיתוח ושימוש בבניה מלאכותית, כפי שהם משתקפים מהספרות האקדמית ומדו"חות שונים, מישראל ומהעולם. הפרק הבא יעסוק בדרכי פעולה אפשריות להתמודדות עם הסוגיות, הסיכונים והאתגרים שייסקרו בפרק זה.

מטבע הדברים, הסקירה אינה ממצה את כל הסוגיות, הסיכונים והאתגרים שמתעוררים בקשר לבניה מלאכותית, ואופן הצגתם אינו ממצה את כל שידוע ונכתב לגביהם וכן משקף זווית מבט וניתוח שעשויה להיות חלקית. בנוסף, סדר הבאת הדברים אינו בהכרח משקף חשיבות של אתגר אחד למול משנהו או שכיחותו. עם זאת, מטרת הפרק היא להציף נושאים נפוצים העולים בשיח כחלק מהדיון בנוגע למדיניות הבניה מלאכותית בישראל. הסקירה נועדה לאפשר לגורמים השונים להכיר את הדברים, למפות אותם ולעקוב אחריהם.

יובהר כי פרק זה לא עוסק בפרשנות הדין הקיים בישראל, או בדין שנכון לאמץ באמצעות כללים או פרשנות בישראל. בנוסף, אין בהצגת הסוגיות, הסיכונים והאתגרים בפרק זה כדי לקרוא לפעולה ממשלתית, רגולטורית או משפטית מיידית, שאינה היכרות ראשונית של הרגולטורים עם התחום. על אף קיומו של שיח אקדמי גובר בנושאים אלה, ופעילות רבה בתחום המדיניות, יש לזכור כי אנו מצויים בשלב ראשוני יחסית של ההתפתחות הטכנולוגית, ובהתאמה של המענים המשפטיים והרגולטוריים. ביחס לחלק ניכר מהסוגיות, הסיכונים והאתגרים שייסקרו להלן, הדיון מצוי בשלב מוקדם, וסביר להניח שככל שיחלוף הזמן הדברים יתבהרו ויתחדדו. לפיכך, בשלב זה מוצע לראות בסקירה זו מצע ראשוני לדיון ממשלתי, רגולטורי ומשפטי וכאמצעי לשיתוף הציבור בדיון זה.

בפרט, עבור גורמי הממשלה והרגולטורים הרלוונטיים, ולאור ההצעה להלן בעניין "מיפוי השימושים בבניה מלאכותית והאתגרים הנלווים להם בענפים מוסדרים", מוצע להסתייע בפרק זה עבור הבנה ומיפוי של הסוגיות, הסיכונים והאתגרים הכרוכים בשימושים קונקרטיים במערכות מבוססות בינה מלאכותית הנעשים בענף שעליו הם אמונים.

יודגש כי פרק זה נמנע מעיסוק מיוחד בשימושים של רשויות מינהליות (משרדי ממשלה, יחידות סמך, חברות ממשלתיות וכד') במערכות מבוססות בינה מלאכותית. אמנם חלק ניכר מהסוגיות, הסיכונים והאתגרים הנסקרים בו, עשוי להתעורר גם ביחס לשימושים כאמור. אולם כפי שצוין לעיל, בכל הנוגע להיבטים המשפטיים של שימושי רשויות מינהליות וגופים דו-מהותיים בבניה מלאכותית מתקיימת עבודת מטה של הגורמים המשפטיים הרלוונטיים בממשלה, בהובלת מחלקת ייעוץ וחקיקה במשרד המשפטים, תוך עמידה על הדמיון והשוני בין שימושים כאמור לבין שימושים

⁹⁶ כך, שימוש בבניה מלאכותית בחקלאות עלול לעורר סיכונים בתחום הבטיחות, בעוד ששימוש בבניה מלאכותית בתחום החלטות אשראי עלול לעורר סיכונים בתחום הפרטיות ואיסור האפליה.

במערכות מבוססות בינה מלאכותית על ידי גופים פרטיים. בהתאם, דוגמאות משימושים של גופים מנהליים מהעולם, שיובאו להלן מפעם לפעם, יהיו למטרת המחשה של תופעות מסוימות או מורכבות טכנית, אך אינן מיועדות לעסוק באופן מיוחד באתגרים הנוגעים לגופים מנהליים.

4.1. אפליה

בבתי משפט במספר מדינות בארה"ב נעשה שימוש במערכות מבוססות בינה מלאכותית להערכת רמת הסיכון לרצידיביזם (פשיעה חוזרת) של עצורים ונאשמים, כדי לסייע בהחלטות שיפוטיות, כגון החלטה על הארכת מעצר או הפניה לחלופות מאסר. במחקר משנת 2016 בנוגע למערכת מסוג זה, COMPAS (Correctional Offender Management Profiling for Alternative Sanctions), נמצא כי המערכת קבעה באופן סיסטמטי שהסיכוי של עצורים ונאשמים שחורים לרצידיביזם גבוה יותר מאשר זה של לבנים, וזאת גם כאשר לא היה בסיס לכך. למעשה נמצא כי כמות "השגיאות" של המערכת (שמשמעותן היא שהמערכת העריכה סיכוי גבוה לרצידיביזם שלא התממש בתוך תקופת הבקרה) הייתה כמעט כפולה ביחס לעצורים ונאשמים שחורים לעומת לבנים.⁹⁷

אחד החששות הנפוצים והבולטים בשיח בנוגע לשימוש במערכות מבוססות בינה מלאכותית, וכפועל יוצא מכך, אחד מההיבטים אשר יש מקום לבחון את אופן ההתמודדות עמם, הוא החשש שהפיתוח של מערכות אלו והשימוש בהן עלול להביא לאפליה אסורה, במיוחד לרעת קבוצות חלשות וקבוצות מיעוט. זאת, אף שהשימוש במערכות מבוססות בינה מלאכותית עשוי למתן את האפליה הנובעת מהטיות אנושיות ומבנים חברתיים, באמצעות קבלת ההחלטות על ידי אלגוריתם על בסיס מאפיינים אובייקטיביים ורלוונטיים, עד שיש שסבורים שניתן לחשוב על זכות אדם שמכונה תקבל החלטה בעניינו, או על קביעת רגולציה שתחייב במקרים מסוימים להשתמש בקבלת החלטות על ידי מכונה כדי להימנע מאפליה כתוצאה מהטיה אנושית.⁹⁸

הסיכון לאפליה אסורה בבינה מלאכותית נובע בעיקר מהאפשרות להטיה בנתונים המשמשים לאימון המערכת או שימוש במאפיין רגיש באחד המשתנים הנאספים באופן שעלול להביא לתוצאות מפלות, או מהאפשרות להטיות מושרשות בתחום המומחיות המשמש לאימון הבינה המלאכותית.⁹⁹ סיכונים אלה נדונו בעבר בהקשרים של שימוש ביכולות מחשוב ונתונים מתקדמות, ובעיקר "נתוני עתק" ("big data"),¹⁰⁰ וקבלת החלטות אלגוריתמיות,¹⁰¹ אולם רמת המורכבות של מערכות מבוססות בינה מלאכותית, כמפורט להלן, מעלה את החשש שסיכונים אלה יבואו לידי ביטוי ויתממשו ביתר שאת, ותוך יכולת פחותה לזהות את העיוות שנוצר את מקורו.

להלן יוצגו סיכונים נפוצים העולים בספרות האקדמית ליצירת חשש לאפליה אסורה במערכות מבוססות בינה מלאכותית, והאתגרים המשפטיים והטכנולוגיים בהתמודדות עם חשש זה.

⁹⁷ Julia Angwin, Jeff Larson, Surya Mattu & Lauren Kirchner, *Machine Bias*, PROPUBLICA (May 23, 2016).
⁹⁸ ONLY LOBEL, THE EQUALITY MACHINE – HARNESSING DIGITAL TECHNOLOGY FOR A BRIGHTER, MORE INCLUSIVE FUTURE (2022).

⁹⁹ מאפיין רגיש הינו או מאפיין כגון דת, גזע, מין, וכד', או נתון אחר שיש לו זיקה לאותו מאפיין (פרוקסי).
¹⁰⁰ Solon Barocas & Andrew D. Selbst, *Big Data's Disparate Impact*, 104 CALIF. L. REV. 671, 671 (2016).
¹⁰¹ Tal Zarsky, *Understanding Discrimination in the Scored Society*, 89 WASH. L. REV. 1 (2014); 677-87 (2016).
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2550248

¹⁰¹ ראו למשל Danielle Keats Citron & Frank Pasquale, *The Scored Society: Due Process for Automated Predictions*, 89 WASH. L. REV. 1 (2014).

4.1.1 סיכונים לאפליה אסורה במערכות מבוססות בינה מלאכותית

כפי שיפורט להלן, וכפי שעולה מהכתיבה האקדמית בתחום, תוצאה מפלה המתקבלת עקב השימוש במערכות מבוססות בינה מלאכותית עלולה להיגרם מסיבות שונות, ובהתאמה ההתמודדות עם כל אחד מגורמי האפליה שונה.

כך לדוגמה, אפשרות אחת היא אפליה מכוונת (intentional discrimination). במקרה זה מפתחי המערכת או מפעיליה מעצבים אותה באופן מכוון כך שתיווצר אפליה. מערכות מבוססות בינה מלאכותית יכולות להקל על ביצוע אפליה מכוונת בהיותן כלי אפקטיבי לזיהוי ואפיון של השתייכות אדם לקבוצות אוכלוסייה שונות גם בנסיבות בהן שיקול זה לא אמור להיות גלוי. זאת, באמצעות יכולת מערכות אלה להסיק את ההשתייכות לקבוצה מסוימת על בסיס מידע ממגוון מקורות המלמדים על כך באופן עקיף. ניתן לעשות שימוש במידע זה על מנת להפלות במתכוון כנגד (או בעד) אותה קבוצה, למשל על ידי חסימת קבוצות אוכלוסייה ספציפיות מקבלת שירותים מסוימים.¹⁰²

אפשרות שניה היא אפליה הנובעת מאופן בחירת הנתונים ועיצוב האלגוריתם המשמש לקבלת החלטות. החשש בהקשר זה הוא שעל אף שלמעצבי הבינה המלאכותית או הארגון המשתמש בה אין כוונה להפלות, אופן בחירת המשתנים (feature selection) המשמשים את המודל לקבלת ההחלטה, ועיצוב המודל, מייצרים אפליה הלכה למעשה.¹⁰³ לדוגמה, חוקרים גילו כי מערכת מבוססת בינה מלאכותית שנועדה לסייע לבתי חולים וחברות ביטוח בארה"ב לאתר מטופלים לתכנית טיפול ייעודית לחולים בסיכון גבוה, איתרה שלא במתכוון פחות מטופלים שחורים לתכנית. בדיעבד התברר כי אחד האינדיקטורים המרכזיים עליהם הסתמכה המערכת כדי לקבוע את רמת הסיכון נגע להיקף ההוצאות הרפואיות של המטופל עד אז. במחקר הסתבר כי מאחר שככלל היו למטופלים שחורים הוצאות רפואיות נמוכות יותר, הם קיבלו באופן שיטתי ציון נמוך מזה שניתן למטופלים לבנים, מה שפגע בזכאותם לתכנית הטיפול. כך, למרות שמשנתה ההוצאות הרפואיות של המטופל אינו מפלה כשלעצמו, השימוש בו יצר אפליה דה-פקטו באופן לא מכוון.¹⁰⁴

¹⁰²-Fredrik Zuiderveen Borgesius, Discrimination, Artificial Intelligence, and Algorithmic Decision-Making, Study for the Counsel of Europe, 22 (2018) <https://rm.coe.int/>; (להלן: "המועצה האירופית"); במחקרו של פרופסור בורגיוס עבור מועצת אירופה הוא מציין את הסיבות הנפוצות הבאות, שחלקן יובהרו בטקסט, "To sum up, AI decision-making can lead to discrimination in at least six ways, which relate to (i) the definition of the target variables and the class labels; (ii) the labelling and (iii) collecting of the training data; (iv) the selection of the features; (v) proxies. And (vi) **organisations could use AI systems to discriminate on purpose**. AI can also lead to other types of unfair differentiation, or to errors." פורסם מחקר מקיף מטעם מרכז המחקר של הפרלמנט האירופי, הסוקר באופן מקיף סוגים שונים של הטיות שעולות להביא לאפליה אסורה; ראו EPRS Study 2022, לעיל ה"ש 74 בעמ' 5-16.

¹⁰³ שם, בעמ' 12.
¹⁰⁴ Ziad Obermeyer, Brian Powers, Christine Vogeli & Sendhil Mullainathan, *Dissecting Racial Bias in an Algorithm Used to Manage the Health of Populations*, 366 SCIENCE 447 (2019).

אפשרות שלישית קשורה לסיכון הנובע מהטיות המשקפות אפליה קיימת בחברה, אשר מקבלת ביטוי בנתונים המשמשים לאימון מערכת הבינה המלאכותית. אפליה מסוג זה מתרחשת כאשר הנתונים עליהם המערכת "מתאמנת" או המדגם בו היא משתמשת משקפים תפיסות מוטות, דעות קדומות ואפליה קיימת בחברה.¹⁰⁵ אלגוריתם בינה מלאכותית "מתאמן" ו"לומד" על סמך הנתונים המוזנים לו. כאשר נתונים אלה משקפים אפליה קיימת, המערכת תשכפל ותנציח אפליה זו בהחלטותיה. דוגמה מוכרת נוגעת למערכת מבוססת בינה מלאכותית שהופעלה בחברה מסחרית לצורך גיוס עובדים. המערכת אומנה על בסיס דפוסי גיוס העובדים של החברה ב-10 השנים הקודמות, ובהמלצותיה היא שיקפה את הבולטות של גברים במשרות טכניות בתעשייה בתקופה זו בכך שהעדיפה באופן שיטתי גיוס של גברים למשרות אלו.¹⁰⁶ באופן דומה, נמצא בבריטניה שתוכנת מחשב שנועדה לסייע בבחינת מועמדותם של סטודנטים לרפואה שיקפה את העובדה שבעבר התקבלו ללימודי רפואה בעיקר גברים מקומיים, והפלתה לרעה נשים ומהגרים.¹⁰⁷

השימוש במערכת עלול לא רק להנציח את האפליה הקיימת אלא אף להחריף אותה כתוצאה מהיזון חוזר של המערכת (feedback loop). כך למשל, תואר החשש שמערכת המשמשת את המטרה לחיזוי פשע תעשה שימוש בנתונים היסטוריים מפלים על מנת להמליץ על אופן פריסת ניידות מטרה. מכיוון שפריסה רחבה יותר תוביל בסבירות גבוהה לעלייה בהיקף המעורבות המשטרתית באזור, לא מן הנמנע כי המערכת תלמד שצריך להגדיל אף יותר את הנוכחות המשטרתית באותם אזורים. כך למעשה תיווצר לולאה המזינה את עצמה, שתעצים את האפליה הקיימת בנתונים.¹⁰⁸

אפשרות רביעית קשורה לכך שהמערכת "מתאמנת" על מאגרי מידע שאינם ייצוגיים מספיק (sample bias).¹⁰⁹ במצבים אלו, המערכת תתקשה לנתח בצורה מדויקת את ההשפעה של אלמנטים המאפיינים קבוצות אוכלוסייה שונות על ההחלטה. דוגמה בולטת לסוג זה של אפליה מצויה בתחום המחקר הרפואי. באופן היסטורי מחקרים רפואיים בוצעו בעיקר על גברים, ולכן היקף המידע הנוגע לרפואת גברים גדול משמעותית מזה הנוגע לרפואת נשים. כתוצאה מכך, כלי אבחון רפואיים המבוססים על בינה מלאכותית נחשפו למאגר מידע המכיל בעיקר מידע הנוגע לרפואת גברים, והם עלולים לעיתים ליצור אפקט מפלה בכך שיפעלו בצורה פחות יעילה ביחס לנשים.¹¹⁰

אפשרות חמישית היא שעשויה להיווצר אפליה כתוצאה משינוי הייעוד של המערכת, הקשרה או קהל היעד שלה. במקרים אלו, המערכת פותחה מתוך הנחה כי היא תשמש לקבלת החלטות בהקשר מסוים או בנוגע לאוכלוסייה מסוימת, ולכן פעילותה תהיה מותאמת לקבלת החלטות בהקשרים אלו או ביחס לפלחי אוכלוסייה זו. לכן, ניסיון לעשות שימוש במערכת זהה ביחס לאוכלוסיות אחרות עלול לגרום לתוצאות מוטות, שנובעות מכך שהיא לא תדע לנתח בצורה מדויקת את המאפיינים השונים של ההקשר החדש או קבוצת האוכלוסייה החדשה.¹¹¹ ייתכנו אף מצבים בהם

¹⁰⁵ המועצה האירופית, לעיל ה"ש 102, בעמ' 17.

¹⁰⁶ שם, בעמ' 15.

¹⁰⁷ Frederik Zuiderveen Borgesius, *Strengthening Legal Protection Against Discrimination by Algorithms and Artificial Intelligence*, 24 INT'L J. HUM. RTS. 1572, 1575 (2020).

¹⁰⁸ Lindsey Barrett, *Reasonably Suspicious Algorithms: Predictive Policing at the United States Border*, 41 N.Y.U. REV. L. & SOC. CHANGE 327, 337 (2017).

¹⁰⁹ Kirsten Lloyd, *Bias Amplification in Artificial Intelligence Systems*, 2, CORNELL UNIVERSITY (Sep. 20, 2018), <https://arxiv.org/abs/1809.07842>.

¹¹⁰ Isabel Straw, *The Automation of Bias in Medical Artificial Intelligence (AI): Decoding the Past to Create a Better Future*, 110 ARTIFICIAL INTELLIGENCE MED. 101965, 101966 (2019).

¹¹¹ David Leslie, *Understanding artificial intelligence ethics and safety: A guide for the responsible design and implementation of AI systems in the public sector*, THE ALAN TURING INSTITUTE, 32 (Jun. 10, 2019),

המערכת "תתאמן" על אוכלוסייה ממדינה מסוימת, אולם בפועל היא תוטמע במדינה אחרת, ואף לכך עלולה להיות השפעה מפלה על תוצאותיה של המערכת. תרחיש דומה עלול לנבוע משימוש ברכיב בינה מלאכותית "גנרי", כגון מודל שפה, בהקשר ספציפי.

כאמור לעיל, התיאור של סיכונים אלה לאפליה אסורה אינו ממצה, אולם הוא נועד להציג את הצורך במודעות לאפליה אסורה עקב מאפייני הבינה המלאכותית.

4.1.2. האתגרים בהתמודדות עם אפליה במערכות מבוססות בינה מלאכותית

כמתואר לעיל, על אף ההבטחה הגלומה במערכות מבוססות בינה מלאכותית לצמצם אפליה שמקורה בהטיות המוכרות בשיקול דעת אנושי, קיים חשש כי השימוש במערכות עלול להביא דווקא להנצחת תופעות של אפליה קיימת ואולי אף להחרפתן או להתגבשות תופעות של אפליה מסוג חדש. חלק זה עוסק בקשיים המבניים, הן בצד הטכנולוגי והן בצד המשפטי, בהתמודדות עם אפקט מפלה שנוצר על ידי מערכות מבוססות בינה מלאכותית.

א. אתגרים טכנולוגיים:

על פניו, ניתן לחשוב על פתרון לחשש מאפליה באמצעות חסימה טכנית של המערכת מלהתאמן, ללמוד או להתחשב במסגרת תהליך קבלת ההחלטה בנתונים "אסורים" המבטאים שיקולים מפלים כגון מין, גזע, מוצא, דת, השקפה וכיו"ב. אולם, כפי שיפורט להלן, וכפי שתואר לעיל בקצרה, היכולת למנוע את ההכללה של מאפיינים מפלים כרוכה באתגרים לא מבוטלים.

ראשית, אף אם תהיה אפשרות לחסום את המערכת מלהתחשב באופן ישיר במאפיינים מפלים, עדיין ישנו חשש משמעותי כי המערכת תגיע לשקול אותם באמצעות התחשבות בפרמטרים אחרים אשר מתואמים עם המאפיין המפלה (proxy). כך למשל, מקום המגורים של הפרט שבעניינו מתקבלת ההחלטה (לפי כתובת או מיקוד), הוא נתון שלעיתים קרובות מעיד על שיוך קבוצתי, מאפיינים אתניים או מעמד סוציו-אקונומי. ישנו קושי אינהרנטי לאפיין מראש את כל הגורמים המקיימים קורלציה עם משתנים מפלים ולמנוע מהמערכת לשקול אותם, שכן החשש הוא שבאמצעות יכולות הניתוח הגבוהות של המערכת היא תמצא את הדרך להתחשב בקריטריונים המפלים. מעבר לכך, במקרים רבים הוצאתם גם לא תאפשר לשמר את יתרונות הבינה המלאכותית בניתוח כמויות עצומות של מידע. המערכת מסתמכת על קורלציות אלו, שפעמים רבות הן אף בלתי ניתנות לזיהוי על ידי גורמים אנושיים, על מנת לבסס ניבויים מדויקים יותר.¹¹² חשוב לציין שלעיתים למרות הרלוונטיות של סוג מסוים של נתונים מתקבלת החלטה ערכית של המוחק (בעת שהוא מעצב את החוק) שסוג כזה של נתונים לא יילקח בחשבון במסגרת ההסדר, בשל פגיעתו האפשרית באינטרס או זכות בעלות חשיבות. כך נעשה למשל בהקשר של נתוני אשראי והדבר נכון גם לענייננו.

שנית, מכיוון שבמקרים רבים קיים קושי מבני להסביר מנגנון קבלת ההחלטות של מערכות מסוימות המבוססות על בינה מלאכותית או כיצד הן הגיעו להחלטה מסוימת, יהיה קשה לאתר מקרים בהם המערכת הביאה בחשבון שיקולים מפלים, בין אם באופן ישיר ובין אם באופן עקיף, ואת המשקל שיוחס להם במסגרת קבלת ההחלטה (להרחבה ראו פרק "הסברתיות" להלן).¹¹³

Leslie, *Understanding artificial intelligence*: (להלן: https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3403301)

Any E. R. Prince & Daniel Schwarcz, *Proxy Discrimination in the Age of Artificial Intelligence and Big Data*, 105 IOWA L. REV. 1257, 1263-1265 (2020).

Bert Heinrichs, *Discrimination in the Age of Artificial Intelligence*, 37 AI & SOC'Y 143, 150-151 (2022) ¹¹³

שלישית, במקרים רבים מאגרי מידע מקיפים ומגוונים ביחס לכל קבוצות האוכלוסייה פשוט אינם בנמצא. הדבר מפחית כאמור מאפקטיביות המערכת ביחס לקבוצות אוכלוסייה בלתי מיוצגות. חשוב להבהיר כי אין מדובר בבעיה נקודתית, אלא בבעיית רוחב שנובעת מהעובדה שיש פער היסטורי באיסוף נתונים ביחס לקבוצות מסוימות באוכלוסייה. אמנם, לעיתים ניתן להתמודד עם הבעיה באמצעות איסוף נתונים משמעותי על קבוצות אלו, אולם אפשרות זו מוגבלת ולא תמיד קיים תמריץ כלכלי מספק לביצועה.¹¹⁴ תהליך איסוף הנתונים, טיוב והכנת מאגר נתונים הם לעיתים קרובות חלק הארי בעלות הפיתוח בתחום הבינה המלאכותית, והצבת דרישות בהקשר זה עשויה להוות חסם כלכלי משמעותי.

רביעית, מורכבות מיוחדת נובעת מכך שלעיתים הפרמטרים "המפלים" עצמם נדרשים לטובת יישום שמטרתו תועלת לציבור, כגון רפואה מותאמת אישית. מחלות ומצבים רפואיים עשויים להיות אופייניים לקבוצות אתניות שונות, או לחול באופן אחר על גברים ונשים, או על צעירים ובוגרים.

ב. אתגרים משפטיים:

אף אם מבחינה טכנולוגית ניתן יהיה להתגבר על האתגרים שתוארו לעיל, ישנם מספר אתגרים משפטיים שעלולים להקשות על ההתמודדות עם אפליה במערכות מבוססות בינה מלאכותית.

ראשית, כיום דיני איסור האפליה אוסרים הן על אפליה ישירה, דהיינו שימוש בקריטריונים מפלים כשלעצמם, והן על אפליה עקיפה, דהיינו פרקטיקות שיוצרות אפקט מפלה גם ללא כוונת מכוון (דוגמת שירות צבאי כתנאי לקבלה לעבודה, שמוביל לאפליה של אוכלוסיות שאינן משרתות בצבא).¹¹⁵ מבחינה משפטית, ההסתמכות על מערכות מבוססות בינה מלאכותית לקבלת החלטות עלולה לאתגר את ההכרה באפליה עקיפה. כמתואר לעיל, במקרים רבים המערכת עשויה להתבסס על משתנים בעלי ערך חיזויי הרלוונטיים מבחינה טכנולוגית לקבלת ההחלטה שהוגדרה עבורה, המגלמים למעשה אינדיקציה (proxy) לשיוך קבוצתי שאסור להשתמש בו.

אמנם על פני הדברים ובלי להביע עמדה באשר למצב המשפטי המצוי או הרצוי, נראה כי במקרים בהם הזיקה בין משתנה הפרוקסי לקריטריון המפלה היא מובהקת יותר, יהיה מדובר בשימוש אסור לפי דיני איסור אפליה. אולם, ככל שהזיקה רחוקה יותר, כך יהיה קשה לקבוע כי השימוש במשתנה גורם לאפליה עקיפה בהתאם לדיני איסור האפליה, ולפיכך אסור. במובן זה, השימוש בבינה מלאכותית עשוי לאתגר את דיני איסור אפליה, שכן המערכות המתבססות על תבניות והקשרים לא נוגעות בהכרח למאפיין המפלה, ולעיתים רחוקות ממנו, ועדיין עשויות להפלות. ייתכן כי ניתן יהיה למתן קושי זה באמצעות כלים משפטיים שונים כגון היפוך נטלי הוכחה, קביעת חזקות ועוד.

שנית, דינים שונים המגבילים שימוש במידע עלולים למנוע יצירת מאגרי נתונים מגוונים ומאוזנים, על אף שמאגרים אלה מהווים תנאי הכרחי למערכת שאינה מפלה. דינים מגבילים כאמור, עשויים להשתייך לשני סוגים. **סוג ראשון** נוגע למצבים בהם הנתונים שמבוקש לכלול במאגר יהיו מוגנים. כך למשל, דיני הגנת הפרטיות, זכויות יוצרים, חסיונות וסודות מסחר מגבילים שימוש בנתונים

¹¹⁴ להשוואה ראו (Eur. Comm'n, DATA COLLECTION IN THE FIELD OF ETHNICITY (2017).

¹¹⁵ חוק איסור הפליה במוצרים, בשירותים ובכניסה למקומות בידור ולמקומות ציבוריים, התשס"א-2000; חוק שוויון ההזדמנויות בעבודה, התשמ"ח-1988. להתייחסות לשני סוגי האפליה במשפט הישראלי ראו למשל דג"ץ 4191/97 **רקנט נ' בית-הדין הארצי לעבודה**, נד(5) 330, פס' 27-21 לפסק הדין של הנשיא ברק (2000).

מסוימים במאגרים ולכן עלולים למנוע את השימוש במידע הנדרש. לדוגמה, בניסיון ליצור אלגוריתם שיסייע בתהליך קבלתם של עובדים למקום העבודה, התבססות על נתוני עבר השייכים לחברה הקולטת בלבד עלולים לגרום להנצחת הטיות קיימות. ואולם, בהיעדר חובה רגולטורית לשיתוף מידע, יתעורר קושי להרחיב את מאגר הנתונים הקיים מגורם חיצוני שכן מידע כזה עשוי להיות מוגן באמצעות דיני הפרטיות (ביחס לעובדים) ודיני הקניין הרוחני (ביחס למעסיקים). **הסוג השני** של ההגבלות עניינו הגבלת השימוש במאגרי הנתונים עצמם. מאגרי נתונים רבים המשמשים ללמידת מכונה יהיו מוגנים בזכויות קניין רוחני, כגון סודות מסחריים או זכויות יוצרים, גם אם כל פריט מידע לא יהיה מוגן בפני עצמו. משכך, שימוש במאגר עצמו על ידי גורם אחר ללא אישור עלול לעלות כדי הפרת הזכות של בעל הקניין הרוחני.¹¹⁶ הדבר עלול להוות בעיה מכיוון שיצירת מאגר נתונים שכבר קיים בשוק משמעותה בזבוז משאבים וחוסר יעילות ואולי גם פגיעה במוצר הסופי. על רקע זה, נדרש לבצע איזון בין דרישה להרחבה וגיוון של מאגרי המידע ל"אימון" המערכת כדי לצמצם את החשש לאפליה ולדייק את המערכת, לבין האתגרים שבגיבוש וניהול מאגרי מידע נרחבים שכאלה, ובכלל זה שמירה על פרטיות המידע והגנה על סודות מסחריים וזכויות קניין רוחני.

4.2. מעורבות אנושית

אחד הציטוטים המפורסמים ביותר בסדרה "הממלכה הקטנה" ששודרה בבריטניה בתחילת המילניום הוא "The computer says no". אישה הכורעת ללדת; ילדה המגיעה לניתוח; עובדת בשיחת משוב על ביצועיה; או לקוח בסוכנות נסיעות – כולם נתקלו בפקידה המקישה באיטיות על מקלדת המחשב ומגיבה באדישות "המחשב אומר לא". הציטוט הזה, שהפך לשגור בבריטניה, נועד לבקר אנשי שירות לקוחות ופקידים ממשלתיים שמסתמכים אך ורק על המידע המופיע במחשב על מנת לקבל החלטות, לעיתים באופן שמנוגד לגמרי להיגיון הברי. ¹¹⁷ חשש זה מלווה את שילוב המחשבים בארגונים כחלק מתהליך איסוף ועיבוד מידע וייצור תחזיות או החלטות מורכבות. כעת, כ-20 שנים לאחר שהסדרה שודרה, הביטוי קם לתחייה ומשמש לתיאור החשש שהשימוש במערכות המבוססות על בינה מלאכותית בתהליכי קבלת ההחלטות עלול להוביל למציאות דומה; מציאות שבה יסתמכו כמעט באופן בלעדי על המערכות בלי מעורבות אנושית בתהליך קבלת ההחלטה.

סוגיה מרכזית שמתעוררת בקשר לשימוש במערכות המבוססות על בינה מלאכותית, נוגעת למידת המעורבות האנושית הנדרשת במסגרת תהליך קבלת החלטה הנסמך על מערכת מסוג זה.

מאפיין מרכזי שמייחד מערכות מבוססות בינה מלאכותית הוא יכולתן לפעול באופן אוטונומי המדמה בינה אנושית ובמקרים רבים מסוגל לבצע פעולות חישוב ביעילות רבה יותר מבינה אנושית. בין היתר, הן מאפשרות לספק תחזיות או למצוא קורלציות בין פרטי מידע ללא מעורבות אנושית ולעיתים רבות אף באיכות גבוהה יותר. מערכות אלה יכולות באופן עקרוני להתאים עצמן לסביבה משתנה ולאתגרים בלתי צפויים (הגם שלא תמיד באופן מלא, ראו פרק "אמינות, עמידות, אבטחה ובטיחות" להלן).¹¹⁸

Amanda Levendowski, *How Copyright Law Can Fix Artificial Intelligence's Implicit Bias Problem*, 93 ¹¹⁶ WASH. L. REV. 579 (2018).

Doa A. Elyounes, "Computer Says No!": *The Impact of Automation on the Discretionary Power of* ¹¹⁷ Public Officers, 23 VAND. J. ENT. & TECH. L. 451, 453 (2021).

Establishing a Pro-innovation Approach to Regulating AI, GOV.UK (Jul. 18, 2022), ¹¹⁸ <https://www.gov.uk/government/publications/establishing-a-pro-innovation-approach-to-regulating-ai/establishing-a-pro-innovation-approach-to-regulating-ai-policy-statement>.

העיסוק במעורבות אנושית בבינה מלאכותית נעשה בשני הקשרים מרכזיים. הקשר אחד נובע מעיקרון ערכי-מוסרי שלפיו בעניינים שיש להם השפעה (משמעותית) על אדם, ההחלטה תתקבל על ידי אדם, כגון החלטה בעלת משמעות כלכלית כבדה או החלטות שיש בהן מימד ערכי או אתי.¹¹⁹ יש שיגידו כי קבלת החלטה אוטומטית בידי מכונה בענין משמעותי עבור אדם יכולה בנסיבות מסוימות להשפילו, בכך שהיא משקפת לכאורה התייחסות מכנית "קרה" המתעלמת מיחודיותו של אותו אדם ואשר שוללת יחס אנושי נדרש. הקשר שני הוא תועלתני, ומטרתו להתמודד עם החשש מפני כשלים בטכנולוגיה עקב בעיות אינהרנטיות (כגון תכנון או ביצוע לקוי, אירוע לא צפוי, שימוש לרעה), באמצעות שילוב אדם בתהליך קבלת ההחלטות. מובן כי בין שני עקרונות אלה קיימים קשרי גומלין. מן העבר השני, וכפי שיתואר להלן, המעורבות האנושית עלולה גם לפגוע באיכות וביעילות תהליך קבלת ההחלטה, ומכאן שאינה חפה מקשיים. על כן, הצורך במעורבות אנושית ומידת המעורבות הרצויה הן סוגיות מורכבות אשר צריכות להיות מוכרעות בהתאם לנסיבות השימוש ועל בסיס הקשרים אלה.

לצורך הצגה של היבטים שונים של סוגיה זו יסקור הפרק את המאפיינים השונים של מעורבות אנושית בקבלת החלטות של בינה מלאכותית, ולאחר מכן את היתרונות במעורבות אנושית, ואת האתגרים במעורבות אנושית.

4.2.1. מהי מעורבות אנושית?

מעורבות אנושית מתייחסת לשילוב אדם בתהליך קבלת ההחלטה או הפעולה של מערכת מבוססת בינה מלאכותית שיש לה מאפיינים אוטונומיים, בשלב שבו היא פועלת.

ניתן לחלק את מידת המעורבות בתהליך לשלוש קטגוריות: בקטגוריה הראשונה נמנים מצבים בהם המערכת מקבלת את ההחלטה בעצמה, וגורמים אנושיים לא לוקחים כל חלק בתהליך קבלתה ואף לא מפקחים עליו; נהוג לכנות אותם כ-**human off the loop**. בקטגוריה השנייה מצויים מצבים שבהם המערכת מקבלת את ההחלטה בעצמה, אולם לגורם אנושי יש סמכות לפקח על ההחלטה ולהתערב במידת הצורך, כאשר הפיקוח יכול להתרחש בזמן אמת (תוך כדי ההחלטה) או בדיעבד (כגון בערעור על ההחלטה); נהוג לכנות מצבים אלו כ-**human on the loop**. הקטגוריה השלישית עוסקת במצבים שבהם נדרשת השתתפות של הגורם האנושי בדרך של קבלת ההחלטה עצמה, כך שהגורם האנושי מפעיל שיקול דעת עצמאי, אך ההחלטה מתקבלת בסיוע המערכת; נהוג לכנותם כ-**human in the loop**.

חלוקה אפשרית אחרת למידת המעורבות האנושית בתהליך קבלת החלטה מבחינה בין שישה מצבים: **אנושי (manual)** – כשהגורם האנושי מקבל את ההחלטה בעצמו ופועל ללא כל סיוע מהמערכת; **עצה (advice)** – כשהגורם האנושי מקבל את ההחלטה בעצמו, אך הוא עושה זאת תוך הסתייעות בהמלצה של המערכת; **הסכמה (consent)** – כשההחלטה מתקבלת על ידי המערכת עצמה אולם היא יכולה לצאת לפועל רק לאחר אישור של גורם אנושי; **וטו (veto)** – כשהמערכת מחליטה ופועלת בעצמה באופן אוטומטי, אולם לגורם האנושי יש הזדמנות להטיל וטו על ההחלטה של המערכת; **ערעור (appeal)** – כשהמערכת מחליטה בעצמה ופועלת באופן אוטומטי, אולם יש

¹¹⁹ ביטוי פוזיטיבי לעיקרון זה בדין הקיים בעולם הינו הוראת סעיף 22 ל-GDPR העוסק בקבלת החלטות אוטונומיות לגבי אדם, ראו Regulation 2016/679 of the European Parliament and of the Council of 27 April 2016 on the Protection of Natural Persons With Regard to the Processing of Personal Data and on the Free Movement of Such Data, and Repealing Directive 95/46/EC, 2016 O.J. (L 119) 1, 46.

למושא ההחלטה אפשרות לערער עליה בפני גורם אנושי; **אוטונומי** (autonomous) – כשהמערכת מחליטה ופועלת בעצמה בלי ליידע את הגורם האנושי ובלי שניתן לערער על ההחלטה.¹²⁰

4.2.2. יתרונות המעורבות האנושית

מעורבות אנושית עשויה לסייע, בהקשרים מסוימים, בהתמודדות עם חלק מהאתגרים הכרוכים בשימוש במערכות מבוססות בינה מלאכותית. כפי שיפורט להלן, בהתאם להקשר השימוש ולסוג הסיכונים, היא עשויה לשפר את איכות תהליך קבלת ההחלטה ולצמצם טעויות של המערכת; לשפר את מידת האחריות של מערכות; ולהגביר את הלגיטימציה של ההחלטות ולהפחית את הפגיעה בכבודו של מושא ההחלטה.¹²¹

ראשית, מעורבות אנושית (וכאמור למעורבות אנושית יכולות להיות דרגות שונות) עשויה לעיתים לשפר את האיכות והדיוק של תהליך קבלת ההחלטה ולמנוע או לצמצם טעויות בהווה ובעתיד. טיעון מקובל הוא שאנשים מבינים בצורה טובה יותר את המטרה של ההחלטה, ויש להם אינטואיציה שחסרה אצל מערכות המבוססות על בינה מלאכותית. בהתאם לכך, במקרים מסוימים, לאנשים יש יתרונות יחסיים על פני מערכות המבוססות על בינה מלאכותית ומעורבותם יכולה לשפר את איכות תהליך קבלת ההחלטה. כאשר כתוצאה מכשל בפיתוח המערכת או בשימוש בה היא מגיעה להחלטה המסתמנת כשגויה ואדם יכול היה לזהות שגיאה זו, התרומה האנושית ברורה, שכן אדם יוכל לסטות מהחלטת המערכת ולתקנה במקרה נתון ולמקרים הבאים; אולם גם כאשר המערכת פועלת כהלכה, מעורבות גורם אנושי בתהליך קבלת ההחלטה עשויה לשפרו בטווח הקצר והארוך, בין היתר, על בסיס תובנות ביחס לעתיד שאינן מצויות בנתוני העבר עליהם מסתמכת המערכת. על רקע החשש מטעויות והטיות הכרוכות בשימוש במערכות אלה (ראו פרק "אפליה" לעיל וכן פרק "אמינות, עמידות, אבטחה ובטיחות" להלן), ובמיוחד בשלבים הראשונים של הטמעת השימוש במערכות, המעורבות האנושית צפויה לסייע להתמודד עם אתגרים אלה.¹²²

שנית, הדרישה למעורבות אנושית עשויה לשפר ולהרחיב, כמו גם להביא לידי ביטוי, את האחריותיות (accountability) (ראו בעניין זה פרק "אחריותיות"); כאשר גורם אנושי מעורב בהחלטה של המערכת, ויש לו את הזמן, את המידע הרלוונטי או המומחיות ואת הסמכות לקבל את ההחלטה הסופית בעצמו או לפקח עליה, ניתן להביא באופן זה לידי ביטוי את האחריותיות של מפתחי או מפעילי הבינה המלאכותית לאופן פעולתה כמו גם לתוצאות של אותה החלטה.

שלישית, הצדקה נוספת למעורבות אנושית בתהליך קבלת ההחלטה היא שהחלטות שבהן מעורב גורם אנושי עשויות להיתפס כלגיטימיות יותר ולסייע בהגברת אמון הציבור בשימוש במערכות. מחקרים אמפיריים העלו כי אנשים סבורים שהחלטות שמתקבלות באופן אוטונומי על ידי מערכות מבוססות בינה מלאכותית, נתפסות ככאלה שמידת הלגיטימציה שלהן פחותה, בהשוואה להחלטות שמתקבלות על ידי גורמים אנושיים.¹²³ הסבר אחד לממצאים אלו, הוא שתפיסת הלגיטימיות של

Rebecca Crotoft, Margot E. Kaminski & W. Nicholson Price II, *Humans in the Loop*, 76 VAND. L. 429, 434, 437-438 (2023) (להלן: **Crotoft**); Mica R. Endsley, *The Out-of-the-Loop Performance Problem and Level of Control in Automation*, 37 HUM. FACTORS 381, 385 (1995).

Kiel Brennan-Marquez, Daniel Susser & Karen Levy, *Strange Loops: Apparent versus Actual Human Involvement in Automated Decision-Making*, 34 BERKELEY TECH. L. J. 745, 746 (2019).

Kiel Brennan-Marquez & Stephen Henderson, *Artificial Intelligence and Role-Reversible Judgment*, 109 J. CRIM. L. CRIMINOLOGY 137, 146-147 (2019); כפי שיפורט להלן, ישנם גם מקרים בהם המעורבות האנושית תפגע באיכות ההחלטה שתקבל ולכן לא ניתן לקבוע באופן גורף כי מעורבות אנושית משפרת את איכות תהליך קבלת ההחלטה (להלן: **Brennan-Marquez & Henderson**).

Ari Waldman & Kirsten Martin, *Governing Algorithmic Decisions: The Role of Decision Importance and Governance on Perceived Legitimacy of Algorithmic Decisions*, BIG DATA & SOC'Y 1 (2022).

החלטה מסוימת מושפעת, בין היתר, מאופי ההליך שהוביל לה, ומהזכויות שמוקנות למושא ההחלטה בהליך. החלטה מסוימת תיתפס כלגיטימית יותר כשהיא התקבלה לאחר הליך הוגן (Due process), תוך מתן אפשרות למושא ההחלטה להישמע ולהשפיע על ההליך ושמירה על כבודו.¹²⁴ השימוש במערכות מבוססות בינה מלאכותית לקבלת החלטות מציב אתגרים משמעותיים לאפשרות להשיג יעדים אלו. זאת, הן משום שבמקרים רבים תהליך קבלת ההחלטה או הנימוקים שבבסיסה לא יהיו שקופים למושא ההחלטה (ראו פרק "הסברתיות" להלן),¹²⁵ הן משום שיכולת מושא ההחלטה להשמיע את קולו ולהשפיע על הליך קבלת ההחלטה מוגבלת.¹²⁶ הסבר נוסף הוא שהערעור בתפיסת הלגיטימיות נובע מכך שאנשים רגילים שהחלטות בעניינם בהקשרים מסוימים מתקבלות על ידי גורמים אנושיים, ומעדיפים שהליך קבלתן יביא בחשבון גם תכונות אנושיות כמו רגישות וחמלה.¹²⁷ מעורבות אנושית, ובפרט קבלת ההחלטה הסופית על ידי גורם אנושי, מאפשרת להתמודד עם קשיים אלו. הגורם האנושי יוכל לשמוע את מושא ההחלטה, לבחון מחדש את ההחלטה ולנמק אותה (בכפוף לאילוצים הטכנולוגיים, כמפורט בפרק "הסברתיות").

רביעית, יש הטוענים כי הכפפת אנשים להחלטות שמתקבלות על ידי מערכת מבוססת בינה מלאכותית בלבד, ללא מעורבות אנושית, עלולה לגרום, בהקשרים מסוימים, לפגיעה בכבודם של הגורמים שבעניינם מתקבלת ההחלטה. הסיבה המרכזית לכך היא שכאשר החלטות מתקבלות אך ורק על ידי מערכת ישנו חשש לפגיעה במרכיב האינדיבידואלי של מושא ההחלטה. למעשה ברוב המקרים הצורה שבה המערכת מקבלת החלטה היא באמצעות מציאת מתאם סטטיסטי בין המאפיינים של האדם לבין נתוני העבר, תוך ניסיון לשבץ את המקרה והאדם שבפניה לתוך קטגוריה של מקרי עבר. מכיוון שההליך שמבצעת המערכת לא מאפשר לאינדיבידואל להפגין את הייחודיות שלו, נטען כי למעשה מדובר בראייתו של הפרט כסט תכונות בלי לבחון את מכלול הזהות האנושית שלו באופן פרטני.¹²⁸ מעורבות גורם אנושי יכולה לספק למושא ההחלטה הזדמנות להביע את עצמו, ולוודא כי ההחלטה תתקבל על בסיס בחינת האדם כמכלול ולא כסט תכונות.

4.2.3. אתגרים ביחס למעורבות אנושית

לצד היתרונות המשמעותיים שניתן להשיג באמצעות מעורבות אנושית, שעיקרם הוצגו לעיל, החלה של דרישה למעורבות אנושית אינה מובנת מאליה. זאת, שכן כפי שיפורט להלן, קיים חשש כי במקרים רבים המעורבות האנושית לא תהיה אפקטיבית ואף עלולה להסב נזק.

ראשית, הציפייה שאנשים יוכלו להתערב באופן אפקטיבי, לפקח על פעילות המערכת ולתקן את טעויותיה – בפרט כשמדובר בטעויות שאינן נורמטיביות, אלא פרקטיות – אינה בהכרח ריאלית בהקשרים מסוימים. עצם השימוש במערכות הבינה המלאכותית מתבסס על ההערכה שהן מסוגלות בהקשרים מסוימים לחזות בצורה טובה יותר מבני אדם מהי ההחלטה הנכונה. על כן, ההנחה שאנשים יוכלו באותם הקשרים לשפוט בצורה טובה יותר מהמערכת האם ההחלטה שהתקבלה היא נכונה, אינה מתבקשת. הדבר נכון במיוחד ביחס למקרים שבהם אין הסבר מלא כיצד המערכת הגיעה להחלטה (ראו פרק "הסברתיות" להלן), אך רלוונטי גם למקרים אחרים

¹²⁴ Ari Ezra Waldman, *Algorithmic Legitimacy*, in THE LAW OF ALGORITHM 107, 110 (Woodrow Barfield eds., 2020); Aziz Z. Huq, *A Right to Human Decision*, 106 VA. L. REV. 611, 656-671 (2020).

¹²⁵ Brennan-Marquez & Henderson, לעיל ה"ש 122, בעמ' 148; Crotoft, לעיל ה"ש 120, בעמ' 55-54.

¹²⁶ Ric Simmons, *Big Data and Procedural Justice: Legitimizing Algorithms in the Criminal Justice System*, 15 OHIO ST. J. CRIM. L. 573, 579-580 (2018).

¹²⁷ שם, בעמ' 573-574.
¹²⁸ Margot E. Kaminski, *Binary Governance: Lessons From the GDPR's Approach to Algorithmic Accountability*, 92 S. CAL. L. REV. 1529, 1541-1545 (2019).

שבהם ככלל המערכת מתבססת על מספר עצום של נתונים ועל האינטראקציות ביניהם כדי לקבל החלטה, באופן שהמוח האנושי מתקשה לבצע.¹²⁹ אם הגורם האנושי המעורב לא יודע את הנימוקים להחלטה, או לא מבין כראוי את מכלול הנתונים והאינטראקציה שהובילו לקבלת ההחלטה, הוא עלול להתקשות בתיקון החלטות שגויות של המערכת.

שנית, לא אחת השימוש במערכות מבוססות בינה מלאכותית נעשה במטרה להתמודד עם העובדה שתהליך קבלת ההחלטות האנושי אינו אופטימלי. בפרט, הגם שמתעוררים חששות לא מבוטלים לאפליה בשימוש במערכות מבוססות בינה מלאכותית (ראו פרק "אפליה" לעיל), יש לזכור שמנגד קיימים חששות מקבילים ביחס לאפליה על ידי גורמים אנושיים. בנוסף, מחקרים מתחום הכלכלה ההתנהגותית מלמדים כי ישנן הטיות רבות שמפיעות על תהליך קבלת ההחלטות האנושי, וגורמות לכך שחלק ניכר מההחלטות שמתקבלות על ידי אדם בסופו של דבר אינן אופטימליות. על-כן, במקרים מסוימים, הותרת ההחלטה הסופית בידי הגורם האנושי עלולה לרוקן מתוכן את אחד היתרונות המשמעותיים שניתן להשיג באמצעות השימוש במערכות מבוססות בינה מלאכותית והוא היכולת לנטרל, למצער בצורה חלקית, אפליה והטיות אלה.¹³⁰

שלישית, ישנה שאלה כבדת משקל בנוגע למשמעות המעשית של מעורבות אנושית בתהליך קבלת החלטה, כאשר ניתנת המלצה ברורה של מערכת מבוססת בינה מלאכותית. מחקרים מתחום הכלכלה ההתנהגותית הראו כי ישנן הטיות ייחודיות המאפיינות את האינטראקציה בין מקבלי החלטות אנושיים לבין מערכות המבוססות על בינה מלאכותית. אחת ההטיות המרכזיות לעניין זה היא הטיית האוטומציה (automation bias), במסגרתה אנשים נוטים לאמץ את המסקנה אליה הגיעה מערכת אוטומטית (ובכלל זה מערכת מבוססת בינה מלאכותית), בלי להפעיל את אותה מידה של שיקול דעת כמו ביחס להחלטות שלא התקבלו על ידי מערכת. כתוצאה מכך, אנשים נוטים לפעול בהתאם להמלצות האלגוריתם גם כשיש ראיות סותרות, בגלל תפיסה כללית ש"האלגוריתם צודק" (the machine knows), וכפועל יוצא מתקשים בתיקון טעויות הנעשות על ידי המערכת.¹³¹ מצדו השני של המתרס, ישנם ממצאים שלפיהם אנשים נוטים להימנע באופן לא רציונלי מאימוץ ההחלטה שהתקבלה על ידי מערכת, ומעדיפים את שיקול דעתם גם בהיעדר הצדקה ברורה לכך. הטיה זו, ההופכית לקודמת שצוינה, מכונה בספרות סלידה אלגוריתמית (algorithm aversion).¹³² על רקע זה, ישנו חשש כי אותם גורמים אנושיים המעורבים בתהליך יתקשו להעריך בצורה נכונה את טיבה של ההחלטה שהתקבלה על ידי המערכת, וכתוצאה מכך יטעו בהחלטה אם לאמץ את המסקנה אליה הגיעה המערכת או לסטות ממנה.¹³³ כך שבפועל ספק אם המעורבות האנושית תהיה אפקטיבית בתיקון טעויות, והיא אף עלולה לגרום מאיכות תהליך קבלת ההחלטות. נושא יצירת

Adrien Bibal, Michael Lognoul, Alexandre de Stree & Benoît Frénay, *Legal Requirements on Explainability in Machine Learning*, 29 ARTIFICIAL INTELLIGENCE & L. 149, 156 (2021) (להלן: **Bibal**, **Lognoul, de Stree & Frénay**).

Jeremy Waldron, It's All for Your Own Good, THE NEW YORKER REVIEW (Oct. 9, 2014),¹³⁰ <https://www.nybooks.com/articles/2015/05/27/machine-learning-and-the-law/>; איל זמיר ודורון טייכמן "ניתוח התנהגותי של החלטות שיפוטיות: הישגים ואתגרים" **משפט ועסקים** 57 (2015).

Marina Chugunova & Daniela Sele, *We and It: An Interdisciplinary Review of the Experimental Evidence on How Humans Interact with Machines*, 99 J. BEHAV. EXPERIMENTAL ECON. 1, 19-20 (2022).

Hasan Mahmud, A.K.M. NajmulIslam, Syed Ishtiaque Ahmed & Kari Smolander, *What Influences Algorithmic Decision-Making? A Systematic Literature Review on Algorithm Aversion*, 175 TECH. FORECASTING AND SOC. CHANGE 121390, 121390-121391 (2022).

Ben Green, *The Flaws of Policies Requiring Human Oversight of Government Algorithms*, 45 **Computer L. Security Rev.** 1, 14-18 (2022) (להלן: **Green**).

ממשק עבודה אפקטיבי בין אדם ומכונה באופן שימצה את היכולות השונות לכדי החלטה מיטבית נמצא במוקד של מחקר רב באקדמיה ובחברות פרטיות שונות.

רביעית, ובהמשך לדיון בנקודה הקודמת, קיים חשש שגם במקרים המתאימים וללא הטיות ייחודיות האדם שמעורב בתהליך קבלת ההחלטה לא יסטה מהמלצות המערכת, מפאת החשש לנשיאה באחריות במקרה של טעות בשיקול דעתו (ראו גם פרק "אחריות" להלן). כך למשל, רופא עלול לחשוש לסטות מהמלצה של המערכת בנוגע לטיפול רפואי מומלץ על ידי המערכת מתוך חשש כי המטופל לא יגיב טוב לטיפול החלופי ויטען שהרופא פעל ברשלנות. באופן דומה, פקיד בבנק עלול לחשוש במיוחד לאשר מתן אשראי למי שהמערכת המליצה שלא לאשר לו, למרות שהפקיד סבור שיש הצדקה לאשר לו, וזאת בגין החשש כי במקרה שיהיה קושי בפירעון הוא יאלץ לשאת באחריות בגין התוצאה השלילית מכיוון שסטה מהמלצת המערכת. חשש זה מהווה תמריץ משמעותי עבור הגורם האנושי המעורב להיצמד להמלצות המערכת ולא לסטות מהן.

חמישית, קיים חשש שהצורך במעורבות אנושית יפחית באופן משמעותי את התועלת הנובעת משימוש במערכות מבוססות בינה מלאכותית. אחת המטרות המרכזיות של השימוש במערכות אלה, הוא האופן שבו ניתן לייעל את תהליך קבלת ההחלטות באמצעותן. ככלל, מערכות אלה מסוגלות לפעול לאורך כל שעות היממה בהיקף נרחב (scale), הן לא "מתעייפות" ועלותן בטווח הארוך עשויה להיות נמוכה מתשלום לעובדים.¹³⁴ בהתאם, הדרישה למעורבות אנושית – שמחייבת להמשיך ולהעמיד תשומות אנושיות על מגבלותיהן ועלויותיהן – עלולה לפגוע ביעילות המושגת באמצעות השימוש במערכות אלה, וככל שתהיה רחבה יותר, הפגיעה ביעילות עלולה להיות משמעותית יותר.¹³⁵ לשם המחשה, מערכת בינה מלאכותית בתחום רפואי (כמו לדוגמה דימות), עשויה להיות מסוגלת לפענח תוצאות בדיקות רפואיות במהירות גבוהה, ובכך לסייע להנגשת בדיקות לכלל האוכלוסייה בזול ובמהירות. קביעה שיש הכרח לבחון את המלצות המערכת על ידי רופא תאט את קצב הבדיקות ותעלה את מחירן.

שישית, חשש נוסף הוא שההנחה כי אנשים מסוגלים לפקח באופן אפקטיבי על המערכת תגרום לתחושה לא מוצדקת של ביטחון. השימוש במערכות מבוססות בינה מלאכותית עלול להוביל לתוצאות לא רצויות במקרים שבהם המערכת אינה מושלמת, סובלת מהטיות וכדומה. ההנחה כי אנשים מסוגלים לפקח באופן אפקטיבי על המערכות באמצעות מנגנון המעורבות האנושית, יכולה לשמש כתמריץ לא מוצדק וגדול מדי לאימוץ של מערכות מבוססות בינה מלאכותית בשלב מוקדם מדי או במקרים שבהם הנזק שצפוי להיגרם כתוצאה מטעות בהחלטה של המערכת גבוה.¹³⁶ זאת בעוד שאלמלא מעורבות אנושית, שכאמור אינה יכולה בהכרח לתת מענה לחסרונותיה של המערכת, ייתכן שהשימוש היה נמנע, נדחה, מוגבל או מפקח יותר.

לבסוף, יש חשש כי המעורבות האנושית תוביל לכך שהאחריות במקרה של טעות תעבור מכתפי מקבלי החלטות שהחליטו להשתמש במערכת או מפתחי המערכת, לגורם האנושי המעורב בתהליך קבלת ההחלטה, גם כשהדבר אינו מוצדק. למשל, במקרה שבו מקום עבודה ימנע מהעסקת עובדים מקבוצת אוכלוסייה מסוימת, יתכן כי הנטייה תהיה להאשים את הגורם האנושי המעורב גם אם המערכת באופן שיטתי דירגה מועמדים מאותה קבוצה ככאלה שמתאימים למקום העבודה במידה

¹³⁴ Dominique Hogan-Doran, *Computer Says "No": Automation, Algorithms and Artificial Intelligence* 13 JUD. L. REV. 1, 3-14 (2018).
¹³⁵ Green, לעיל ה"ש 133 בעמ' 21-23.
¹³⁶ שם, בעמ' 18-21.

פחותה; זאת, גם אם לא בוצעו בדיקות נאותות מקובלות לצמצום החשש לאפליה או שלא הקפידו שהמערכת תוכל להסביר את החלטותיה וכפועל יוצא מידת השליטה של הגורם האנושי המעורב בתהליך הייתה חלקית.¹³⁷ יודגש כי לאו דווקא מדובר באחריות משפטית (פלילית או נזיקית), ובין היתר למשל, ניתן לעשות שימוש בגורמים אנושיים כדי לשמור על תדמית המערכת באמצעות ניסיון להראות כי הכשל נגרם בגלל הגורם האנושי שהיה מעורב בתהליך, ולא בגלל תקלה במערכת.¹³⁸

4.3. הסברתיות

בשנת 2016 חוקרים מאוניברסיטה בארה"ב פיתחו מערכת מבוססת בינה מלאכותית שיועדה להבחין בין כלבים מסוג "האסקי" לבין זאבים, כחלק מניסוי. כאשר נבחנה הפעילות של המערכת ביחס לנתונים עליהם היא "התאמנה" היא הפגינה ביצועים טובים, אולם כאשר נאלצה להתמודד עם דוגמאות אחרות הביצועים של המערכת נפגעו בצורה משמעותית. התברר שבמקום להתבסס על קריטריונים כמו גודל, צבע פרווה ומאפיינים פיזיולוגיים נוספים על מנת להבחין בין הכלבים לזאבים, הקריטריון המרכזי שעליו התבססה המערכת היה אם הופיע ברקע של התמונות שלג או שלא. תוצאות אלה, הגם שהמערכת פותחה במסגרת ניסוי, ממחישות כיצד היכולת להסביר איך המערכת פועלת יכולה לסייע להבין מתי השימוש בה עלול לגרום לטעויות ואי-דיוקים. במקרים בהם יהיה מדובר בהחלטות משמעותיות יותר, לטעויות מסוג זה יכולה להיות השלכה משמעותית שאולי לא הייתה מתגלה, או שהייתה מתגלה באיחור רב.

4.3.1. על הסברתיות ותופעת "הקופסא שחורה"

כמתואר לעיל, בינה מלאכותית מבוססת על יכולות חישוביות היוצרות מודל תחזיות וקבלת החלטות שלעיתים לא ניתן "לחלץ" מתוך המערכת. הדבר נובע מכך שבשונה מאלגוריתם רגיל, בבינה מלאכותית אופן "אימון" המערכת מוביל לכך שהלוגיקה התפעולית נגזרת מתהליך האימון ומנתוני האימון ואינה מוגדרת באופן מפורש על ידי המפתח. במקרים אלו, גם מעצבי המערכת לא בהכרח יכולים להתחקות אחר החלטה פרטנית שלה. יכולות טכנולוגיות מתקדמות אלה מחריפות תופעה המתוארת כ-"קופסא שחורה" (black box) – מושג שנועד להמחיש מצב בו המידע המוזן למערכת (הקלט) ידוע, וכך גם התוצאה אליה הגיעה (הפלט); אולם דרך פעולת המערכת הפנימית אינה ידועה למי שמושפע ממנה, או אף אינה ניתנת להבנה פשוטה.

המשמעות של קבלת החלטה שלא ניתן להבין או להסביר על ידי מכונה, עלולה לעורר חשש לפגיעה בכבוד האדם ובאוטונומיה שלו, בכך שהוא לכאורה נתון להחלטה שרירותית, כמו גם פגיעה נסתרת בשוויון. יתר על כן, קבלת החלטה בדרך זו מונעת, או לפחות מקשה, על האדם הנפגע להיאבק בהחלטה שנתקבלה לרעתו ובכך עלולה להיגרם פגיעה במעגל נוסף של זכויותיו. במקרים אלה מתעוררת שאלה כללית והיא באילו מצבים ראוי לדרוש כי יינתן הסבר ביחס לאופן פעילות

¹³⁷ להרחבה ראו, Madeleine Clare Elish, *Moral Crumple Zones: Cautionary Tales in Human-Robot Interaction*, 5 ENGAGING SC. TECH. & SOC'Y. 40 (2019). כך לדוגמה, יש הטוענים כי הנטייה להסיק את האחריות לגורם האנושי שהיה מעורב בתהליך על חשבון גורמים אחרים באה כבר לידי ביטוי בתחום התעופה. שם, הנטייה היא להאשים את הטייס האנושי ולא את המפתחים של מערכת הטייס האוטומטי או המעצבים של הממשק שלה, גם במקרים בהם מידת השליטה של הטייס הייתה חלקית. לעניין זה, ולאופן שבו הוא צפוי להשפיע גם על השימוש במכונות אוטונומיות ראו Tim Hwang & Madeleine Clare Elish, *Praise the Machine! Punish the Human! The Contradictory History of Accountability in Automated Aviation*, DATA & SOCIETY RESEARCH INSTITUTE (May. 18, 2015), https://papers.ssrn.com/sol3/papers.cfm?abstract_id=2720477.
¹³⁸ Green, לעיל ה"ש 133, בעמ' 21-23.

המערכת ובעיקר מה צריכה להיות רמת הפירוט של ההסבר שניתן בכל תרחיש רלבנטי. דרישה זו מכונה הסברתיות (explainability; לעיתים מתייחסים אליה גם כ-interpretability).

הסברתיות היא היכולת להציג בצורה שניתנת להבנה על ידי בני אדם את אופן פעולת המערכת או ההחלטה שלה.¹³⁹ הגם שתכונה זו מאפיינת חלק מהמערכות, היא לא תמיד מתקיימת. יש מערכות שמטבען ניתנות להסברה בקלות (explainable), למשל כאלה המשתמשות במודלים פשוטים יחסית על מנת לקבל החלטה או כאשר ניתן להציג באופן פשוט את התהליך שביצעה המערכת. עם זאת, כשהמערכת מבוססת על מודל מורכב יותר, מתבססת על מקורות מידע מרובים, או "לומדת" ומשנה באופן תדיר את אופן פעילותה, קיים קושי להסביר את התהליך שביצעה המערכת.¹⁴⁰ ביחס למערכות מסוג זה, לעיתים נדרש תוסף או אמצעי חיצוני למערכת, שפעמים רבות יפותח במקביל לה ויבוא על גביה, אשר יהא בכוחו להסביר את תהליך קבלת ההחלטה.¹⁴¹

אפשרות אחת לחלוקה של הדרישה להסברתיות, מבחינה בין התחומים השונים שלגביהם נדרש לספק הסבר. כך, במסגרת סקירה שחברה על ידי מכון אלן טיורינג,¹⁴² הוצע לחלק את דרישת ההסברתיות לשישה תחומים שהדרישה להסברתיות עשויה להיות רלוונטית אליהם. ראשית, ניתן לדרוש הסברתיות ביחס לנימוקים (rational explanation), משמע לתת הסבר לגבי הסיבות שהובילו לכך שהתקבלה ההחלטה הקונקרטית על ידי המערכת, בצורה שתהיה נגישה. שנית, ניתן לדרוש הסברתיות ביחס לאחריות (responsibility explanation), משמע להסביר מי לקח חלק בפיתוח המערכת, איך הטמיעו את המערכת, מי מנהל אותה ולמי ניתן לפנות כדי לערער על ההחלטה שהתקבלה. שלישית, ניתן לדרוש הסברתיות ביחס למידע (data explanation), קרי להסביר מה המידע עליו התבססה המערכת בהחלטה קונקרטית ואיך היא השתמשה בו, הן ביחס למידע שהפרט שבעניינו התקבלה החלטה סיפק והן ביחס למידע חיצוני שהמערכת נחשפה אליו. רביעית, ניתן לדרוש הסברתיות ביחס להוגנות (fairness explanation), כלומר להסביר מהם הצעדים שננקטו בתכנון המערכת ובהטמעה שלה על מנת לוודא כי ההחלטות שלה, ככלל, אינן מוטות, הוגנות ומובילות לתוצאות שוויוניות. חמישית, ניתן לדרוש הסברתיות ביחס לבטיחות ולביצועים (Safety and performance explanation), שמשמעותה הסבר לגבי הצעדים שננקטו על מנת למקסם את הדיוק, האמינות, הבטיחות והחוסן של המערכת. שישית, ניתן לדרוש הסברתיות ביחס להשפעה (impact explanation), משמע להסביר מה הצעדים שננקטו על מנת לפקח על ההשפעות של השימוש במערכת וההחלטות שלה על האינדיבידואל ועל החברה בכללותה.¹⁴³

אפשרות נוספת לחלוקה של הדרישה להסברתיות, מבחינה בין מספר קטגוריות: (1) מתן הסבר כללי ביחס לפרמטרים המרכזיים עליהם בנוי המודל, אופן פעילות המערכת, ובאיזה שלב מעורב גורם אנושי (אם בכלל), בלי להתייחס להחלטה ספציפית; (2) מתן הסבר פרטני ביחס לפרמטרים

¹³⁹ Gabriel Nicholas, *Explaining Algorithmic Decisions*, 4 GEO. L. REV. 711, 715 (2020) (להלן: Nicholas).
¹⁴⁰ נמחיש את ההבדל באמצעות דוגמה פשוטה של מערכת שמטרתה להחליט אם לאשר הלוואה. הפיתוח של מערכת פשוטה יבוצע על בסיס כללים מוגדרים – לדוגמה ייקבע כי ההלוואה תאושר אם למבקש אין הלוואות נוספות, סכום ההלוואה אינו עולה על המשכורת החודשית ויש למבקש הון עצמי מינימלי של 10,000 ₪; המדובר בכללים שניתן להסביר. לעומת זאת, מערכות מורכבות המבוססות על למידה מכונה או רשת עצבית מלאכותית מסוגלות לבנות בעצמן את הכללים לאישור הלוואה מבלי להתבסס על מודל מבנה.

¹⁴¹ Bibal, Lognoul, de Streel & Frénay, לעיל ה"ש 129, בעמ' 157-159.

¹⁴² מכון טיורינג הוא מכון מחקרי בתחומי מדע הנתונים ובינה מלאכותית שהוקם על ידי מספר אוניברסיטאות בבריטניה ונתמך על ידי הממשלה, ראו אתר האינטרנט של המכון בכתובת: <https://www.turing.ac.uk/>.

¹⁴³ David Leslie, *Explaining Decisions Made with AI*, THE ALAN TURING INSTITUTE, 20 (May. 1, 2020), https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4033308.

הרלוונטיים שנלקחו בחשבון בגיבוש החלטה ספציפית; (3) מתן הסבר פרטני כיצד האינטראקציה בין הפרמטרים שנלקחו בחשבון השפיעה על ההחלטה הספציפית (ה"לוגיקה" שבבסיס ההחלטה).

4.3.2. יתרונות וחסרונות הדרישה להסברתיות

להסברתיות, בין שנעשית באופן וולונטרי ובין שנדרשת מכוח הדין, מספר יתרונות עיקריים:

ראשית, ההסברתיות מסוגלת להקל על זיהוי ותיקון טעויות של המערכת כשהן מתרחשות. כאשר המערכת ניתנת להסברה, מתאפשר להבין בצורה פשוטה ומהירה יותר מתי היא מגיעה לתוצאות שגויות ומה גורם לכך, ובהתאם לדרוש את תיקונה או את הפעלתה תחת פיקוח של גורם אנושי שיוכל לסטות במידת הצורך מהחלטה שקיבלה המערכת (ראו פרק "מעורבות אנושית" לעיל).

שנית, ההסברתיות עשויה לסייע לשפר את תהליך קבלת החלטות של הגורמים המשתמשים במערכת. למשל, רופא שנעזר במערכת בינה מלאכותית לשם אבחון יוכל לקבל ממנה מידע שיסייע לו בעתיד לקבל החלטות טובות יותר גם כאשר הוא לא נעזר במערכת.¹⁴⁴ בהקשר זה, הסברתיות רלוונטית גם ליחסים שבין מפתח מערכת המבוססת על בינה מלאכותית לבין הארגון המשתמש בה, שבהם עשוי להיות פער מידע אודות מאפיינים אלה. סגירת פער מידע זה מאפשר לארגון המשתמש להבין טוב יותר את אופן השימוש בטכנולוגיה ואת חלוקת הסיכונים הקשורה בה, ולעמוד מצדו בחובות החלות עליו כלפי מי שמושפע מהחלטותיו.

שלישית, במקרים רבים בלי הסברתיות לא תהיה אפשרות אפקטיבית להשתמש בכלים משפטיים על מנת לפקח אחר השימוש במערכות בינה מלאכותית. מבחינת מושא ההחלטה, בלי לקבל הסבר ביחס לסיבות שהובילו להחלטה, הוא יתקשה לדעת שהשימוש במערכת הביא להפרה של חובה בדין, להוכיח זאת ולנקוט בצעדים, ובכללם צעדים משפטיים, בהתאם. כך למשל, צרכן שהבנק סירב להעניק לו אשראי יתקשה להראות כי הסירוב היה בלתי-סביר בנסיבות העניין.¹⁴⁵ באופן דומה, במקרים של אפליה יהיה קשה לנקוט בהליכים לפי חוק איסור הפליה בשירותים ומוצרים.¹⁴⁶ מבחינת הרגולטורים, ללא הסברתיות יהיה קושי לוודא כי הגורמים המפוקחים עומדים בדרישות החוק,¹⁴⁷ למשל בהקשר של דיני איסור אפליה.¹⁴⁸ לבסוף, מבחינת בתי המשפט, ללא קבלת הסבר הולם על אופן פעולת המערכת, יתקשו השופטים לבחון את השימוש במערכות בינה מלאכותית ואת תוצאותיו בהיעדר מידע על טיב השימוש, הגורם לתוצאות והסיבה שהחליטה כפי שהחליטה.¹⁴⁹

רביעית, ההסברתיות עשויה לשפר את היכולת של מושאי ההחלטה להתאים את ההתנהגות שלהם. כך, צרכן שמערכת המבוססת על בינה מלאכותית סירבה לתת לו אשראי יוכל להבין את הנימוקים שהובילו להחלטתה (למשל, פיגור מסוים בהחזרי הלוואה קודמת), ובהתאם להחליט אם לשנות את ההתנהגותו כך שיוכל לקבל בעתיד שירות, לפנות למתחרה או לוותר.¹⁵⁰

¹⁴⁴ Nicholas, לעיל ה"ש 139, בעמ' 717.

¹⁴⁵ ס' 2 לחוק הבנקאות (שירות ללקוח), התשמ"א-1981; Bibal, Lognoul, de Streel & Frénay, לעיל ה"ש 129, בעמ' 156-153.

¹⁴⁶ חוק איסור הפליה במוצרים, בשירותים ובכניסה למקומות בידור ולמקומות ציבוריים, התשס"א-2000.

¹⁴⁷ William Magnuson, *Artificial Financial Intelligence*, 10 HARV. BUS. L. REV. 337, 353-354 (2020);

Bibal, Lognoul, de Streel & Frénay, לעיל ה"ש 129, בעמ' 153-151.

¹⁴⁸ אחיעז, חמדני, עמירם וקסטיאל, לעיל ה"ש 25, בעמ' 20.

¹⁴⁹ Ashley Deeks, *The Judicial Demand for Explainable Artificial Intelligence*, 119 COLUM. L. REV. 1829 (2019) (להלן: Deeks).

¹⁵⁰ Carlos Zednik, *Solving the Black Box Problem: A Normative Framework for Explainable Artificial Intelligence*, 34 PHIL. & TECH. 265, 274 (2021).

לבסוף, ההסברות צפויה להגביר את אמון הציבור ולהשפיע על נכונותו להשתמש במערכות בינה מלאכותית. הדבר יבוא לידי ביטוי הן ביחס לגורמים המשתמשים במערכת אשר ייטו לסמוך על ההחלטה שלה, ובהתאם לאמץ אותה כשזה נדרש,¹⁵¹ והן ביחס למושאי ההחלטה של המערכת שכל הנראה יהיו נכונים יותר לאפשר למערכת להחליט בעניינם, ולקבל עליהם את ההחלטה.¹⁵²

על אף יתרונותיה, בהקשרים שונים קיימת מורכבות בהחלט הדרישה להסברות באופן רחב. אמנם קיימות ומפותחות כיום טכנולוגיות מתקדמות שייעודן להסביר החלטות שהתקבלו על ידי מערכות מבוססות בינה מלאכותית, ולמעשה לפענח את אותה "קופסא שחורה".¹⁵³ ואולם, כפי שתואר לעיל, לא כל מערכת בינה מלאכותית ניתנת להסברה בקלות, אם בכלל. ככל שהמערכת מורכבת יותר, מסתמכת על נתונים רבים יותר ומבצעת קישורים רחבים יותר, כך גובר הקושי הטכנולוגי בהסברת התוצאה אליה הגיעה. על כן, אחד החששות המרכזיים שנוגעים להחלט דרישה להסברות הוא שקיום הדרישה באופן רחב כתנאי סף לשימוש בטכנולוגיה עלול להגביל את השימוש במערכות מתקדמות ומדויקות יותר. ככל שהמערכת מתחשבת ביותר סוגי פריטי מידע (ולעיתים מדובר במאות ואלפי סוגי פריטי מידע) כך גדל הקושי בהסברות מלאה שלה. מכאן שעשוי להתקיים מתח (trade-off) בין איכות המערכת, לבין היכולת להסביר את החלטותיה ופעילותה. למשל, כיום יש סוגים של מערכות מבוססות בינה מלאכותית שמפאת מורכבותן והמודל המיושם באמצעותן, יש קושי במתן הסבר מהסוג שהוצג לעיל להחלטותיהן.¹⁵⁴

בהמשך לאמור, הדרישה להסברות מעוררת באופן בלתי נמנע את השאלה למי נועד ההסבר – האם ההסבר צריך להיות מובן ונהיר לאנשי מקצוע בתחום המחשוב או בתחום הפעילות של כל יישום ספציפי, או שעליו להיות מובן לכל משתמש קצה.

זאת ועוד, הצבת דרישה להסברות באופן רחב וגורף עלולה לפגוע בהתקדמות הטכנולוגית. כשמדובר במערכות מורכבות, אפילו אם ניתן יהיה בסופו של דבר להסביר את האופן בו פועלת המערכת, הדבר יהיה כרוך בהשקעה של משאבים רבים. על כן, ובפרט כשמדובר בדרישה להסברות ברמה גבוהה, יתכן שמיזמים רבים ייפגעו בכדאיותם הכלכלית או שכדי להוציא אותם לפועל יהיה צורך להסיט משאבים מפיתוח מערכות מתקדמות יותר (בהקשרי יעילות ודיוק למשל), לצורך פיתוח האפשרות להסביר כיצד התקבלה ההחלטה. הדבר עלול בסופו של דבר להיות לא יעיל מבחינה חברתית. כך גם, דרישת ההסברות עלולה להוות חסם משמעותי עבור חברות קטנות המבקשות להיכנס לתחום הבינה המלאכותית, ואשר כתוצאה מהדרישה ייאלצו לפתח לצד המערכת גם את היכולת להסביר את האופן בו היא פועלת.¹⁵⁵

¹⁵¹ Marina Chugunova & Daniela Sele, *We and It: An Interdisciplinary Review of the Experimental Evidence on How Humans Interact with Machines*, 99 J. BEHAV. EXPERIMENTAL ECON. 1, 32-33 (2022).

¹⁵² Donghee Shin, *The Effects of Explainability and Causability on Perception, Trust, and Acceptance: Implications for explainable AI*, 146 INT'L. J. HUM-COMPUT. STUD. 102551 (2020); Andrea Ferrario, Michele Loi, *How Explainability Contributes to Trust in AI*, ACM CONFERENCE ON FAIRNESS, ACCOUNTABILITY, AND TRANSPARENCY (Jan. 28, 2022), https://papers.ssrn.com/sol3/papers.cfm?abstract_id=4011111; יש לציין, כי ביחס לשימוש בבינה מלאכותית על ידי גופים ציבוריים עולה לעיתים נימוק נוסף דאונטולוגי שמתייחס לכך שלתת הסבר לאדם אודות השימוש בכוח שמופעל עליו זה לכבד אותו בתור סובייקט, אולם נימוק זה אינו מובן מאליו כשמדובר בהסדרת פעילות פרטית. להרחבה ראו Jerry Mashaw, *Reasoned Administration: The European Union, the United States, and the Project of Democratic Governance*, 76 GEORGE WASH. L. REV. 99 (2007).

¹⁵³ כך למשל, חלק מהמערכות הללו מסוגלות להצביע על החלקים הספציפיים בתמונה ששימוש את המערכת לניתוח תמונה לשם הגעה לזיהוי כלשהו, או יכולות לייצר "מפת חום" של השיקולים המאפשרת להבין את המשקל היחסי שניתן לשיקולים שונים בתהליך קבלת ההחלטה.

¹⁵⁴ Yavar Bathaee, *The Artificial Intelligence Black Box and the Failure of Intent and Causation*, 31 HARV. J. L. & TECH. 889, 929-930 (2018).

¹⁵⁵ שם, בעמ' 893-894 ו-930; Deeks, לעיל ה"ש 149 בעמ' 1834.

לבסוף, מאחר שהסברתיות מצריכה שקיפות ומסירת מידע אודות הנתונים והמודל, עלולות להתעורר שאלות ביחס לזכויות מפתחי המערכת בתחום הקניין הרוחני, ובפרט בתחום הסודות המסחריים וזכויות היוצרים.¹⁵⁶ עם זאת, הסברתיות המבוססת על הצגת שיקולים מרכזיים, או הסברת החלטת ביחס לאדם מסוים, ללא חשיפת קוד המקור, אינה בהכרח מחייבת חשיפת סודות מסחריים.¹⁵⁷ במקרים אחרים, עשויות לעלות טענות ביחס לחשיפת סודות מסחריים של החברות המפתחות הקשורים לאופן שבו המערכת פועלת, שיצריכו התמודדות ואיזון בין האינטרסים.¹⁵⁸

4.4. גילוי

בניסוי שנערך בשנת 2014 התכנס חבר שופטים במטרה לבחון האם "יוג'ין גוסטמן" (Eugene Goostman) הוא ילד אוקראיני בן 13 (שאינו דובר אנגלית כשפת אם), או שמא מדובר במערכת מבוססת בינה מלאכותית. לאחר כ-300 שיחות עם יוג'ין שבוצעו על ידי 30 שופטים, כל אחת במשך 5 דקות, לראשונה בהיסטוריה נטען כי מערכת מבוססת בינה מלאכותית הצליחה לעבור את מבחן טיורינג, לאחר שהצליחה להטעות 30% מהשופטים שסברו כי הם מדברים עם אדם.¹⁵⁹ ההתקדמות הטכנולוגית בתחום שנים האחרונות, מעוררת חשש כי היכולת של משתמש הקצה לדעת כי הוא בא באינטראקציה עם מערכת מבוססת בינה מלאכותית תהיה תלויה ברצון המפעיל לחשוף זאת.

עם התקדמות הטכנולוגיה והשימוש ההולך והגובר במערכות מבוססות בינה מלאכותית לשם קבלת החלטות, כמו גם השימוש במערכות צ'אט-בוט (chat-bot) המבוססות על בינה מלאכותית,¹⁶⁰ או במערכות יוצרות ליצירת תכנים ופרסומים, גוברת האפשרות שאנשים לא יהיו מודעים לכך שלמערכות אלה תפקיד משמעותי "מאחורי הקלעים" בקבלת ההחלטות בעניינים, שהם משוחחים עם מערכות כאלה (ולא עם אנשים),¹⁶¹ או שהם נחשפים לתכנים ופרסומים שנוצרו על ידן. כפי

¹⁵⁶ Rita Matulionyte, *Reconciling Trade Secrets and Explainable AI: Face Recognition Technology as a Case Study*, 44 EUR. INTELL. PROP. REV. 1, 14 (2022) (להלן: **Matulionyte**).

¹⁵⁷ שם, בעמ' 10: "Overall, while the implementation of AI explainability principle will be compatible with trade secret protection in some cases, in other cases the application of more intrusive explainability techniques might interfere with trade secret protection".

¹⁵⁸ Matulionyte, לעיל ה"ש 156.
¹⁵⁹ *Turing Test Success Marks Milestone in Computing History*, UNIV. READING (June. 08, 2014), <https://archive.reading.ac.uk/news-events/2014/June/pr583836.html>.

¹⁶⁰ תוכנות המתמחות בניהול שיחות טקסטואליות עם בני אדם, שנועדו להיחנות כטבעיות ככל האפשר.

¹⁶¹ Nika Mozafari, Maik Hammerschmidt & Welf H. Weiger, *The Chatbot Disclosure Dilemma: Desirable and Undesirable Effects of Disclosing the Non-Human*, ICIS 2020 PROCEEDINGS. (Nov. 04, 2020), https://aisel.aisnet.org/icis2020/hci_artintel/hci_artintel/6/.

שיפורט להלן, השימוש במערכות מבוססות בינה מלאכותית נעשה על ידי מגוון רחב של גופים, ובין היתר על ידי חברות מסחריות וחברות טכנולוגיה גדולות, חברות הזנק (Start-up) וגופי ממשל (אם כי מסמך זה אינו נוגע להיבטים המיוחדים של השימוש במערכות על ידי גופים ציבוריים). במקרים מסוימים, לגופים אלו יהיה תמריץ לחשוף למשתמש הקצה כי הם עושים שימוש במערכות מבוססות בינה מלאכותית; ואולם, במקרים אחרים, ואלו המקרים בהם יתמקד חלק זה, יהיה לאותם גופים דווקא תמריץ להסתיר אותו, לעיתים תוך ניסיון להטעות את משתמש הקצה.

על רקע זה, מתעוררת שאלה, מה מידת הגילוי לגבי מעורבות מערכות מבוססות בינה מלאכותית (AI related transparency), הנדרשת מצד גורמים המשתמשים במערכות אלה.

יצוין כי יש העוסקים בהיבט זה כחלק משקיפות (Transparency) במשותף עם הסבריות (Explainability). רב המשותף בין דרישות אלה, שכן גילוי בעניין השימוש בבינה מלאכותית ושקיפות לגבי אותו שימוש יכולים להיכלל כחלק מההסבר על תהליך קבלת ההחלטה, וגם הסבריות (בפרט כשמדובר בחשיפת המודל והנתונים ששימוש לאימון המודל) במהותה היא סוג של שקיפות מצד מפתחי ומפעילי המערכת. אולם במסמך זה נערכה הבחנה בין השניים, בין היתר בשים לב לכך שמסירת מידע מתמקדת בגילוי על האינטראקציה עם המערכת, בעוד שההסבריות עוסקת בדרישה להסביר את דרך פעולתה של המערכת ושל החלטות ותחזיות שהתקבלו על ידה.

4.4.1. מה טיבה של דרישת הגילוי?

ניתן לחלק את האינטראקציה האנושית עם מערכות מבוססות בינה מלאכותית לישירה ועקיפה. אינטראקציה ישירה נוצרת כשהגורם האנושי בא בקשר בלתי אמצעי עם המערכת, כגון שהוא משוחח עם צ'אט-בוט או כשהוא נחשף לתוכן שנוצר על ידי מחשב באמצעות כלי בינה מלאכותית יוצרת. לעומת זאת, אינטראקציה עקיפה מתקיימת במצבים בהם במסגרת תהליך קבלת ההחלטה נעשה שימוש במערכת, אולם הקשר הישיר של מושא ההחלטה הוא עם גורם אנושי, שתפקידו לתווך את ההחלטה של המערכת, או להסתייע בה על מנת לקבל את ההחלטה בעצמו. כפי שיפורט להלן, הנימוקים להצבת דרישה לגילוי ביחס לאינטראקציה ישירה או עקיפה עשויים להיות שונים.

להמחשה, נתמקד באדם המעוניין בהלוואה מהבנק. כידוע, כיום גופים רבים עושים שימוש בצ'אט-בוטים על מנת לנתב פניות ולספק מענה ראשוני של שירות הלקוחות, כך שלא מן הנמנע שבפתח תהליך בקשת ההלוואה יידרש המבקש לנהל אינטראקציה ישירה עם מערכת צ'אט-בוט שתהיה מבוססת בינה מלאכותית. בשלב הבא, ולאחר שהגיש את בקשת ההלוואה, ייתכן שהבנק יעשה שימוש במערכת מבוססת בינה מלאכותית על מנת לבחון את הבקשה ולהחליט אם להיעתר לה ובאלו תנאים. השימוש במערכת יכול להיעשות תוך החלפת שיקול הדעת של הפקיד, כשהמערכת קובעת באופן אוטונומי אם ובאלו תנאים לאשר את הבקשה, או שהוא יכול להיעשות על מנת לסייע לפקיד להחליט בעצמו אם ובאלו תנאים לאשר אותה. בעוד שבשלב הראשון האינטראקציה של המבקש עם המערכת היא ישירה, בשלב השני, האינטראקציה עם המערכת היא עקיפה.

בכל אינטראקציה עם מערכת מבוססת בינה מלאכותית, קיימת האפשרות כי אדם לא יהיה מודע לכך שהוא בא באינטראקציה עם מערכת כזו, ויטעה לחשוב שהוא מתקשר עם אדם או שאדם (בלבד) מחליט בעניינו. אדרבה, לא אחת חברות מסחריות משקיעות משאבים רבים כדי לשוות לבוט מאפיינים אנושיים. כך לדוגמה, מערכות עזרה וירטואלית מבוססת קול, עלולות לגרום לטעות ולמחשבה שמדובר בגורם אנושי. זאת, באמצעות ניסיון לחקות את האופן בו אנשים

מדברים (בין היתר, קיימות מערכות המחקר היסוס ועושות שימוש בביטויים כמו "אמממ").¹⁶² מטרתה של הדרישה לגילוי, ככל שתקיים, היא למעשה להבטיח כי לכל אדם אשר בא באינטראקציה עם מערכת המבוססת על בינה מלאכותית תהיה היכולת לדעת זאת.

4.4.2. היתרונות והחסרונות להצבת דרישת גילוי

ככלל, גילוי למשתמש הקצה הבא באינטראקציה עם מערכת מבוססת בינה מלאכותית יכול לסייע בהשגה של מספר מטרות, הן בשלב שלפני השימוש במערכת והן לאחר השימוש בה.

ראשית, הידיעה כי הגורם העומד בפניו משתמש במערכת מבוססת בינה מלאכותית, עשויה להוות אינדיקציה חיובית או שלילית עבור הפרט, שיכולה להשפיע על בחירתו כיצד לפעול, ובין היתר אם לרכוש מוצר מסוים או ממי לקבל שירות. בעוד שברור מדוע ידיעה זו עשויה להשפיע על המשתמש כאשר הוא בא באינטראקציה ישירה עם המערכת, למשל מכיוון שהוא מעדיף לשוחח עם גורם אנושי, הידיעה אם נעשה שימוש במערכת עשויה להשפיע גם כאשר המדובר באינטראקציה עקיפה. כפי שתואר גם בחלקים קודמים, השימוש במערכות מבוססות בינה מלאכותית יכול לשפר את דיוק ההחלטה; על כן, ייתכן למשל כי צרכנים יעדיפו לקבל שירות מחברה אשר מייצרת תחזיות עסקיות באמצעות שימוש בטכנולוגיית בינה מלאכותית על פני חברות אחרות. לעומת זאת, יהיו מצבים בהם השימוש במערכת דווקא ירתיע צרכנים מלבחור בקבלת שירות מסוים.

שנית, הידיעה כי מערכת מבוססת בינה מלאכותית הייתה מעורבת בקבלת החלטה בעניינו של הפרט, עשויה לאפשר לו לוודא שהמערכת עומדת בדרישות הרגולטוריות הרלוונטיות תוך שזכויותיו נשמרות. במובן זה, הידיעה על אודות האינטראקציה עם מערכת מבוססת בינה מלאכותית "פותחת את השער" עבור מי שמושפע מפעילותה לוודא את תקינות ההחלטה או הפעולה של המערכת, ואף של המערכת כולה, בשים לב לכל האתגרים הנסקרים במסמך זה ולסביבה הרגולטורית בה פועלת המערכת. ברי כי ככל שיתפתח שיח הזכויות והחובות סביב הפעלתן של המערכות, ובכפוף להסתייגויות שהוצגו לעיל, כך לידיעה כי נעשה שימוש במערכת יהיה תפקיד משמעותי יותר במיצוי הזכויות ועמידה על החובות של הגורמים הרלוונטיים.

שלישית, ניתן לראות את הגילוי כאקט של הגינות כלפי משתמש הקצה, כך שבמקרים מסוימים ייתכן כי ראוי להכיר במעין "זכות" שלו לדעת אם הוא בא באינטראקציה עם מערכת מבוססת בינה מלאכותית מטעמים מוסריים, וזאת אף שכנראה לא ינהג אחרת אם יגלה על המערכת.

נוסף לטעמים האמורים לעיל, שאינם נוגעים במישרין ליתרונות אפשריים מבחינת מפעיל המערכת, חשוב לציין כי לעיתים גם לו יהיה אינטרס לחשוף את השימוש שהוא עושה במערכת ולגלות על האינטראקציה. זאת משום שהחשיפה עשויה לשפר את התדמית של חברה בתחום הטכנולוגי, להצביע על יכולות מתקדמות יותר ביחס למתחרים, ולאפשר התנהלות בצורה שמותאמת לכך שנעשה שימוש במערכות מבוססות בינה מלאכותית (למשל, בתקשורת עם צ'אט-בוט ניתן לכתוב בשפה ישירה ופשוטה יותר). בפרט, בכל הנוגע לגילוי על אינטראקציה ישירה עם המערכת, לחברה עשוי להיות אינטרס לחשוף את השימוש בצ'אט-בוט במטרה לתת מענה ללקוחות המעדיפים להיות

¹⁶² Fatimah Ishowo-Oloko, Jean-François Bonnefon, Zakariyah Soroye, Jacob Crandall, Iyad Rahwan & Talal Rahwan, *Behavioural Evidence for a Transparency-Efficiency Tradeoff in Human-Machine Cooperation*, 1 NATURE MACH. INTELLIGENCE 517, 517-520 (2019).

באינטראקציה עם המערכת ולא עם גורם אנושי, למנוע תסכול שיתעורר עקב ציפייה לא ריאלית ביחס ליכולות המערכת ומתוך תפיסה עקרונית כי עדיף להתנהל בשקיפות למול לקוחות.¹⁶³

מנגד, לצד יתרונותיה, דרישה לגילוי עלולה לפגוע ביעילות השימוש במערכת מבוססת בינה מלאכותית ולעיתים לפגוע בנכונות לעשות שימוש במערכות מבוססות בינה מלאכותית. במיוחד בנוגע לאינטראקציה ישירה, חשיפת השימוש במערכת עלולה לפגוע ביעילות. מחקרים מצאו כי לכל הפחות בעת הזו משתמשים נוטים לסמוך פחות על מערכות מבוססות בינה מלאכותית ביחס לגורמים אנושיים, אפילו אם רמת השירות זהה,¹⁶⁴ ובכלל זה לנהל אינטראקציות קצרות יותר עם צ'אט-בוטים; לרכוש פחות מוצרים; ובאופן כללי לא לשתף פעולה עם המערכת.¹⁶⁵ ייתכן עוד שחשיפה הנובעת מחובת הגילוי תפגע ביעילות השימוש במערכת גם במקרים של אינטראקציה עקיפה, למשל בכך שתקל על ניסיונות זדוניים להשפיע על התוצאה שתתקבל על ידי המערכת, בין אם על ידי משתמש הקצה ובין אם על ידי מתחרים (להרחבה ראו פרק "אמינות, עמידות, אבטחה ובטיחות" להלן).

4.5. אמינות, עמידות, אבטחה ובטיחות

במרץ 2016 הוצג לעולם "טאי" (Tay), בוט המתבסס על בינה מלאכותית, שנועד לחקות נערה אמריקאית בת 19 ולנהל שיחות אקראיות עם משתמשים בטוויטר. החברה המובילה שפיתחה אותו קיוותה שבאמצעות האינטראקציה עם משתמשים אחרים טאי יתפתח ויוכל לנהל שיחות מתוחכמות יותר ויותר, ואגב כך יסייע לה בפיתוח מערכות מתקדמות בעתיד. אלא שטאי שרד "באוויר" רק 16 שעות לפני שהחברה מיהרה להשבית אותו. התברר שבפרק הזמן הקצר שבו פעל טאי, הוא "למד" מהגולשים האחרים, ובהתאם התחיל להגיב בצורה מגונה, אנטישמית וגזענית. לטענת החברה המפתחת, מי שהיו אחראיים להידרדרות טאי הם הגולשים אשר לקחו חלק ב"מאמץ מתואם לנצל לרעה את כישורי הדיבור שלו כדי לגרום לו להגיב באופן לא הולם".¹⁶⁶ מקרה זה ממחיש את החשש שחלק מן המערכות המבוססות על בינה מלאכותית אינן מספיק אמינות, עמידות או בטוחות על מנת שניתן יהיה להפעילן בשגרה בלי לבצע התאמות מיוחדות.

מערכות מבוססות בינה מלאכותית נשענות על טכנולוגיות מידע ותקשורת, החשופות לטעויות ותקלות טכניות, או למניפולציות מכוונות. עקב כך, על מנת ליהנות מהתועלות הקשורות בבינה מלאכותית, יש חשיבות רבה לוודא את תפקודה התקין, הבטוח והמהימן של טכנולוגיה זו.

היקף השימוש הצפוי בבינה מלאכותית, במוצרים שכיום אין להם יכולת חישובית (כגון מכוניות, ומוצרים מסוגים שונים), מרחיב את היקף הרלוונטיות של סוגיות בטיחות הנובעות מסיכונים ממוחשבים. כך, כל עוד מדובר במוצרים שיש בהם רכיבים פיזיים בלבד, הסיכון הבטיחותי מהם

Roberta De Cicco, Susana Cristina Lima da Costa e Silva & Riccardo Palumbo, *Should a Chatbot Disclose Itself? Implications for an Online Conversational Retailer*, in CHATBOTS RESEARCH AND DESIGN 3, 5 (Asbjørn Følstad, Theo Araujo, Symeon Papadopoulos, Effie L.-C. Law, Ewa Luger, Morten Goodwin & Petter Bae Brandzaeg eds., 2020).

Nika Mozafari, Welf H. Weiger & Maik Hammerschmidt, *Resolving the Chatbot Disclosure Dilemma: Leveraging Selective Self-Presentation to Mitigate the Negative Effect of Chatbot Disclosure*, HAW. INT'L CONF. ON SYS. SCI., 2916 (2021) <https://scholarspace.manoa.hawaii.edu/>

Xueming Luo, Siliang Tong, Zheng Fang & Zhe Qu, *Frontiers: Machines vs. Humans*; 2917 בעמ' 2917; *The Impact of Artificial Intelligence Chatbot Disclosure on Customer Purchases*, 38 MARKETING SCI. 913 (2019).

Davey Alba, *It's Your Fault Microsoft's Teen AI Turned into Such a Jerk*, WIRED (Mar. 25, 2016), <https://www.wired.com/2016/03/fault-microsofts-teen-ai-turned-jerk/>.

נובע מאופן פעולתם הפיזי. שילוב רכיבים ממוחשבים, במיוחד במערכות כגון כלי רכב, קווי ייצור ומתקני תשתית, מייצר סיכון מסוג חדש.

השימוש במערכות מבוססות בינה מלאכותית עשוי לייצר חשיפה לתוצאות שליליות, שעלולות להיגרם עקב כשלים בפעילות מערכות אלה. כפי שיפורט להלן, תוצאות אלה עלולות להיגרם הן בעקבות התממשות כשלים פנימיים שמשפיעים על אמינות המערכת (Reliability) – מצבים שבהם המערכת סובלת מביצועים ירודים (poor performance) משום שהיא אינה מסוגלת לבצע כהלכה את המשימה שיועדה לה; הן בעקבות התממשותם של איומים חיצוניים שמשפיעים על עמידות המערכת (robustness) – מצבים שבהם גורם חיצוני מצליח לשבש את פעילות המערכת על ידי ניצול של נקודת תורפה הטבועה בה (vulnerability).¹⁶⁷ לעיתים התממשות הכשלים הפנימיים או האיומים החיצוניים אף עלולה לגרום ליצירת סיכונים בטיחותיים וכלכליים משמעותיים. על רקע זה, יש חשיבות לאמינות ודיוק,¹⁶⁸ עמידות¹⁶⁹ ובטיחות¹⁷⁰ מערכות מבוססות בינה מלאכותית.

הכשלים האמורים הם ככלל אתגרים טכנולוגיים, נהלים או תהליכיים, ולרוב מפתחי המערכות והמשתמשים בהן יהיו בעלי האינטרס המרכזי להתמודד עמם, שכן הם אמורים להיות מעוניינים ליצור או להשתמש במערכות איכותיות ומוגנות. עם זאת, במקרים מסוימים, עשויה להיות הצדקה לשקול התערבות רגולטורית, על מנת לוודא כי מערכות מבוססות בינה מלאכותית הן אמינות, עמידות, מאובטחות ובטוחות במידה מספקת. זאת, בפרט בנסיבות בהן מתקיימות הצדקות מקובלות להתערבות רגולטורית, כדוגמת החצנות שליליות הכרוכות בפעילות מערכות אלה.

בהקשר זה, יתכן שיש פגיעות מיוחדת למערכות אשר משלבות פרקים של אימון בתוך מחזור השימוש הרגיל שלהן (באופן אשר מכונה לעיתים "Active learning" או "Continuous Learning") – מערכות אשר משלבות קלט חדש ועדכני אל תוך מאגר הנתונים ומאמנות את המודל בצורה מחזורית. במערכות אלו במיוחד נדרשת בקרת איכות לכל אורך חיי המוצר, ולא ניתן להסתפק בבדיקה אחת "בשערי המפעל".

4.5.1. אמינות המערכת

לא אחת, למרות שמערכת המבוססת על בינה מלאכותית פועלת בתנאים שגרתיים, ואפילו אופטימאליים, התוצאות שאליהן היא תגיע לא יהיו מספיק טובות, עקב כשלים פנימיים. כלומר, מכיוון שאופן פיתוח המערכת או השימוש בה גורם לכך שלא תוכל לבצע כיאות את הפונקציה שיועדה לה. כשלים אלה, המשליכים על אמינות המערכת, עלולים להתרחש בכל שלב ב"מחזור

¹⁶⁷ Ronan Hamon, Henrik Junklewitz & Ignacio Sanchez, *Robustness and Explainability of Artificial Intelligence – From Technical to Policy Solutions*, JRC TECHNICAL REPORT, 14 (2020) (להלן: **Hamon, Junklewitz & Sanchez**).

¹⁶⁸ בטיחות מסגרת התקינה של NIST (מכון התקנים הטכנולוגי האמריקאי) לבינה מלאכותית מוצע להגדיר דיוק, בהתבסס על תקן ISO/IEC TS 5723: 2022 כ- "הקרבה של תוצאות, תצפיות, חישובים או הערכות לערכים אמיתיים או לערכים הנתפסים כאמיתיים". מדידת הדיוק צריכה להיות קשורה ליעדי בדיקה ברורים ופרטים לגבי שיטת הבדיקה, וכן מתועדים. במסגרת המוצעת של NIST מוצע להגדיר אמינות בהתאם לתקן ISO/IEC TS 5723: 2022 כ- "יכולת של חפת לתפקד כמתוכנן, ללא כשל, לפרק זמן נתון בתנאים נתונים". <https://www.nist.gov/itl/ai-risk-management-framework>. Comment.

¹⁶⁹ במסגרת המוצעת של NIST מוצע להגדיר עמידות בהתאם לתקן ISO/IEC TS 5723: 2022 כ- "יכולת של מערכת בינה מלאכותית לשמור על רמה של ביצוע בתנאים שונים". המשמעות היא לא רק ציפיה שמערכת הבינה המלאכותית תפעל כפי שתוכננה, אלא שגם תפעל כפי שתוכננה באופן שמצמצם סיכונים לאנשים גם בתנאים לא צפויים.

¹⁷⁰ בהתאם למסגרת המוצעת של NIST מוצע להגדיר בטיחות על פי תקן ISO/IEC TS 5723: 2022, באופן ש- "מערכות בינה מלאכותית, לא יגרמו, בתנאים מוגדרים, לנזק פיזי או פסיכולוגי או יובילו למבצע שבו יש סיכון לחיי אדם, בריאות, רכוש או הסביבה". הפעלה בטוחה של בינה מלאכותית מצריכה תהליכי תכנון ופיתוח אחראים, מידע למיישמי הבינה המלאכותית לגבי אופן השימוש במערכת, וקבלת החלטות אחראית בידי המיישמים ומשתמשי הקצה.

החיים" שלה – בעיצוב המערכת; באיסוף וניתוח המידע והנתונים שעליהם מתבססת המערכת; ביצירת המודל שלפיו המערכת מקבלת החלטות; ובממשק הסופי בין המערכת לבין המשתמשים.

כך לדוגמה, בנוגע לשלב אימון המערכת (AI training), החשש המרכזי הוא שהאימון לא יעשה בצורה מספיק טובה, מה שיוביל לביצועים ירודים של המערכת ויגרע מאמינותה. אחד המקרים המרכזיים שבהם המערכת עלולה לסבול מביצועים ירודים הוא כאשר הנתונים ששימשו לאימון המערכת אינם מהימנים או מקיפים (תופעה המכונה הטיית ספקטרום (Spectrum bias)). בהקשר זה, החשש הוא שכאשר המערכת תידרש להתמודד עם מצבים או קבוצות אוכלוסייה שמאפייניהם חורגים מהמנעד שעליו היא התאמנה התוצאות שלה יהיו פחות מדויקות, וייתכן שאף שגויות לחלוטין. כך למשל, אם האימון של מכונית עצמאית (אוטונומית) נעשה בישראל, ייתכן שבעת שימוש בה במדינה אחרת, היא לא תהיה ערוכה להתמודד עם נסיעות בתנאי שטח שאינם שכיחים בישראל כמו שלג או סופות חול, או שלא תדע כיצד להתמודד עם חיות על הכביש שאינן נפוצות בישראל כמו הקנגורו השכיח באוסטרליה. בהקשר זה, וכפי שתואר לעיל בפרק "אפליה", קיים גם חשש כי כאשר האימון של המערכת יתמקד בקבוצות מסוימות באוכלוסייה, היא תפעל באופן שיטתי בצורה פחות מדויקת ביחס לקבוצות אחרות באוכלוסייה. לדוגמה, במחקר שבחן מערכות לזיהוי פנים נמצא כי מספר מערכות, שפותחו על ידי חברות מובילות שונות, נטו לזהות ביתר קלות פנים של גברים לבנים לעומת פנים של נשים שחורות, ככל הנראה בגלל ייצוג נמוך יותר לנשים שחורות במאגרי המידע והנתונים שעל בסיסן התאמנו המערכות.¹⁷¹

מאפיין נוסף של מערכות מבוססות בינה מלאכותית שעלול לגרום לתוצאה דומה לזו של הטיית הספקטרום, הוא היותן של חלק מן המערכות שבריריות (Brittle). כלומר, מערכות שעוצבו באופן כזה שאינן מסוגלות לבצע הכללות או להתאים עצמן לנסיבות משתנות. תכונות אלה הכרחיות למערכת מבוססת בינה מלאכותית, כיוון שלא תמיד יש יכולת ממשית לחשוף מערכת לנתונים המשקפים כל סיטואציה אפשרית או לחזות את כלל המצבים שעמם היא תידרש להתמודד בעתיד. כשמערכות שבריריות יידרשו להתמודד עם מצבים לא מוכרים, הדיוק שלהן ייפגע והסיכוי שיגיעו לתוצאה הנכונה יפחת באופן משמעותי.¹⁷² בעיה זו חמורה במיוחד, כאשר יש כוונה לעשות שימוש במערכות מבוססות בינה מלאכותית בתחומים שבהם זמינות המידע והנתונים לוקה בחסר, וניתן לחשוף את המערכת מראש רק לחלק קטן מסוגי המקרים שעמם היא תידרש להתמודד.

בשלב הבא, לאחר כניסת המערכת לשימוש (AI uses), אחד החששות המרכזיים הוא מהשפעות במאפייני משתנה היעד של המערכת (כלומר, שינוי במאפיינים של מה שהמערכת נועדה לחזות). תופעה זו, המכונה "Concept Drift", מתייחסת לחשש כי ככל שעובר זמן רב יותר משלב האימון של המערכת, כך הנתונים ההיסטוריים שעליהם היא התאמנה ישקפו בצורה פחות טובה את המציאות שבה היא פועלת, ולכן ישנו חשש כי המערכת תהיה פחות מדויקת וחשופה לשגיאות בלתי צפויות. לשם ההמחשה, ביחס למערכת שמטרתה לבצע מסחר בניירות ערך (algo-trading), אחד

Joy Adowaa, *Gender Shades: Intersectional Phenotypic and Demographic Evaluation of Face Datasets*¹⁷¹ and *Gender Classifier* (Unpublished B.A. Thesis, MASSACHUSETTS INSTITUTE OF TECHNOLOGY), (2017) <https://dspace.mit.edu/handle/1721.1/114068>

¹⁷² דוגמה מפורסמת לקושי של מערכת לבצע הכללות באה לידי ביטוי בתחום הנהיגה האוטונומית; התברר שמערכות מתקדמות שהצליחו לזהות בצורה מדויקת אוטובוס הסעות, טעו בסיווג ב-97% מהמקרים כאשר בתמונה האוטובוס נטה על צדו. ראו Michel A. Alcorn, Qi Li, Zhitao Gong, Chengfei Wang, Long Mai, Wei-Shinn Ku & Anh Nguyen, *Strike (with) a Pose: Neural Networks Are Easily Fooled by Strange Poses of Familiar Objects*, CORNELL UNIVERSITY (Apr. 18, 2019), <https://arxiv.org/abs/1811.11553>; Mary L. Cummings, *Rethinking the Maturity of Artificial Intelligence in Safety-Critical Settings*, 42 AI MAG. 6, 7-8 (2021).

החששות הוא ששינויים בשוק ניירות הערך עלולים לפגוע במידת הדיוק של המערכת; כך למשל, שהשימוש ההולך וגובר בשנים האחרונות בחברות רכישה למטרות מיוחדות (חברות SPAC),¹⁷³ עלול להוביל לעיוותים ולתוצאות שגויות של המערכת, מכיוון שייתכן שהמערכת לא נחשפה בתהליך האימון לחברות מסוג זה, שהשימוש בהן לא היה שכיח בעבר.¹⁷⁴

חשש נוסף שקיים בשלב השימוש במערכת הוא מ"שכחה קטסטרופלית" (Catastrophic Forgetting). חשש זה נוגע לכך שכאשר מנסים ללמד מערכת לבצע מספר משימות, לעיתים הלימוד של משימות נוספות גורם לה "לשכוח" איך לבצע את המשימה הראשית.¹⁷⁵ כך לדוגמה, חוקרים מדרום קוריאא הקימו אתר אינטרנט שנועד לסייע, באמצעות מערכת מבוססת בינה מלאכותית, להבחין בין תמונות אמיתיות לבין זיופים שנעשו באמצעות טכנולוגיית זיוף עמוק (deep-fake). בהתחלה האתר עבד בצורה טובה, אולם לאחר מספר חודשים נעשה שימוש בטכנולוגיות חדשות ליצירה של זיופים עמוקים. כאשר החוקרים אימנו את המערכת להתמודד עם הסוגים החדשים של הזיוף, התברר להם כי המערכת "שכחה" איך לזהות את הסוג המקורי של הזיוף העמוק.¹⁷⁶

חשוב לציין כי אין מדובר ברשימה ממצה של כשלים פנימיים ושל סוגי המקרים שמעוררים חשש לפגיעה במידת האמינות של מערכות מבוססות בינה מלאכותית. הסקירה לעיל אך נועדה להציג כמה כשלים מרכזיים, ולהמחיש כיצד עלולה להיגרם פגיעה באמינות מערכות אלה.

4.5.2. עמידות המערכת ואבטחתה

לצד החשש מפני כשלים פנימיים, קיים חשש כי מערכות מבוססות בינה מלאכותית יגיעו לתוצאות שגויות או שפעולתן תגרום לנזק עקב התערבות גורם חיצוני (adversary), לרוב בעל כוונות זדוניות, שינצל חולשה של המערכת על מנת לשבש את פעילותה או להשפיע על אופן פעילותה. חשש זה נוגע בעיקר לתחום המחקר של למידת מכונה אדברסרית (Adversarial Machine Learning) העוסק בטכניקות המאפשרות לנצל, לשבש או לשטות במודלים חכמים, עם או בלי גישה למודל המערכת עצמו ולמידע שעליו המודל מתבסס, ובדרכים שבאמצעותן ניתן להתגונן מפני מתקפות כאלה.¹⁷⁷

הגם שהחשש מתקיפה ושיבוש קיים כמעט ביחס לכל סוג של תוכנה, ובהתאם התפתחו ענפים באקדמיה, בתעשייה וברגולציה העוסקים בהגנה על מערכות תוכנה, בכל הנוגע למערכות מבוססות בינה מלאכותית חשש זה מחריף ואף עלול לבוא לידי ביטוי בסוגי תקיפה ושיבוש ייחודיים. זאת, משום שמגוון הדרכים שבהן ניתן לפגוע בפעילות של מערכות אלה רחב יותר.¹⁷⁸

¹⁷³ תאגיד אשר נרשם לבורסה כשלד בורסאי ללא פעילות במטרה לרכוש חברה פרטית ולהפוך אותה לציבורית בלי שהחברה הנרכשת תידרש לנהל הנפקה ראשונה לציבור. Johannes Kolb & Tereza Tykvova, *Going Public Via Special Purpose Acquisition Companies: Frogs Do Not Turn Into Princes*, 40 J. CORPORATE FIN. 80, 83-84 (2016).

¹⁷⁴ Leslie, *Understanding artificial intelligence*, לעיל ה"ש 111, בעמ' 37.

¹⁷⁵ Ian J. Goodfellow, Mehdi Mirza, Da Xiao, Aaron Courville & Yoshua Bengio, *An Empirical Investigation of Catastrophic Forgetting in Gradient-Based Neural Networks*, 1 CORNELL UNIVERSITY (Mar. 4, 2015), <https://arxiv.org/abs/1312.6211>; James Kirkpatrick, Razvan Pascanu, Neil Rabinowitz, Joel Veness, Guillaume Desjardins, Andrei A. Rusu, Kieran Milan, John Quan, Tiago Ramalho, Agnieszka Grabska-Barwinska, Demis Hassabis, Claudia Clopath, Dharshan Kumaran & Raia Hadsell, *Overcoming Catastrophic Forgetting Inneural Networks*, CORNELL UNIVERSITY (Jan. 25, 2017), <https://arxiv.org/abs/1612.00796>.

¹⁷⁶ Shahroz Tariq, Sangyup Lee & Simon S. Woo, *One Detector to Rule Them All: Towards a General Deepfake Attack Detection Framework*, CORNELL UNIVERSITY (May. 1, 2021), <https://arxiv.org/abs/2105.00187>; Charles Q. Choi, *7 Revealing Ways AIs Fail*, IEEE SPECTRUM (Sept. 21, 2021), <https://spectrum.ieee.org/ai-failures>.

¹⁷⁷ המכון הישראלי למדיניות טכנולוגיה *למידת מכונה אדברסרית: התפתחויות במחקר, סכנות והשלכות* 2 (2021) (להלן: "למידת מכונה אדברסרית").

¹⁷⁸ שם.

ניתן לחלק את הדרכים השונות לתקוף מערכת המבוססת על בינה מלאכותית לשלוש קטגוריות שונות: מתקפות שמטרתן לגרום למערכת **ללמוד** את הדבר הלא-נכון (Learn the wrong thing); מתקפות שמטרתן לגרום למערכת **לעשות** את הדבר הלא-נכון (Do the wrong thing); ומתקפות שמטרתן לגרום למערכת **לגלות** לתוקף את הדבר הלא-נכון (Reveal the wrong thing).

בנוגע למתקפות שמטרתן לגרום למערכת **ללמוד** את הדבר הלא-נכון, שיטה נפוצה לדוגמה, מכונה "תקיפת הרעלת מידע" (Data Poisoning Attacks). תקיפה זו מתייחסת למצבים שבהם התוקף משפיע על התנהגות המערכת על ידי מניפולציה של הנתונים המשמשים לאימון המערכת. ככל שהמערכת משתמשת במידע רב יותר לאימון, כפי שלמשל נדרש במערכת המבוססת על לימוד עמוק (deep learning), וככל שהמידע שבו משתמשת המערכת מגיע ממקורות לא מאובטחים או לא ידועים, כך גדלה החשיפה למתקפות מסוג הרעלת מידע.¹⁷⁹ מתקפות הרעלת מידע יכולות לשמש תוקף פוטנציאלי על מנת לפגוע בביצועים של המערכת (להפחית את הדיוק), לנטרל את המערכת, או ליצור דלת אחורית (back-door) שתאפשר לנצל את הפרצה בעתיד.¹⁸⁰ חלק משמעותי מהאתגר בהתמודדות עם תקיפות הרעלת מידע הוא שבעוד שדרך ההגנה המסורתית מפני מתקפות היא באמצעות יצירת חץ בין פעילות המערכת לבין מידע חיצוני, עבור מערכות מבוססות בינה מלאכותית השימוש במידע חיצוני הוא הכרחי על מנת לאמן את המערכת לפעול כהלכה.¹⁸¹ מטבע הדברים, הסיכון האמור נכון במיוחד במערכות אשר הליך האימון בהן נמשך באופן רציף, כפי שתואר לעיל.

באשר למתקפות שמטרתן לגרום למערכת **לעשות** את הדבר הלא-נכון, שיטה נפוצה למשל, מכונה "תקיפת התחמקות" (Evasion Attacks). תקיפה זו מתייחסת למצבים שבהם התוקף מבצע מניפולציה על הקלט המוזן למערכת במטרה לגרום לה לסווג את הקלט באופן שגוי. במקרים רבים המטרה של התקיפה היא לגרום לכך שהמערכת לא תצליח לסווג קלט מסוים (למשל, למנוע מתיבת הדואר האלקטרוני לסווג הודעה מסוימת כ"ספאם").¹⁸² אחת השיטות לבצע תקיפת התחמקות היא באמצעות דוגמה אדברסרית (Adversarial Example), קלט שעוצב באופן ייעודי כך שמצד אחד הוא יגרום למערכת לזהות אותו באופן שגוי, אבל שמצד שני יהיה בלתי ניתן לגילוי על ידי גורם אנושי שיבצע בקרה על המידע.¹⁸³ כך למשל, חוקרים הצליחו למצוא כי אפילו שינוי של פיקסל אחד בלבד בתמונה עלול לגרום לטעות בסיווג התמונה על ידי מערכות לזיהוי פנים.¹⁸⁴

לבסוף, ובהתייחס למתקפות שמטרתן לגרום למערכת **לגלות** לתוקף את הדבר הלא-נכון, הכוונה היא למקרים שבהם התקיפה נעשית על מנת לחלץ מהמערכת מידע שהיא לא הייתה אמורה לגלות. להמחשה, שיטה מרכזית לגרום למערכת לגלות את הדבר הלא-נכון מכונה "היפוך מודל" (Model Inversion). במסגרתה התוקף, שאינו חשוף לכלל המידע המצוי בפני המערכת, יכול להסיק באמצעות ניתוח הפלטים של המערכת, תכונות הנוגעות למידע שמשמש את המערכת. החשש

Avi Schwarzschild, Micah Goldblum, Arjun Gupta, John P Dickerson & Tom Goldstein, *Just How Toxic is Data Poisoning? A Unified Benchmark for Backdoor and Data Poisoning Attacks*, 1 CORNELL UNIVERSITY (Jun. 17, 2021), <https://arxiv.org/abs/2006.12557>.

¹⁷⁹ Hamon, Junklewitz & Sanchez, לעיל ה"ש 167, בעמ' 17.

¹⁸⁰ Jacob Steinhardt, Pang Wei Koh & Percy Liang, *Certified Defenses for Data Poisoning Attacks*, 1 CORNELL UNIVERSITY (Nov. 24, 2017), <https://arxiv.org/abs/1706.03691>.

¹⁸¹ Zeinab Khorshidpour, Sattar Hashem & Ali Hamze, *Learning a Secure Classifier Against Evasion Attack* IEEE 18th INT'L CONF. ON DATA MINING WORKSHOPS, 295 (2016).

¹⁸² למידת מכונה אדברסרית, לעיל ה"ש 177, בעמ' 3.

¹⁸³ Jiawei Su, Danilo Vasconcellos Vargas & Sakurai Kouichi, *One Pixel Attack for Fooling Deep Neural Networks*, CORNELL UNIVERSITY (Oct. 17, 2019), <https://arxiv.org/abs/1710.08864>.

המרכזי לגבי סוג זה של מתקפות נוגע לפגיעה בפרטיות, שעלולה להיגרם עקב האפשרות שיתגלו פרטי מידע ממאגרי מידע שבמקרים רבים הם חסויים (להרחבה ראו פרק "פרטיות" להלן).¹⁸⁵ חשוב לציין, כי גם כאן אין מדובר ברשימה ממצה של איומים חיצוניים ושל החששות המתעוררים בקשר לעמידות ואבטחת מערכות בינה מלאכותית בפניהם. כמו כן, חשוב להבהיר, שחלק זה אינו עוסק באמצעים הטכנולוגיים הקיימים והמפותחים על מנת להתמודד עם איומים אלה. הסקירה לעיל אך נועדה להמחיש איומים חיצוניים מרכזיים, את החשיבות לעמידות המערכות בפניהם ואת הצורך לבחון התערבות רגולטורית על מנת להתמודד עם איומים אלו.

4.5.3. בטיחות המערכת

השימוש במערכות מבוססות בינה מלאכותית לא בהכרח יניב תמיד את התוצאה הרצויה. מדובר בחשש לא מבוטל, הנובע מהאפשרות שיעשה שימוש במערכת שככלל אינה מדויקת מספיק, וגם מהאפשרות שהמערכת לוקה בכשלים פנימיים או חשופה לאיומים חיצוניים. במקרים מסוימים החשש משגיאות של מערכות אלה לא גדול, אך במקרים אחרים הן עלולות לגרום פגיעה משמעותית בזכויות הפרט או באינטרסים ציבוריים, ובפרט בבטיחות המשתמשים בהן והמושפעים מהן.

ביחס לפגיעה בזכויות הפרט, חשוב להדגיש כי החשש לפגיעה כזו לא מתעורר רק במקרים של שימוש במערכות מבוססות בינה מלאכותית במגזר הציבורי.¹⁸⁶ גם שימוש במערכות אלו במגזר הפרטי, כגון על מנת להחליט אם לקבל עובד מסוים לעבודה או אילו עובדים לפטר, או על מנת לקבוע אם פרט מסוים זכאי לקבל אשראי ובאילו תנאים, עלול לגרום, במקרה של התממשות כשל פנימי או איום חיצוני, לפגיעה משמעותית בזכויות ובאינטרסים של הפרט.

באשר לפגיעה באינטרסים ציבוריים, למשל, קיים חשש כי השימוש במערכות מבוססות בינה מלאכותית בתחום השירותים הפיננסיים עלול לגרום לפגיעה ביציבות השווקים ובאינטרסים כלכליים רחבים. כך לדוגמה, בשנת 2013 הופץ פרסום כוזב כי התרחש פיצוץ בבית הלבן שהביא לצניחה של המדדים המרכזיים בבורסות בארה"ב. אחת הסיבות המרכזיות לצניחה המהירה הייתה כי חברות השקעה רבות עשו שימוש במערכות מבוססות בינה מלאכותית אוטונומיות למסחר בשוק ההון (Algo-trading). המערכות, שנחשפו לידיעה החדשותית הכוזבת, הגיבו באופן אוטומטי ומיידי לידיעה וגרמו לצניחה ניכרת של המדדים. אמנם השוק חזר לעצמו במהרה לאחר שהתברר שמדובר בדיווח כוזב,¹⁸⁷ אך דוגמה זו ממחישה כיצד שגיאות של המערכת, במקרה זה על רקע התממשות איום חיצוני, עלולות להוביל להשלכות כלכליות משמעותיות.¹⁸⁸

לבסוף, התממשות כשלים פנימיים במערכת או איומים חיצוניים עלולה לגרום לסיכונים בטיחות משמעותיים. כך למשל, הגם שבאופן עקרוני שימוש במכוניות עצמאיות צפוי להפחית את מספר התאונות בכבישים, כשלים של המערכת ואיומים עליה עלולים לגרום לתאונות ולפגיעה בבטיחות משתמשי הדרך.¹⁸⁹ לשם ההמחשה, נמצא כי הדבקה של מדבקה קטנה על תמרור "עצור" עלולה

¹⁸⁵ Michael Veale, Reuben Binns & Lilian Edwards, *Algorithms That Remember: Model Inversion Attacks and Data Protection Law*, 376 PHIL. TRANSACTION ROYAL SOC'Y A. 1, 4-6 (2018).

¹⁸⁶ להרחבה על השימוש במערכות המבוססות על בינה מלאכותית בתחום הרווחה ראו גם סיון תמיר **בינה מלאכותית בשירותי ממשל: הטמעת מערכות לקבלת החלטות מבוססות אלגוריתם בשירותי הרווחה** (המכון הישראלי למדיניות טכנולוגית וגיוונוס-אלכא 2020).

¹⁸⁷ Patti Domm, *False Rumor of Explosion at White House Causes Stocks to Briefly Plunge; AP Confirms Its Twitter Feed Was Hacked*, CNBC (Apr. 23, 2013), <https://www.cnbc.com/id/35781111>.

¹⁸⁸ להתייחסות רחבה יותר להשלכות של השימוש במערכות המבוססות על בינה מלאכותית במגזר הפיננסי ראו אחיעז, חמדני, עמירם וקסטיאל, לעיל ה"ש 25.

¹⁸⁹ Zachary Arnold & Helen Toner, *AI Accidents: An Emerging Threat - What Could Happen and What to Do*, CSET POLICY BRIEF, 16 (2021) <https://cset.georgetown.edu> (להלן: **Arnold & Toner**).

לגרום לכך שהמערכת לא תצליח לזהות את התמרור, מה שעלול לשבש את מודל הנהיגה האוטונומי ולגרום למפגע בטיחותי באקראי או בתקיפה מכוונת (ראו דיון בתקיפת התחמקות (Evasion Attack) לעיל).¹⁹⁰ באופן דומה, קיימים סיכונים בשימוש במערכות מבוססות בינה מלאכותית בתחום הבריאות. כבר כיום נעשה בעולם שימוש במערכות אלה על מנת לסייע באבחון של מצבים רפואיים ועל מנת להמליץ לרופאים על אפשרויות טיפול. לנוכח רגישות הטיפול הרפואי, במקרים מסוימים כשלים של המערכת ואיומים עליה עלולים להביא לפגיעה ממשית בבריאות מטופל.

4.6. אחריותות

ב-18 למרץ 2018 התנגשה מכונית באיליין הרצברג, כשחצתה את הכביש שלא במעבר חציה, וכתוצאה מהתאונה היא נפטרה מפציעה. זו הייתה הפעם הראשונה שבה נהרג הולך רגל בתאונה שבה מעורבת מכונית עצמאית (אוטונומית). התאונה התרחשה בארה"ב במסגרת ניסוי, כאשר המכונית פעלה באופן עצמאי, כלומר לא בשליטתו של נהג הבטיחות. בעוד שרשויות האכיפה בארה"ב מיהרו להודיע כי לא ינקטו הליכים פליליים כנגד החברה שהפעילה את הניסוי אף שהמערכת לא הצליחה לסווג ולחזות כנדרש את מסלול ההליכה של הולכת הרגל, כשנתיים לאחר התקרית הוגש כתב אישום כנגד נהגת הבטיחות, הנהגת האנושית שתפקידה להתערב במקרה שהרכב אינו פועל כנדרש. לפי הנטען בכתב האישום, הנהגת גרמה ברשלנות למותה של הולכת הרגל מכיוון שהייתה עסוקה בטלפון הסלולארי שלה בזמן התאונה.¹⁹¹ המשפט הפלילי בעניינה, שעדיין מתנהל נכון לכתובת שורות אלה, מעורר עניין רב ברחבי העולם. שכן באופן עקרוני, וכפי שיפורט להלן, מקרים מסוג זה מעלים שאלות בנוגע לאפשרות להטיל אחריות בכלל, ואחריות משפטית (פלילית או אזרחית) בפרט, במקרה של שימוש במערכות מבוססות בינה מלאכותית.

האחריותיות (Accountability), ובכלל זה האחריות המשפטית (Legal liability), נמנית עם האתגרים המרכזיים הקשורים בשימוש במערכות מבוססות בינה מלאכותית. המאפיין האוטונומי של מערכות אלה, והעובדה כי הן מסוגלות לחקות שלבי "מחשבה" מסוימים המזהים עם פעילות אנושית ולפעול בהתאם, עלולים במקרים מסוימים לערער את המבנה של האחריותיות, המושתת על הימצאות אדם במוקד ההתרחשות שכלל נושא באחריות למעשיו. ככל שהמערכות האנושיות בהחלטה או בפעולה מצטמצמות ורחוקה יותר, אתגר זה מתחדד, ומתעוררות שאלות בנוגע לנשיאה באחריות בגין טעויות שנגרמו אגב השימוש במערכות אלה, ובכלל זה, מי נושא באחריות? מה סוג האחריות שניתן לייחס לו? ומה משטר האחריות המתאים לבחינת שאלת האחריות?

המונח אחריותיות מתייחס ככלל לצורך בכך שיהיו גורמים מתאימים שיוודאו כי המערכת מפותחת ופועלת בהתאם לתפקיד שהוגדר לה ולסביבה הרגולטורית הרלוונטית. ניתן להגביר את מידת האחריותיות באמצעות קידום נורמות ארגוניות המבטאות נטילת אחריות, וזאת לצד הטלת אחריות משפטית, מוסרית או חברתית על הגורמים המתאימים כאשר הם לא עומדים בחובתם.

¹⁹⁰ למידת מכונית אדברסרית, לעיל ה"ש 177, בעמ' 3; Kevin Eykholt, Ivan Evtimov, Earlence Fernandes, Bo Li, Amir Rahmati, Chaowei Xiao, Atul Prakash, Tadayoshi Kohno & Dawn Song, *Robust Physical-World Attacks on Deep Learning Visual Classification* IEEE/CVF CONF. COMPUT. VISION AND PATTERN RECOGNITION (2018), <https://ieeexplore.ieee.org/document/8578273>.

¹⁹¹ Matt McFarland, *Uber Self-Driving Car Operator Charged in Pedestrian Death*, CNN (Sept. 18, 2020), <https://edition.cnn.com/2020/09/18>.

בנוגע לקידום נורמות ארגוניות, המטרה היא להביא לכך שמפתחי המערכות והמשתמשים בהן יבינו כיצד הן משליכות על זכויות יסוד ואינטרסים ציבוריים, ויכולו לנקוט בצעדים הנדרשים על מנת למנוע או לצמצם את החשש שהשימוש במערכות אלה יגרום לתוצאות שליליות.

במסגרת זו, לדוגמה, מפתחי מערכות מבוססות בינה מלאכותית או המשתמשים בהן יכולים לבחון את השלכות השימוש במערכת באופן עיתי במסגרת **תהליך של הערכת השפעה** (impact assessment) או הערכת סיכון (risk assessment). זאת, בין היתר, ביחס לסיכונים והאתגרים שנדונו לעיל. ביצוע תהליכי הערכת השפעה וסיכון יאפשר להקטין את החשש להתממשות סיכונים לא צפויים או לשימוש במערכת כאשר הנזק הצפוי עולה על התועלת. ככל שהערכות אלו יפורסמו לעיון הציבור או יועברו לגורמי פיקוח, תגדל גם האפשרות לעשות שימוש בכלים משפטיים (לדוגמה, סנקציות רגולטוריות) או חברתיים (למשל חרם צרכני) כדי להטיל אחריות על רקע תוצאות שליליות שהתרחשו עקב השימוש במערכת.¹⁹²

באופן דומה, מפתחי מערכות מבוססות בינה מלאכותית והמשתמשים בהן יכולים **למנות אחראיים ארגוניים** שתפקידם יהיה להבטיח כי הארגון נוקט באמצעים המתאימים כדי להתמודד עם האתגרים והסיכונים שמעורר השימוש במערכות אלה. למשל, הם יבדקו ויוודאו כי הארגון עומד בסטנדרטים מקצועיים מקובלים, כי תהליכי הערכת ההשפעה נעשים כנדרש וכי מקבלי ההחלטות ערים לחששות הנוגעים לשימוש במערכת ובוחנים אותם כחלק מתהליך קבלת ההחלטות בארגון. זאת לדוגמה, בדומה להמלצת הרשות להגנת הפרטיות למנות ממונה על הגנת הפרטיות.¹⁹³

בנוגע להטלת אחריות, היא יכולה להיעשות בשני מישורים עיקריים – במישור המשפטי ובמישור המוסרי-חברתי. אחריות משפטית (legal liability), מתייחסת לשימוש בכלים משפטיים או רגולטוריים שנועדו להתמודד עם מצב שבו יחיד או תאגיד אינו עומד בסטנדרטים המחייבים הנדרשים או מפר חובה שמוטלת עליו. לצד האחריות המשפטית, עלולה להתקיים גם אחריות מוסרית או חברתית (responsibility), כשהמערכת לא פותחה או פועלת כראוי.¹⁹⁴ למשל, אחריות חברתית יכולה להתבטא במוניטין תקשורתי שלילי או בכך שהציבור יפסיק להשתמש במערכת, מכיוון שהוא רואה בה את הסיבה להתרחשותה של תוצאה שלילית ולא מעוניין בהישנותה. ברם, פרק זה יתמקד בסוגיות הקשורות לאפשרות להטיל אחריות משפטית, ובעיקר אחריות נזיקית או פלילית, על רקע השימוש במערכות מבוססות בינה מלאכותית. זאת, לנוכח האתגרים הניצבים בפני המערכת הרגולטורית והמשפטית שמתעוררים בפרט בהקשרים אלה, ומשום חשיבותם של ההסדרים הללו ליצירת מבנה תמריצים הולם ולהגנה על זכויות ואינטרסים.

בהקשר זה, ראוי להקדים לדיון מספר דגשים: ראשית, יודגש כי מטרתו של פרק זה היא אך להציג ברמה כללית של הפשטה את האתגרים שהשימוש במערכות מבוססות בינה מלאכותית עלול לעורר ביחס לאחריות המשפטית, וכן כיווני חשיבה מקובלים שהובעו בעניין זה. בהתאם, אין בפרק זה משום הבעת עמדה ביחס לתחולת הדין הקיים או להסדרים הרצויים.

¹⁹² Bernd Carsten Stahl, *Responsible Innovation Ecosystems: Ethical Implications of the Application of the Ecosystem Concept to Artificial Intelligence*, 62 INT'L J. INFO. MGMT. 102441, 102448 (2022).

¹⁹³ הרשות להגנת הפרטיות, **מינוי ממונה הגנה על פרטיות בארגון ותפקידיו** (2022).

¹⁹⁴ *AI Principles Accountability (Principle 1.5)*, OECD (Sept. 19, 2022), <https://oecd.ai/en/dashboards/ai-principles/P9>.

שנית, יודגש כי ייתכן שחלק מההיבטים הנסקרים להלן זוכים למענה עקרוני במסגרת הדין הקיים בישראל, או באמצעות פרשנות מתבקשת להסדרים קיימים, אולם למיטב הידיעה אין קביעה פוזיטיבית עקרונית בעניינים אלה, הנדונים בספרות, דו"חות ומסמכים מישראל ומהעולם.

שלישית, יודגש כי כמובן ישנו פער יסודי בתכליות ובהסדרים השונים בנוגע לאחריות נזיקת ולאחריות פלילית. אולם לשם הניתוח ולנוחות הקריאה, ומאחר שהסקירה להלן נועדה לבחון את הסוגיות הרחבות הקשורות בהטלת אחריות משפטית בגין תוצאות שליליות שנגרמו אגב שימוש במערכות מבוססות בינה מלאכותית, תיאור הסוגיות הדומות ייעשה ביחס לאחריות נזיקת ופלילית במשותף, תוך התייחסות להיבטים פרטניים של הסוגיות, הדוגמאות וההבדלים בנפרד. מובן כי ניתוח מלא של הסוגיה מחייב לבחון כל אחד מתחומים אלו – התחום הפלילי והתחום הנוזקי – באופן נפרד, ואופן הצגת הדברים להלן הוא אך לשם הניתוח ונוחות הקריאה כאמור.

4.6.1. האתגרים בשימוש בהסדרים הקיימים להטלת אחריות

כפי שיפורט להלן, השימוש במערכות מבוססות בינה מלאכותית עלול לאתגר את היכולת לעשות שימוש בהסדרים הקיימים להטלת אחריות, נזיקת או פלילית. במקרים מסוימים המאפיינים הייחודיים של מערכות אלה לא עולים בקנה אחד עם העקרונות העומדים בבסיס ההסדרים המשפטיים בתחומים הללו, וכפועל יוצא, נוצרים קשיים מהותיים ביחס לאפשרות להטיל אחריות כאשר שימוש במערכת הביא לתוצאה שלילית. נוסף על כך, גם במקרים שבהם ראוי להטיל אחריות, לא אחת השימוש במערכות עלול לעורר קשיים פרוצדוראליים בהוכחת האחריות.

קשיים אלו עשויים להתעורר ככלל על רקע המורכבות של המערכות והיכולת הייחודית שלהן "ללמוד" ולשנות את מאפייני הפעולה בהתאם למידע חדש שהן נחשפות אליו, לעיתים באופן אוטונומי. כתוצאה ממאפיינים אלו, וכפי שיפורט להלן, לא תמיד ברור כיצד לייחס את התוצאות השליליות הנובעות מהשימוש במערכות מבוססות בינה מלאכותית לגורם אנושי, ולמי.

א. קשיים מהותיים בהטלת אחריות

ראשית, קושי שעשוי להתעורר בקשר לשימוש במערכות מבוססות בינה מלאכותית, נוגע לסוגיית ייחוס האחריות לגורם ספציפי, ובפרט לשאלה על מי ניתן להטיל אחריות בגין התרחשות תוצאה שלילית בעקבות השימוש. לא אחת, המעורבות של מערכת מבוססת בינה מלאכותית בהתרחשות, גורמת לכך שלא ניתן בנקל לבחור את הגורם שראוי לייחס לו את האחריות. כשבאופן עקרוני, ניתן לחשוב על מספר גורמים מרכזיים שניתן לייחס להם אחריות, לבדם או ביחד עם אחרים.

בהקשר הנוזקי, שאלה זו רלוונטית הן במסגרת מערכת היחסים שבין המזיק לניזוק, הן (וככל הנראה ביתר שאת) במסגרת מערכת היחסים שבין גורמים שונים העשויים לחוב ביחד. **בהקשר הפלילי**, השאלה רלוונטית בקשר לחבות של חשוד או נאשם בביצוע עבירה לבד או עם אחר.

אפשרות אחת, היא להטיל אחריות על מפתח המערכת. על פניו, ובמיוחד כשמדובר במערכות מורכבות, המפתח הוא בעל הידע כיצד המערכת אמורה לפעול, איך נכון להשתמש בה ומהם הסיכונים שעלולים להתרחש בעקבות השימוש. אולם אפשרות זו אינה חפה מקשיים, שכן מערכות מבוססות בינה מלאכותית ממשיכות להתפתח ולשנות את מתווה הפעילות שלהן גם לאחר שיצאו מידיהם של המפתחים, כמתואר לעיל. בנוסף, האחריות עשויה להתחלק בין מספר גורמים – מפתחי בינה מלאכותית (AI Developers), פורסי בינה מלאכותית (AI Deployers) המטמיעים את מערכת הבינה המלאכותית לשימוש או גורמים אחרים שהיו מעורבים במידת מה בשרשרת ייצור המערכת.

בהקשר זה, קושי נוסף הוא שאחד הגורמים המרכזיים לכך שמערכת לא פועלת כהלכה הוא מאגרי מידע לא מספיק איכותיים ששימשו לאימון שלה. הקושי מתעורר כשמפתח המערכת הוא לא הגורם האחראי על השגת המידע, כאשר לא הייתה אפשרות לדעת כי מאגר המידע לא איכותי מספיק, כאשר האפשרות להרחיב את מאגר המידע כדי לשפר את איכותי היא מוגבלת או כאשר מערכות, מסוגים מסוימים, ממשיכות להתאמן על בסיס מידע שהוזן אליהן גם לאחר סיום הפיתוח (להרחבה בעניין זה ראו פרק "אמינות, עמידות ובטיחות" להלן). במקרים אלו, יהיה מקום לבחון אם להטיל את האחריות על ספק המידע או על המפתח.¹⁹⁵

אפשרות שניה, היא לקבוע כי האחריות תוטל על מפעיל המערכת, קרי על הגורם, ככלל עסקי, שמשמש במערכת לטובת פעילותו. היתרון בהטלת האחריות על המפעיל נובע מכך שבמרבית המקרים המפעיל הוא מי שמפיק את עיקר התועלת מהשימוש במערכת, ובהתאם ראוי גם להיות מי שישא באחריות במקרה שנגרמה תוצאה שלילית בעקבות השימוש.¹⁹⁶ מעבר לכך, הטלת האחריות על המפעיל תתמרץ אותו להבין כיצד המערכת פועלת, לבצע ניתוח עלות-תועלת ביחס לשימוש בה, ולהפעיל אותה תחת מעטפת הולמת למניעת תוצאות שליליות (ראו למשל פרק "מעורבות אנושית" לעיל, שיכולה לשמש ככלי לצמצום החשיפה). החיסרון המרכזי הוא, שלא תמיד יש למפעיל אפשרות אמיתית להבין כיצד פועלת המערכת, וחשוב מכך, פעמים רבות לא יהיו בידיו כלים איכותיים כדי לנקוט באמצעי הזהירות הנדרשים.

אפשרות שלישית, במקרים שבהם מעורב גורם אנושי, וכתלות במידת המעורבות שלו (להרחבה ראו פרק "מעורבות אנושית" לעיל), עשויה להיות נטייה לראות אותו כאחראי לתוצאות השליליות שנגרמו במסגרת השימוש במערכת, שכן מתפקידו לכאורה למנוע התרחשות תוצאות שליליות כאלה. ואולם, כפי שכבר תואר בהרחבה בפרק הנוגע למעורבות אנושית, נטייה זו לא בהכרח תהיה מוצדקת בגלל מגבלות הכרוכות במעורבות האנושית בפעילות מערכות מבוססות בינה מלאכותית. זאת, בפרט, במקרים שבהם אותו גורם אנושי נעדר יכולת אפקטיבית להתערב בפעילותה של המערכת, למשל כשהיא לא ניתנת להסברה או כשהממשק שלה אינו מאפשר בקרה יעילה.

אפשרות רביעית, ונעדרת ביסוס בעת הזו, היא שתפתח גישה לפיה בהקשרים מסוימים ניתן יהיה להכיר באישיות משפטית נפרדת למערכת מבוססת בינה מלאכותית, בדומה לזו שמכירים בה לגבי תאגידים, ולהטיל אחריות על המערכת עצמה. כמובן, אפשרות זו מצריכה לעשות שימוש בקונסטרוקציה משפטית שעלולה לעורר קשיים לוגיים ופרקטיים (למשל, אם למערכת מבוססת בינה מלאכותית עשויות להיות זכויות ככל שהיא עלולה לשאת באחריות משפטית; מה היחס בין המערכת לבין מפתחיה ומפעיליה; ואם וכיצד ניתן לייחס למערכת יסוד נפשי). נוסף על כך, אפשרות זו מעוררת קושי לגבי השתתפות סנקציה על המערכת.¹⁹⁷ מכל מקום, הדבר אינו על הפרק לעת הזו.

¹⁹⁵ Herbert Zech, *Liability for AI: Public Policy Considerations*, 22 ERA F. 147, 148-149 (2021).
European Commission, *Liability for Artificial Intelligence and Other Emerging Digital Technologies*,¹⁹⁶

¹⁹⁶ **(European Commission)** PUBLICATIONS OFFICE, 39-42 (2019) <https://op.europa.eu/en/>
¹⁹⁷ יש שהציעו כי את הסנקציות הנהוגות כיום ניתן להחיל, בשינויים המחויבים, גם על המערכות. כך למשל, הוצע כי במסך תקופת "המאסר" המערכת תושבת ולא ניתן יהיה לעשות בה שימוש; או שכלופה לעבודות שירות, המערכת תפסיק לפעול לטובת המטרות המסחריות שלה ותפעל לטובת הציבור; אפילו עלתה ההצעה שבמקום עונש מוות הקוד של המערכת ימחק לצמיתות. אולם ספק אם הצעות אלו מהוות כלי ענישה אפקטיבי אשר יספק מענה הולם להתמודדות עם מקרים בהם השימוש במערכת גרם לתוצאה שלילית. להצעות בכיוון זה ראו, Gabriel Hallevy, *I, Criminal: When Science Fiction Becomes Reality: Legal Liability of AI Robots Committing Robot - I, Criminal: When Science Fiction Becomes Reality: Legal Liability of AI Robots Committing Criminal Offenses* 22 SYRACUSE SCI. & TECH. L. REP. 1, 29-35 (2010).

אפשרות חמישית, היא, במקרים המתאימים, לבחון את האחריות של גורמים נוספים אשר היו מעורבים בתהליך כניסת המערכת לשימוש או בשימוש שגרם לתוצאה השלילית. כך למשל, במקרים מסוימים ייתכן כי יהיה ראוי להטיל אחריות על היבואן של המערכת, ככל שלא ביצע את הבדיקות הנדרשות לפני שהכניס אותה לשימוש בשוק המקומי; על המפיץ של המערכת, ככל שלא עשה את הבדיקות הנדרשות לפני תחילת ההפצה; או על הבעלים של המערכת, קרי הגורם שלטובתו פותחה המערכת. עם זאת, חשוב לציין כי לעיתים רבות תהיה זהות בין חלק מהגורמים הנ"ל.¹⁹⁸

שנית, עקרון יסוד הוא שככלל אין מטילים אחריות על גורם שאינו נושא באשמה ביחס לתוצאה השלילית שהתרחשה. וכפי שיפורט להלן, השימוש במערכות מבוססות בינה מלאכותית עלול להקשות על האפשרות לייחס אשמה בגין החלטותיהן ופעילותן לגורמים אנושיים. לפיכך הטלת אחריות עלולה לעורר שאלות שונות, והדבר מלמד על האתגר המשמעותי הגלום בהיבט האחריותיות. כך למשל, עלולות להתעורר שאלות בנוגע לבחינת עולות הרשלנות **בדיני הנזיקין** ולאפשרות להוכיח התרשלנות או הפרה של חובת הזהירות במסגרתה. כך גם **במשפט הפלילי**, עלולות להתעורר שאלות הנוגעות, בין היתר, לקיומו של היסוד הנפשי, בהיבטי מודעות וכוונה.

שלישית, קושי דומה לזה שמתעורר בקשר לייחוס אשמה, עשוי לעלות כאשר נדרש לבחון אם קיים קשר סיבתי משפטי בין המעשה לבין התוצאה השלילית. מכיוון שמבחן מרכזי המשמש לבחינת הקשר הסיבתי המשפטי בדיני הנזיקין ובמשפט הפלילי הוא מבחן הצפיות,¹⁹⁹ ומאחר שמאפיין בולט של מערכות מבוססות בינה מלאכותית הוא יכולתן לנתח מידע ולמצוא קשרים מעבר ליכולות אנושיות ולפעול על בסיס קשרים אלה, עלול להתעורר קושי לייחס לגורם אנושי יכולת לצפות את התרחשות התוצאה השלילית, ובהתאם להיווצר אתגר ייחודי בזיהוי קשר סיבתי משפטי.

לבסוף, ישנם אף שיקולי מדיניות שעשויים להיות רלוונטיים לנושא הטלת האחריות. כך למשל, הטלת אחריות בנזיקין עלולה לגרום לאפקט מצנן ביחס לנכונות של מפתחים להשתתף בקידום תוכנות "קוד פתוח", ובכך לפגוע בתועלת החברתית שהן מביאות.

ב. קשיים פרוצדוראליים בהטלת אחריות

לצד הקשיים המהותיים שנסקרו לעיל ביחס ליישום ההסדרים הקיימים בדיני הנזיקין ובמשפט הפלילי במקרים של שימוש במערכות מבוססות בינה מלאכותית, עלולים להתעורר גם קשיים פרוצדוראליים שבכוחם לפגום ביכולת לעשות שימוש בהסדרים הקיימים בתחומים הללו.

ראשית, השימוש במערכות מבוססות בינה מלאכותית עלול לעורר קשיי הוכחה משמעותיים, בפרט כאשר נעשה שימוש במערכות שאינן ניתנות להסברה (להרחבה ראו פרק "הסברתיות" לעיל). מערכות אלה פועלות במקרים רבים באופן בלתי צפוי שעלול להביא לתוצאות שליליות, כאשר הסיבות לכך שמערכת מסוימת לא פועלת כנדרש הן רבות, ולא אחת קשה לדעת מה הוביל לכשל (להרחבה ראו גם פרק "אמינות, עמידות ובטיחות" לעיל). מכיוון שבמשפט האזרחי, ככלל, נטל ההוכחה מוטל על הניזוק, ובמשפט הפלילי נטל ההוכחה מוטל על רשויות האכיפה, הרי שהשימוש במערכות מבוססות בינה מלאכותית צפוי להקשות על היכולת לאכוף במקרים הרלוונטיים. לצד

¹⁹⁸ Proposal for a Regulation of the European Parliament and of the Council Laying Down Harmonised Rules on Artificial Intelligence (Artificial Intelligence Act) and Amending Certain Union Legislative Acts, COM (2021) 206 Final (Apr. 21. 2021).

¹⁹⁹ לדרישת הקשר הסיבתי המשפטי בנזיקין ראו למשל ע"א 576/81 **בן שמעון נ' ברדה**, פ"ד לח(3) 001 (1984); לדרישת הקשר הסיבתי המשפטי במשפט הפלילי ראו למשל ע"פ 8827/01 **שטרייזנט נ' מדינת ישראל**, פ"ד נו(5) 506, פס" 27 לפסק הדין של השופט חשין (2003).

זאת, יצוין כי גם במשפט הפלילי וגם בדיני הנזיקין, לעיתים קרובות נעשה שימוש בחזקות משפטיות, בין היתר לעניין הקשר הסיבתי, מה שעשוי להקל על אכיפת הדינים הרלוונטיים.

שנית, אף כאשר ניתן להוכיח את יסודות התביעה, השימוש במערכות מבוססות בינה מלאכותית עלול להיות כרוך בעלויות ניכרות. במקרים רבים, על מנת לפענח את הקוד של המערכת, להבין איך היא פועלת באופן כללי וכיצד פעלה במקרה הנתון, ולהסיק מכך כי היה כשל בפעילות המערכת שמצדיק הטלת אחריות, נדרשים משאבים רבים והסתמכות על חוות דעת של מומחים טכנולוגיים. גם זאת אפשר, רק במקרים שבהם הקוד גלוי או שהוא נחשף במסגרת ההליך המשפטי, דבר שאינו מובן מאליו בשים לב לכך שמדובר לא אחת בסוד מסחרי. כך גם, ייתכן שיתעורר קושי לאתר מומחה מתאים, שכן חלק גדול מהמומחים מועסקים על ידי חברות מסחריות שעוסקות בתחום.²⁰⁰ כתוצאה מכך, הליך שינוהל בעקבות שימוש במערכת צפוי להיות במקרים רבים ארוך ויקר יותר באופן משמעותי לעומת הליך שגרתי. ייקור ההליך והימשכותו, עלולים ליצור "אפקט מצנן" מפני הגשת תביעה או נקיטה בהליכי אכיפה, גם במקרים שבהם הדבר רצוי מבחינה חברתית.²⁰¹

יודגש כי קשיי ההוכחה האמורים אינם נוגעים רק לביסוס האשמה או לקשר הסיבתי המשפטי, אלא נוגעים גם לחלוקת האחריות בין הגורמים השונים שהיו מעורבים בהתרחשות התוצאה השלילית. בדיני הנזיקין הדבר עלול להקשות על חלוקת האחריות בין המזיקים (במצב של ריבוי מעוולים)²⁰² וכן על חלוקת האחריות בין המזיק לבין הניזוק (במצב של אשם תורם);²⁰³ ובמשפט הפלילי על קביעת סוג האחריות של הצדדים לעבירה,²⁰⁴ וכן על קביעת מתחם העונש ההולם בהתאם לאשם היחסי של החשודים או הנאשמים.²⁰⁵ האחריות בגין התוצאה השלילית שנגרמה עקב השימוש במערכות מבוססת בינה מלאכותית יכולה להתחלק בין מספר גורמים שחלקם מנויים לעיל. במקרים מסוימים חלוקת האחריות ביניהם תהיה מורכבת להוכחה, ובין היתר, תלויה גם בשיקולי מדיניות. חשוב לציין כי אמנם הקושי בחלוקת האחריות קיים גם כשלא נעשה שימוש במערכות מבוססות בינה מלאכותית, אך כאשר נעשה שימוש במערכות אלה האתגר מתעצם.

4.7. פרטיות

במקרים רבים פיתוח ושימוש במערכות מבוססות בינה מלאכותית מצריך שימוש בנתונים רבים, לעיתים בהיקפים עצומים. בפרט, נעשה שימוש בנתונים רבים לצורך פיתוח מודלים (קרי, לשם למידה), לצורך בדיקת המודלים לאחר שפותחו לשם תיקונם או דיוקם, או לצורך עיבוד הנתונים הפרטניים לשם הפקת המלצה, תחזית או החלטה קונקרטית. כאשר הנתונים המהווים בסיס לפעילות זו מהווים מידע אישי, המוגן על ידי חוק הגנת הפרטיות, התשמ"א-1981 (להלן: "חוק הגנת הפרטיות"), הוראות החוק חלות על הפעילות. כך, במקרים רבים המידע המוזן לתוך המערכת לצרכים אלו (הקלט) וגם התוצר, קרי, ההמלצה, התחזית או ההחלטה המתגבשות כתוצאה משימוש בבינה מלאכותית (הפלט) עשויים לכלול מידע אישי. מכאן, שבכל אותם מקרים בהם נעשה שימוש במידע אישי בכל אחד מהשלבים שתוארו לעיל, יחולו הוראות חוק הגנת הפרטיות.²⁰⁶

²⁰⁰ אסדרת השימוש ברכבים אוטונומיים - פנייה לקבלת עמדת הציבור לעיל ה"ש 50, בעמ' 9 (2021).

²⁰¹ European Commission, לעיל ה"ש 196, בעמ' 20-21.

²⁰² ס' 84 לפקודת הנזיקין.

²⁰³ ס' 68 לפקודת הנזיקין.

²⁰⁴ ס"י ב' לחוק העונשין התשל"ז-1977.

²⁰⁵ ס' 40 לחוק העונשין.

²⁰⁶ השימוש במידע אישי יכול שיעשה במגוון של אופנים וסוגי קבצים, לרבות שמירה, עיבוד והעברה של קבצי טקסט, תמונות, קבצי וידאו, הקלטת קול, מידע ביומטרי ועוד.

בעת קביעת אופן תחולת החוק והחובות החלות על מי שמבצע את הפעילות הכרוכה באיסוף או עיבוד מידע אישי, יש לתת את הדעת לאתגרים הייחודיים לפרטיות. אתגרים אלה נובעים מהמאפיינים הטכנולוגיים של "נתוני עתק" (Big data) ובינה מלאכותית.

4.7.1. האתגרים לפרטיות בהם מתאפיין השימוש במערכות מבוססות בינה מלאכותית

כפי שיפורט להלן, פיתוח ושימוש במערכות מבוססות בינה מלאכותית מעוררים מספר אתגרים ייחודיים הנוגעים לצורך להבטיח את ההגנה הנדרשת לזכות לפרטיות.

ראשית, השימוש במערכות אלה עלול לאתגר את היכולת לעמוד **בעיקרון צמידות המטרה**. אחד מעמודי התווך של דיני הגנת הפרטיות, הקבוע במפורש בחוק הגנת הפרטיות, הוא עקרון צמידות המטרה. בהתאם לעקרון זה, ניתן לאסוף מידע אישי ולהשתמש בו אך ורק למטרה אשר לשמה נמסר המידע, אלא אם נושא המידע עצמו הסכים לשימושים נוספים במידע אודותיו או אם יש הסכמה חוקית המאפשרת זאת.²⁰⁷ עקרון זה מאתגר בשלב פיתוח הבינה המלאכותית שכן במקרים רבים נדרש לעשות שימוש במידע שנאסף למטרות אחרות. כך למשל, דוגמא מובהקת לאי-התאמה של המטרה, היא מקרה של שימוש במידע שמתקבל לטובת הפעלה קולית לצורך אימון אלגוריתם לאיתור בעיות במצב רפואי של בני אדם. היכולת לקבל את הסכמתם של נושאי המידע לצורך השימוש הנוסף המתבקש, מעוררת קשיים מעשיים שונים, כגון הקושי לשוב ולפנות אל נושאי המידע הרלוונטיים ולקבל הסכמה מדעת לכך, בפרט כשמדובר במידע שנאסף בעבר הרחוק ובנושאי מידע רבים.

שנית, השימוש במערכות מבוססות בינה מלאכותית עלול לאתגר את היכולת לעמוד **בדרישת השקיפות ובחובת היידוע**. דיני הגנת הפרטיות נועדו בין היתר להבטיח את היכולת של נושאי המידע לשלוט במידע המתייחס אליהם, ולקבוע כיצד ייעשה בו שימוש. אחת הזכויות המרכזיות שחוק הגנת הפרטיות מקנה לאדם היא זכות יידוע שמאפשרת לקבל פירוט בנוגע למטרות ולשימושים שעבורם האדם מתבקש למסור את המידע אודותיו, למי יימסר המידע, והאם חלה עליו חובה חוקית למסור את המידע, או שמא מסירת המידע תלויה ברצונו או בהסכמתו.²⁰⁸ חובה חוקית זו חלה בכל מקרה של פניה לאדם לקבלת מידע אודותיו, גם כשאיסוף המידע נעשה מכוח הסכמה בדין וגם כאשר הפניה לאדם נעשתה בעקבות בקשתו לקבלת שירות מסוים.²⁰⁹ שקיפות משפרת גם את האפשרות של נושא המידע לממש את זכות העיון במידע, המוקנית לו בחוק הגנת הפרטיות,²¹⁰ באופן משמעותי ואפקטיבי, וכן את אחריותיות הארגון כפי שיוסבר בהמשך. אולם, במקרים רבים בהם נעשה שימוש במערכות מבוססות בינה מלאכותית שאינן ניתנות להסברה (ראו פרק "הסברתיות" לעיל), עלול להיווצר קושי בעמידה בדרישה זו. זאת מכיוון שחלק מההיבטים הנוגעים לשימוש במערכת לא יהיו ידועים למפעיל, ובהתאם לא ניתן יהיה למסור אותם לנושא המידע.

שלישית, השימוש במערכות מבוססות בינה מלאכותית עלול להוביל לכרסום **בעיקרון ההסכמה מדעת**. חוק הגנת הפרטיות קובע כי הסכמת אדם לאיסוף ולשימוש במידע אישי אודותיו, נדרשת להיות "הסכמה מדעת".²¹¹ הדרישה להסכמה מדעת מחייבת את מבקש המידע להציג לאדם נתונים

²⁰⁷ ראו ס' 92 ו-8 (ב) לחוק הגנת הפרטיות.

²⁰⁸ ראו ס' 11 לחוק הגנת הפרטיות; ראו גם עמדת הרשות להגנת הפרטיות בעניין חובת יידוע במסגרת איסוף ושימוש במידע אישי, לעיל ה"ש 27.

²⁰⁹ שם, בפס' 4-5.

²¹⁰ ס' 13 לחוק הגנת הפרטיות.

²¹¹ בעקבות תיקון שנערך לחוק הגנת הפרטיות בשנת 2007, "ההסכמה" בסעיף 3 לחוק מוגדרת כ-"ההסכמה מדעת, במפורש או מכללא".

בנוגע לאיסוף המידע עליו והשימוש בו, כך שבפני נושא המידע תעמוד תמונה מלאה בטרם יחליט האם להסכים למסירת המידע אודותיו.²¹² הקושי בעמידה בדרישת השקיפות שתואר לעיל, והקושי לקיים את חובת היידוע בשל אופיין של חלק ממערכות בינה מלאכותית שאינן ניתנות להסברה, עלולים, במקרים מסוימים, לאתגר את היכולת לקבל ההסכמה מדעת מנושא המידע.²¹³

רביעית, השימוש במערכות מבוססות בינה מלאכותית עלול ליצור מתח מובנה בין הצורך בנתוני עתק לבין מחיקת מידע עודף בהתאם לעיקרון צמצום המידע. האפשרות לקבל גישה ל"נתוני עתק" (Big Data), והיכולת לנתח מידע בהיקפים גדולים, טומנת בחובה פוטנציאל רב לקידום ופיתוח השימוש במערכות מבוססות בינה מלאכותית. שכן, כאמור לעיל, בשלב פיתוח המודל המערכת "מתאמנת" על המידע, כאשר בשלב איסוף המידע לא בהכרח ניתן לדעת מהן התובנות והתועלות שהמערכת תדע להפיק מהמידע הנאסף. המשמעות היא שאיסוף וניהול כמויות גדולות של מידע יכולים לסייע בפיתוח המערכות, אף שחלק מהמידע הנאסף כלל אינו חיוני למתן השירות עצמו. מנגד, איסוף ושמירה של מידע אישי עודף מגדילים את הסיכון לפגיעה בפרטיות, שכן המידע עשוי לדלוף ולהיחשף לגורמים לא מורשים. משכך, מחייבות תקנות הגנת הפרטיות (אבטחת מידע), התשע"ז-2017, כל בעל מאגר מידע אישי לבחון אחת לשנה אם הוא מחזיק במאגר מידע עודף, שאינו דרוש לו עוד.²¹⁴ מכאן, שכאשר המידע בו נעשה שימוש לצורך המערכת הוא מידע אישי, מתקיים מתח מובנה בין עיקרון צמצום המידע, הקובע כי יש לאסוף ולהשתמש אך ורק במידע הרלוונטי וההכרחי לשם הגשמת המטרה אשר לשמה המידע נאסף, לבין הרצון לספק למערכת מידע בהיקפים גדולים ככל הניתן. כמו כן, שמירה לאורך זמן של מידע מעלה את השאלה האם יש לאדם יכולת לשנות את המידע שנאסף עליו, האם יש מקום להעניק לפרט זכות לבקש את מחיקתו של המידע האישי שנאסף בעניינו או להתנגד להמשך עיבוד המידע שנאסף לגביו, והאם זו היא זכות שניתנת ליישום – שכן מערכות בינה מלאכותית ממשיכות את תהליך הלמידה שלהן גם בשלב שלאחר פיתוח המודל.

חמישית, ניתן לעשות שימוש במערכות מבוססות בינה מלאכותית לטובת הסקת מידע רגיש. יכולת הסקת המסקנות וההיקשים של שימוש בנתוני עתק ולמידת מכונה מאפשרת לשלב ולנתח מרכיבי מידע שונים, וכך להסיק מסקנות הכוללות מידע רגיש ממידע לא רגיש (למשל מצב בריאותי מהרגלי רכישה). לדוגמה, מחקרים מראים, כי בנסיבות מסוימות בינה מלאכותית מסוגלת לזהות ממידע ביומטרי, כגון מבנה עצמותיו או אופן התנהגותו של אדם, דפוסים החוזים את נטיותיו והתנהגותו של אותו אדם ואשר משליכים על מצבו הבריאותי, נטייתו המינית ועוד. יכולת זו מעלה שאלה האם יש מקום להגביל את סוג המידע עליו רשאיות חברות לבסס החלטות.

שישית, יש חשש שיעשה שימוש במערכות מבוססות בינה מלאכותית לשם זיהוי חוזר של מידע שעבר התממה. בשל היכולת של מערכות מבוססות בינה מלאכותית לעבד מגוון רחב של נתונים ממגוון מקורות, השימוש בהן עלול להגביר את הסיכון לזיהוי נושאי המידע במערכי נתונים, לרבות

²¹² על היותה של חובת היידוע חלק מעקרון ההסכמה מדעת, עס"ק (ארצי) 7541-04-14 הסתדרות העובדים הכללית החדשה מרחב המשולש הדרומי - עיריית קלנסווה, פס' 65-66 ו-78 לעמדת היועץ המשפטי לממשלה, פס' 144 לפסק הדין של השופט איטח (15.3.17).

²¹³ ראו למשל את הסקירה המקיפה בתחום הכלכלה ההתנהגותית: Acquisti, Alessandro, Brandimarte, Laura, & George Loewenstein, *Secrets and Lies: The Drive for Privacy and the Difficulty of Achieving It in the Digital Age*, 30 J. OF CONSUMER PSYCH. 736 (2020); Daniel J. Solove, *The Limitations of Privacy Rights*, 98 NOTRE DAME L. REV. (2023) https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3688497.

²¹⁴ תקנה 2(ג) לתקנות הגנת הפרטיות (אבטחת מידע), תשע"ז-2017. ראו גם טיוטת מסמך מדיניות של הרשות להגנת הפרטיות בנושא צמצום מידע. המסמך זמין כאן.

כאלו המשמשים לאימון מערכות, אשר עברו תהליכי התממה (אנונימיות). החשש בהקשר זה גובר ככל שמדובר בשילוב של מידע ממקורות שונים ובהיקפים גדולים.

שביעית, השימוש במערכות מבוססות בינה מלאכותית עלול לעורר קושי **בהגדרת תפקידים ותחומי אחריות ביחס למידע אישי**. כאמור, מידע אישי עשוי לשמש בשלב פיתוח מערכת הבינה המלאכותית, בשלב בדיקת המערכת ובשלב השימוש בה. לעיתים בשלב פיתוח המערכת המפתחים מסתייעים במאגרים של צד שלישי, הכוללים מידע אישי, לצורך אימון המערכת. במקרים אלו הגדרת תחומי האחריות של כל אחד מהצדדים עלולה להיות סבוכה ומורכבת.

יצוין בהקשר זה כי ניתן להבחין בין שלבי הפיתוח והשימוש השונים במערכת. כך למשל, בכל הנוגע לעוצמת הפגיעה בפרטיות שעלולה להיווצר, כאשר נעשה שימוש במידע פרטי במסגרת תהליך ה"אימון" של מערכת בינה מלאכותית, הפגיעה תהיה פחותה יותר ככל שבסוף התהליך המידע הפרטי יימחק.

מכלל האמור לעיל, עולה כי בשים לב למאפייני טכנולוגיית הבינה המלאכותית, אשר עושה שימוש במידע בהיקף עצום, הן לשם פיתוחה ולמידתה והן במהלך השימוש בה, הפרטיות מהווה אתגר משמעותי, שיש להידרש אליו במסגרת הפיתוח והשימוש במערכות מבוססות בינה מלאכותית, וגם במסגרת גיבוש מדיניות רגולטורית בתחום זה.

5. חלק חמישי: דרכי התמודדות עם סוגיות ואתגרים אלה

פרק זה סוקר באופן כללי דרכי פעולה אפשריות להתמודדות עם הסוגיות, הסיכונים והאתגרים שנסקרו בפרק הקודם, כפי שהן משתקפות מהספרות ומדו"חות שונים, מישראל ומהעולם. כפי שצוין לעיל, גם בנוגע לדרכי ההתמודדות הסקירה להלן אינה מתיימרת למצות את כל דרכי הפעולה האפשריות, או לקבוע קביעות לגבי הדין החל או זה שראוי שיחול, אלא לעסוק בדרכי פעולה עיקריות ולבחון את יתרונותיהן וחסרונותיהן המרכזיים בהתאם למתאפשר, וזאת כפתח לבחינה עתידית. בדומה לאמור לעיל, גם פרק זה מיועד לשמש כבסיס להמשך דיון.

בפרט, עבור גורמי הממשלה והרגולטורים הרלוונטיים, ולאור ההצעה להלן בעניין "מיפוי השימושים בבינה מלאכותית והאתגרים הנלווים להם בענפים מוסדרים", פרק זה יכול לסייע בהבנה ומיפוי של דרכי הפעולה האפשריות להתמודדות עם הסוגיות, הסיכונים והאתגרים הכרוכים בשימושים קונקרטיים במערכות מבוססות בינה מלאכותית בענף שעליו הם אמונים.

5.1. אפליה

כאמור לעיל, ישנו קושי ממשי להתמודד עם החשש שהשימוש במערכות מבוססות בינה מלאכותית יביא להנצחת ואף להחרפת בעיית האפליה. אולם, ישנן מספר דרכים שמאפשרות להתמודד, ולו באופן חלקי, הן עם האתגרים הטכנולוגיים והן עם האתגרים המשפטיים, ובכלל זה:

א. מאגרי מידע מגוונים ואיכותיים

דרך מרכזית אחת להתמודדות עם אפליה, היא לדרוש כי מערכת מבוססת בינה מלאכותית "תתאמן" על בסיס מאגרי מידע רחבים ומגוונים, שיאפשרו לה להתחשב במאפיינים של כל קבוצות האוכלוסייה הרלוונטיות, תוך שמירה כי הנתונים לא משקפים אפליה היסטורית. ניתן לקבוע סטנדרטים שונים לגיוון הנתונים ולאיוכותם, בין היתר, באמצעות אתיקה, רגולציה או תקינה. הדבר מחייב במקביל פתרון לבעיה משפטית של מערכות הדינים המגבילות שימוש במידע. פתרון חלקי לבעיה זו נמצא בחוות דעת מקיפה שנכתבה במשרד המשפטים בנוגע לשימוש בתכנים מוגנים בזכויות יוצרים, שקובעת כי ברוב המקרים שימוש כאמור יהווה שימוש הוגן בתכנים.²¹⁵ עם הזמן, יהיה צורך לייצר פתרונות משפטיים שיאפשרו שימוש בפרטי מידע גם על רקע הגבלות בתחומי משפט אחרים (דיני הגנת פרטיות, חסיונות וכדומה). מעבר לכך, ניתן לבחון לעודד ואף לחייב במקרים מתאימים שיתוף נתונים בין גופים על מנת לייצר מאגרי נתונים מגוונים ואיכותיים.

ב. ביצוע בדיקות מקדימות ועיתיות

דרך נוספת היא לדרוש כי לפני תחילת השימוש במערכת המבוססת על בינה מלאכותית ייבחנו החלטותיה ופעולותיה, ובפרט אם היא נוהגת בצורה דומה ביחס לאנשים בעלי מאפיינים שונים. במילים אחרות, ניתן לדרוש כי בשלב של התנסות המערכת, היא תיבחן על קבוצות אוכלוסייה שונות באופן שיוודא שהיא אינה מפלה לרעה. ככל שמדובר במערכת שניתנת להסברה, ניתן אף לבחון את האינטראקציה בין הפרמטרים שהמערכת משקללת לבין מאפייני שיוך קבוצתי חשודים בתחום בו היא פועלת, בניסיון לאתר משתנים קורלטיביים למאפיינים מפלים (משתני proxy). גם לאחר שהמערכת נכנסת לפעולה, דרישת בדיקה עתית של הנתונים והחלטות המערכת עשויה לאתר בשלב מוקדם נטייה לאפליה. במסגרת זו, ניתן אף לדרוש ביצוע של תהליך הערכת סיכונים

²¹⁵ משרד המשפטים גילוי דעת: שימושים מוגנים בזכויות יוצרים לצורך למידת מכונה (2022).

שמטרתו לבחון מה רמת החשש שהשימוש במערכת יוביל לאפליה ומה מידת הפגיעה המשוערת במקרה שהחשש יתממש. ייתכן שבמקרים המתאימים ראוי שהבדיקה והפיקוח ייעשו על ידי גורמים חיצוניים (פרטיים או ציבוריים) שיוכלו להעריך את מידת החשש מאפליה.

מבחינה משפטית ורגולטורית, ישנן מספר אפשרויות ליישם את הדרישה האמורה לגיוון. דרך אחת היא לחייב חברות להצהיר כי בדיקות כאמור בוצעו או לחייבן לפרסם את קבוצות האוכלוסייה לגביהן נבחנה המערכת. כשמדובר בחברות שפעילותן כפופה לחובת רישוי העושות שימוש במערכות מבוססות בינה מלאכותית שיוצרות סיכון לאפליה, ניתן אף להתנות את הרישיון הרגולטורי בקיומן של בחינות גיוון כאמור. לחלופין, ניתן ליצור "נמל מבטחים" (safe harbor) משפטי, עבור חברות שיתבעו בגין החלטה מפלה של החברה – במקרה שהן יכולות להוכיח שהמערכת אומנה ונוסתה על מאגר נתונים מגוונים – וזאת במטרה לתמרץ אותן לנקוט אמצעי הזהירות הנדרשים.

ג. קביעת דרישה להסברתיות

אמצעי נוסף שבאמצעותו ניתן להתמודד עם החשש מפני אפליה הוא קביעה של דרישה להסברתיות (ראו פרק "הסברתיות" להלן). קבלת הסבר בנוגע למערכת ולתהליך קבלת ההחלטה על ידי המערכת במקרה הקונקרטי, תסייע לדעת האם הביאה בחשבון מאפיינים מפלים או משתנים קורלטיביים אליהם (ככל שניתן לאתרם), ולנקוט באמצעים המתאימים – טכנולוגיים או משפטיים – כדי להתמודד עם האפליה. כך לדוגמה, ניתן לבטל טכנית אפשרות להתחשב במשתנה proxy מסוים; לאכוף את דיני האפליה על ידי הרגולטורים הרלוונטיים; או לאפשר אכיפה פרטית בגין אפליה על ידי מושאי ההחלטה, ככל שייחשפו לנימוקים שעמדו בבסיס ההחלטה, לרבות באמצעות קביעת זכות לנושא המידע לקבל את ההסבר. עם זאת, יש לקחת בחשבון גם את הקשיים הנלווים לדרישה להסברתיות, כפי שהוצגו לעיל.

ד. עריכת מבחנים תוצאתיים

לבסוף, אמצעי שיכול לסייע בהתמודדות עם החשש מפני אפליה במערכות מבוססות בינה מלאכותית, הוא שימוש מוגבר במבחנים תוצאתיים על מנת לבחון אם התרחשה אפליה במקרה נתון, וזאת כמנגנון משלים או כחלופה אפשרית לצורך בהסברתיות.

כך למשל, ניתן לקבוע חזקה כי השימוש במערכת גורם לאפליה, ולכן פסול, כשיש שוני העולה על סף מסוים בין אנשים מקבוצות שונות באוכלוסייה. אפשרות נוספת, היא להתבסס על בחינה תוצאתית לצד יצירת "נמל מבטחים" שישפק הגנה במקרה של שימוש במערכות שנבדקו או במערכות ניתנות להסברה. חזקה או נמל מבטחים כאמור עשויים לתמרץ חברות לבכר שימוש במערכות כאלה כאמצעי להגנה מפני תביעות. לחלופין, ניתן לעשות שימוש בכלים משפטיים על מנת להבטיח יחס שווה לגורמים שונים, למשל על ידי היפוך נטלי הראיה או הרחבת חובות הייצוג ההולם.²¹⁶

5.2. מעורבות אנושית

לאור האמור לעיל, ברי כי לצד היתרונות המשמעותיים שניתן להשיג באמצעות מעורבות אנושית, ישנו צורך להתאים את המעורבות לנסיבות שבהן נעשה השימוש במערכת, על מנת לוודא כי התועלת שבמעורבות האנושית משמעותית. לעניין זה חשוב להדגיש כי אין מדובר רק על השאלה

²¹⁶ ראו למשל בעניין ייצוג הולם לנשים: מרכז המידע והמחקר של הכנסת חקיקה בישראל הקובעת ייצוג הולם לנשים (2005).

מתי צריכה להיות מעורבות אנושית, אלא נדרש גם לחשוב כיצד נכון לעצב את האינטראקציה בין הגורם האנושי לבין המערכת על מנת לנצל את היתרונות היחסיים של כל צד ולמנוע הטרות.

סוגיית הצורך במעורבות אנושית רלוונטית כמעט לכל השימושים בבינה מלאכותית. כך לדוגמה, בתחום הבריאות, עולה השאלה אם יש צורך שרופא יבקר דיאגנוזה או המלצות לטיפול שסיפקה מערכת מבוססת בינה מלאכותית; בתחום הפיננסי, עולה שאלה אם יש צורך בבקרה על שימוש במערכת בינה מלאכותית לייעוץ השקעות באופן שלא תתקבל החלטה על השקעה בהתבסס על מערכת מסחר אלגוריתמי בלבד, או אם יש צורך שפקיד יקבל החלטה על מתן או סירוב אשראי כדי שלא לחשוף את הגוף המלווה לסיכונים אשראי לא מקובלים או לשלול לשווא אשראי מלווה פוטנציאלי; בתחום התעסוקה, עולה שאלה אם יש צורך שהמעסיק יפקח על שימוש במערכת על מנת להעסיק או לפטר עובדים. עם זאת, יש להבחין בין סוגי המעורבות השונים, ולבחון איזה סוג של מעורבות אנושית נדרש (אם בכלל) ביחס לכל תחום ושימוש פרטני בבינה מלאכותית.

על מנת לקבוע את הדרישה למעורבות אנושית ביחס לכל תחום ושימוש בבינה מלאכותית, ואת היקפה במקרים שבהם היא נדרשת, יש לתת את הדעת למגוון הנסיבות והמאפיינים. בכלל זה, והגם שאין מדובר ברשימה ממצה, ניתן להתחשב במאפיינים אלה: ראשית, אופי ההחלטה והשפעתה – ככל שהסיכון שעלול להיווצר בעקבות ההחלטה גדול יותר, ובהתאם גובר החשש לפגיעה משמעותית בזכויות יסוד או באינטרסים ציבוריים ולהשפעה לרעה על מושא ההחלטה, כך מתעצמת, למצער בעת הזו, ההצדקה לשקול דרישה למעורבות אנושית באופן משמעותי יותר. כך למשל, ההצדקה למעורבות אנושית עשויה להיות גבוהה במיוחד כשמדובר בהחלטות בלתי הפיכות (כמו בתחום הרפואה).²¹⁷ שנית, אפקטיביות המעורבות האנושית – כך למשל, כאשר מדובר במערכות שאינן ניתנות להסברה, כפי שתואר לעיל, היכולת של הגורם האנושי להתערב באופן אפקטיבי בהחלטת המערכת נפגעת. שלישית, ובהנחה שיימצא כי נדרשת מעורבות אנושית, יש לבחון אם המעורבות האנושית צריכה להיות מראש (באמצעות קבלת ההחלטה הסופית על ידי אדם, או כשהוא יכול להטיל וטו בזמן אמת), או שניתן להסתפק במנגנון של ערעור בדיעבד. רביעית, יש מקום לבחון את השפעת הדרישה למעורבות אנושית – דרישה כזו, במיוחד אם היקפה יהיה משמעותי ויצריך מעורבות אנושית בהחלטות מרובות, עשויה ליצור "צוואר בקבוק" אנושי ובכך לסכל את התועלת שניתן להפיק מיכולות הבינה המלאכותית.

מכל מקום, אין באמור לעיל כדי לקבוע אם ובאלו מקרים יש לדרוש מעורבות אנושית, אלא להצביע על אפשרויות וכיווני מחשבה בלבד. יחד עם זאת, נראה כי במקרים שבהם יימצא כי יש מקום לדרוש מעורבות אנושית כלשהי בתהליך קבלת החלטה, יש לבחון אם נדרש גם להסדיר בקרה על אותה מעורבות אנושית, כדי לוודא שתצליח להגשים את התכלית שלשמה נדרשה. כאשר מצד אחד, ניתן לבצע בקרה לגבי המקרים שבהם הגורם האנושי החליט להתערב ולסטות מההחלטה של המערכת, ולבחון לדוגמה האם הסטייה הייתה מוצדקת; מה הגורמים שהובילו לסטייה (כשל נקודתי או בעיה רוחבית במערכת); והאם הגורם האנושי הסתמך על המידע הרלוונטי במסגרת החלטתו. מצד שני, אפשר לבצע בקרה לגבי המקרים שבהם הגורם האנושי החליט שלא להתערב, תוך בחינה, בין היתר, האם הוא הפעיל שיקול דעת משמעותי; בחן את הנתונים הרלוונטיים; או

²¹⁷ Brennan-Marquez & Henderson, לעיל ה"ש 122, בעמ' 146-147.

הצליח להבין את האופן שבו המערכת הגיעה להחלטה ונימוקה. כמו כן, יהיה מקום לשקול בחלופ זמן את מידת האפקטיביות של המעורבות האנושית ואת מידת תרומתה לתהליך.

5.3. הסברות

כאמור לעיל, להחלטת הדרישה להסברות ביחס לשימוש במערכות המבוססות על בינה מלאכותית יתרונות משמעותיים. ברם, לנוכח האתגרים הקשורים להחלטת דרישה כזו, ייתכן שיש מקום לעשות זאת באופן יחסי ובהתאם לנסיבות השימוש במערכת. באופן דומה, ניתן להתמודד עם חלק מהקשיים שתוארו לעיל באמצעות צמצום היקף ההסבר שיינתן (ראו לעיל בפרק הקודם בחלק העוסק ב"הסברות"), או הגבלת הגורמים שיהיו חשופים להסבר.

במסגרת בחינת האפשרות להחיל דרישה להסברות באופן יחסי ובהקשר מסוים, יהיה מקום להתחשב בזהות הגוף המשתמש במערכת המבוססת על בינה מלאכותית והגוף המושפע מהשימוש, כמו גם באופי ההחלטה עליה מוחלט חובת ההסברות והשפעתה. לגבי **הגוף המבצע את השימוש במערכת**, הבחנה אפשרית היא בין שימושים שמבוצעים על ידי גופים פרטיים בעלי מאפיינים ציבוריים, לבין שימושים שנעשים על ידי גורמים פרטיים לחלוטין – ככל שהגוף נוטה יותר לכיוון הציבורי, עשויה להתחזק הצדקה להחיל לגבי דרישה להסברות. לעומת זאת, ביחס לגופים פרטיים, ייתכן שיהיה מקום להחיל את הדרישה להסברות באופן רחבי או באופן סקטוריאלי רק בתחומים מסוימים – כך למשל, דרישת ההסברות רלוונטית במיוחד כשמדובר בגופים פרטיים שהחלטות שלהם מפוקחות באופן הדוק על ידי רשויות השלטון לאור השפעתן המשמעותית על חיי הפרט, למשל בתחום הבריאות, הבנקאות או הביטוח.²¹⁸ בנוגע **לאופי ההחלטה והשפעתה**, ככל שהסיכון שעלול להיווצר בעקבות ההחלטה גדול יותר, ובהתאם יש חשש לפגיעה משמעותית יותר בזכויות ולהשפעה לרעה על מושא ההחלטה, כך תגדל חשיבותה של הדרישה להסברות. למשל, יש הצדקה חזקה יותר לדרוש הסבר לגבי מערכת שקובעת אם אדם יוכל לקבל הלוואה או להתקבל למקום עבודה, לעומת אלגוריתם שנועד להתאים תוכן וידיאו או מוזיקה לצרכן.²¹⁹

בעניין הגורמים החשופים להסבר, מצד אחד, ככל שהדרישה להסברות תהיה רחבה יותר, כך הפיקוח על פעילות המערכת צפוי להיות טוב יותר, והיתרונות שתוארו לעיל עשויים להתגשם. לדוגמה, החלטת הדרישה להסברות כלפי כלל הציבור, כך שתשתרע אף מעבר למי שמושפעים מההחלטה, תאפשר לקיים שיח ציבורי רחב במשמעויות השימוש במערכת ולפקח עליה (כמובן שדרישה כזו תהיה אפקטיבית יותר אם ההסבר יהיה נגיש וברור לכל הציבור ולא רק ל"יודעי ח"ן"). באופן דומה, הצבת דרישה להסברות כלפי הרגולטור הרלוונטי תאפשר פיקוח אפקטיבי מצדו. על כן, יש מקום כי בגופים הרגולטוריים הרלוונטיים יהיו אנשי מקצוע בעלי המומחיות המקצועית וההבנה הטכנולוגית הנדרשת על מנת לבחון את ההסברים המתקבלים. **מצד שני**, דרישת הסברות צרה יותר תקל על מפתחי המערכת ומשתמשיה, לרבות במישור ההגנה על הקניין

²¹⁸ Bibal, Lognoul, de Streel & Frénay, לעיל ה"ש 129, בעמ' 151-153. לדוגמאות להחלטה של חובת הנמקה במגזר הפרטי והמלצה להחיל חובת הנמקה רחבה על החלטות שמקבלים גופים פרטיים על בסיס ניתוח טכנולוגי של נתוני מידע ראו מעין פרל "פרטיות, שליטה ופיקוח בעידן של נתוני עתק חובת הנמקה על החלטות אלגוריתמיות" **משפט, חברה ותרבות** ב' 167, 188-192 (2019).

²¹⁹ 44, 167, 188-192 (2019). *Miriam C Buiten, Towards Intelligent Regulation of Artificial Intelligence*, 10 EUR. J. RISK REG. 57-58 (2019).

הרוחני שלהם,²²⁰ ואולי אף תסייע לצרכנים שלרוב יתקשו להבין הסברים טכניים ומפורטים יתר על המידה.

בפועל, בקצה אחד של הסקאלה, יש מי שהציעו לחייב משתמשים במערכות בינה מלאכותית לפרסם דו"ח הערכת השפעה ביחס לכל שימוש באלגוריתם (algorithm impact statement – AIS), בדומה לאופן שבו מתפרסמים דו"חות ביחס להשפעה סביבתית או רגולטורית (Regulatory Impact Assessment – RIA) של פעולות מסוימות. בהתאם לגישה זו, יהיה מקום לפרסם לעיון הציבור מהי המערכת בה מבוצע השימוש; מה היה תהליך הפיתוח שלה; איך היא עובדת (בהתאם לרמות שתוארו לעיל); ומה צפויה להיות ההשפעה שלה. בקצה האחר של הסקאלה, ניתן לסייג את הזכות לקבל הסבר רק לפרט מושא ההחלטה של המערכת, ואפילו רק למקרים שבהם הוא זקוק לכך לשם הליך משפטי או כשנשלל ממנו שירות חיוני בעקבות החלטת המערכת.²²¹

לבסוף, יש לתת את הדעת על כך שלא אחת מפתחי המערכת ומשתמשיה הם בעלי תמריץ משמעותי להשתמש במערכות הניתנות להסבר. כך למשל, לבנקים עשוי להיות תמריץ כלכלי להבין מדוע מערכת אשראי מאשרת או מסרבת לתת הלוואה ללקוח, כדי לוודא שההחלטה לא מגלמת סיכונים בלתי סבירים או מפסידה לשווא את הלקוח. באופן דומה, לחברות השקעות יכול להיות תמריץ כלכלי להבין מדוע מערכת המליצה על השקעה מסוימת, כדי לוודא שההמלצה טובה או ליישם את התובנות ביחס להשקעות אחרות. במקרים אלה, לאור תמריצייהם, מפתחי המערכת ומשתמשיה עשויים ליזום הסברתיות ולרצות בה, גם בהיעדר דרישה רגולטורית להסברתיות. ואולם, בחלק מהמקרים לא בהכרח יהיה לגורמים אלה תמריץ מספק להשקיע בקבלת הסבר, או תמריץ למסור את ההסבר שהתקבל לגורם חיצוני, ואז עשוי להתעורר צורך בהתערבות רגולטורית.

5.4. גילוי

על רקע האמור לעיל, ובפרט במקרים שבהם למפעיל מערכת מבוססת בינה מלאכותית אין תמריץ מספק לגלות למשתמש על אינטראקציה עם בינה מלאכותית, עשויה להיות הצדקה להתערבות רגולטורית בדמות דרישה כי בנסיבות מסוימות יהיה עליו לגלות על השימוש במערכת מסוג זה.

בבחינה השאלה אם ראוי להציב דרישה לגילוי עשויים להיות מספר שיקולים רלוונטיים. בכלל זה, ראשית, ההשפעה על מצב משתמש הקצה. ככל שהשימוש במערכת משפיע בצורה יותר משמעותית על זכויותיו, כך ייתכן שיהיה מקום להכיר באינטרס חזק יותר שלו לדעת כי הוא בא באינטראקציה עם מערכת מבוססת בינה מלאכותית. שנית, ההקשר שבו מתקיימת האינטראקציה. כשהשימוש במערכת נעשה במסגרת מערכת יחסים בה כבר הוכרה חשיבות השקיפות, כמו במסגרת מערכת יחסים צרכנית (מכוח דיני הגנת הצרכן)²²² או של טיפול רפואי,²²³ כך ייתכן שתהיה הצדקה משמעותית יותר לדרוש גילוי. שלישית, יש לבחון כיצד הגילוי צפוי להשפיע על התנהלות משתמש הקצה. במקרים בהם הגילוי יכול לסייע למשתמש להתאים את התנהגותו, למצות את זכויותיו או לאכוף את החובות של המפעיל, כך ייתכן שתהיה הצדקה משמעותית יותר לדרישה זו. צריך בהקשר זה לקחת בחשבון את העדר היכולת לצפות באופן מלא את אופני השימוש על ידי הפרט במידע האמור. רביעית, יש לבחון את ההשפעה של הגילוי על המפעיל. ככל שהדרישה לגילוי תגרור פגיעה

Andrew D. Selbst & Solon Barocas, *The Intuitive Appeal of Explainable Machines*, 87 FORDHAM L. REV. 1085, 1122-1126 (2018).

²²¹ שם, בעמ' 1133-1138.

²²² ראו למשל ס' 4 לחוק הגנת הצרכן, התשמ"א-1981.

²²³ ראו למשל ס' 13(ב) לחוק זכויות החולה, התשע"ו-1996.

גדולה יותר באינטרסים העסקיים של המפעיל, ביעילות השימוש במערכת או בסודות מסחריים, יהיה צורך להביא היבט זה בחשבון גם כן. כך גם כאשר דרישת הגילוי מגדילה באופן משמעותי את עלויות העסקה או הנטל הרגולטורי באופן שאינו מוצדק מבחינת איזון האינטרסים.

5.5. אמינות, עמידות ובטיחות

שימושים רבים במערכות מבוססות בינה מלאכותית אינם מעוררים חשש משמעותי גם כשהמערכת לא מגיעה לתוצאות הרצויות. כך למשל, כשנעשה שימוש במערכות אלה על מנת לאפיין משתמשים ולהציע להם פרסומות שתואמות את העדפותיהם האישיות, שגיאת המערכת ככלל לא תעורר עניין חברתי או כלכלי ממשי. במקרים מעין אלו, כשהחשש משגיאה של המערכת לא משמעותי, ההצדקה להתערבות רגולטורית בנוגע למידת האמינות והעמידות של המערכת אינה גדולה, אם בכלל. לעומת זאת, כשהשימוש במערכת עלול, במקרה של התממשות כשל פנימי או איום חיצוני, לגרום לפגיעה בזכויות יסוד או באינטרסים ציבוריים – ובמיוחד כשקיים חשש לפגיעה בבטיחות המשתמשים במערכת או המושפעים ממנה – ישנה הצדקה משמעותית יותר להתערבות רגולטורית במטרה להגביר את מידת האמינות, העמידות, האבטחה והבטיחות של המערכת. זאת, בפרט בנסיבות בהן מתקיימות הצדקות מקובלות להתערבות רגולטורית, כדוגמת החצנות שליליות או פערי מידע.

מעבר לכך, ניתן להצביע על מספר גורמים שעשויים לחזק את ההצדקה להתערבות רגולטורית, בין היתר עקב הגדלת הסיכוי להתרחשותם של סיכונים בטיחותיים:

ראשית, ככל שהשימוש במערכת נעשה בקנה-מידה (Scale) רחב יותר, הן כשמשמשים בה למגוון גדול יותר של משימות והן כשהיא משפיעה על מספר רב יותר של אנשים, כך שגיאה של המערכת עלולה לגרום לפגיעה משמעותית יותר. למשל, אם כל המכוניות העצמאיות של חברה מסוימת יעשו שימוש באותה מערכת, ויארע בה כשל, יגדל הסיכון לנזק רב שייגרם עקב תאונות.²²⁴ במובן זה, קנה מידה רחב כשלעצמו עשוי לחזק את ההצדקה להתערבות רגולטורית על מנת לוודא שמערכת אמינה, עמידה ובטוחה, וזאת מאחר ששגיאות בהיקף רחב עלולות להיות בעלות משמעות חברתית או כלכלית ולהיתרגם לפגיעה בזכויות הפרט או באינטרסים ציבוריים, ובפרט בבטיחות.

שנית, כאשר נעשה שימוש במערכות שאינן ניתנות להסברה (המכונות "קופסא שחורה"), גדל הסיכון להתרחשות אירועים בטיחותיים. במערכות אלו האפשרות לצפות מראש את כלל הסיכונים, לבחון את הסיכוי שיתממשו ולמנוע אותם מראש כבר בשלב פיתוח המערכת, מאתגרת יותר לעומת מערכות הניתנות להסברה. מעבר לכך, כאשר נעשה שימוש במערכות שאינן ניתנות להסברה, גם היכולת של גורם אנושי להתערב באופן אפקטיבי על מנת שיוכל לתקן שגיאות של המערכת מצטמצמת (להרחבה ראו פרק "הסברתיות" ופרק "מעורבות אנושית" לעיל).²²⁵

שלישית, כאשר נעשה שימוש במערכות שהצלחתן תלויה במהירות הפעולה שלהן, גובר החשש להתרחשות אירועים בטיחותיים, שכן ייתכן שמעורבות אנושית ב"זמן אמת" איננה אפשרית. כך למשל, בתחום המסחר האלגוריתמי, המערכת נדרשת לפעול במהירות כדי שתהיה אפקטיבית. במקרים אלו, לגורם אנושי שיבקש להתערב בהחלטת המערכת לא תהיה אפשרות מעשית לנתח את ההחלטה מספק מהר ולתקן שגיאות במקרה של כשל בפעילות המערכת או איום עליה.²²⁶

²²⁴ Arnold & Toner, לעיל ה"ש 189, בעמ' 17-18.

²²⁵ שם, בעמ' 17.

²²⁶ Thomas G. Dietterich, *Steps Toward Robust Artificial Intelligence*, 38 AI MAG. 3, 9 (2017).

יצוין כי בדומה לאמור בפרקים הקודמים, גם ביחס לסוגיית האמינות, העמידות, האבטחה והבטיחות של מערכות מבוססות בינה מלאכותית, היקף ההתערבות הרגולטורית (אם בכלל), ראוי להיבחן על פי מאפייני המערכת העומדת על הפרק ונסיבות השימוש בה. כאשר במסגרת הבחינה יש להביא בחשבון, בין היתר, את מידת החשיפה לכשלים פנימיים או איומים חיצוניים, ההסתברות להתממשותם והשלכות הצפויות אם יתממשו; את קיומו ועוצמתו של תמריץ עצמאי, לרוב כלכלי, של מפתחי המערכת להתמודד עם כשלים ואיומים אלה; ואת העלויות הכרוכות ביישום הדרישות הרגולטוריות על מנת למתן את הסיכונים הנובעים מכשלים ואיומים אלה. בנוסף, יש לשקול מהו הכלי המתאים ביותר להתערבות רגולטורית ביחס לאותה מערכת ובנסיבות הנתונות.

לדוגמה, אמצעי מרכזי המצוי בידי רגולטורים, לשם התמודדות עם החשש בעניין כשלים פנימיים של מערכות מבוססות בינה מלאכותית ואיומים חיצוניים עליהן, הוא דרישה לביצוע הערכות תפקוד וסיכונים למערכות אלה (ראו להלן הצעה שעניינה "התוויית שפה אחידה באמצעות מסגרת מומלצת (וולונטרית) לניהול סיכונים"). מטבע הדברים, נודעת חשיבות לעיתוי ולתדירות ביצוע הערכות אלה, שיכולות להיעשות לאורך כל "מחזור החיים" של המערכת, דרך שלבי הפיתוח וההטמעה שלה ולאחר הטמעתה. כמו כן, ישנה חשיבות לאופן ביצוע הערכות אלה. כך למשל, ביחס לתהליך הערכת התפקוד של המערכת, בשלב פיתוח המערכת כדאי לבצע את ההערכה בסביבה שמדמה באופן הקרוב ביותר את הסביבה שבה תידרש המערכת לפעול לאחר כניסתה לשימוש. באופן דומה, הערכת התפקוד שתבצע על מאגר נתונים חיצוני ושונה מזה שהמערכת התאמנה עליו, תאפשר לבחון כיצד היא מצליחה להתמודד עם מצבים לא מוכרים ובלתי צפויים.²²⁷

אמצעי נוסף העומד בפני הרגולטורים, הוא מעורבות אפשרית ביצירה של סטנדרטים טכנולוגיים ומקצועיים (best practices) למערכות מבוססות בינה מלאכותית, שיהיו מחייבים או מומלצים. סטנדרטים אלו יאפשרו מצד אחד להנגיש למפתחי המערכות את הכלים שבהם ניתן להשתמש כדי למתן את החשש משגיאה של המערכת, ומצד שני יאפשרו גם לקבוע סטנדרט מחייב שבכוחו לאון בין החששות הכרוכים בשימוש במערכת לבין התועלות שניתן להשיג באמצעותה.²²⁸

אמצעי אחר, הוא קביעת דרישה רגולטורית, במקרים המתאימים, להסברתיות של המערכת או למעורבות אנושית בפעילות המערכת (להרחבה ראו פרקים "הסברתיות" ו-"מעורבות אנושית" לעיל). בין היתר, פיקוח מצד גורם אנושי על פעילות המערכת עשוי לאפשר תיקון שגיאות שלה ב"זמן אמת", בטרם ייפגעו זכויות הפרט או אינטרסים ציבוריים; והסברתיות של המערכת יכולה לסייע לזיהוי ותיקון כשלים פנימיים בפעילותה בשלב מוקדם יחסית, וכן לשפר את אפקטיביות המעורבות האנושית. כך שמעבר ליתרונות שצוינו לעיל, קביעת הדרישות הללו עשויה להועיל גם לחיזוק האמינות, העמידות, האבטחה והבטיחות של מערכות מבוססות בינה מלאכותית.

לבסוף, יצוין כי ניתן להפחית את מספר האירועים הבטיחותיים שייגרמו כתוצאה משימוש במערכות אלה, באמצעות מיסוד מנגנון שיתוף מידע על אודות אירועים בטיחותיים שהתרחשו או שנמנעו. במסגרת זו, התערבות רגולטורית יכולה להיות בקביעת דרישת שיתוף המידע ואופן

²²⁷ Hamon, Junklewitz & Sanchez, לעיל ה"ש 167, בעמ' 15.
²²⁸ Gary Marchant, "Soft Law" Governance of Artificial Intelligence, AI PULSE (2019), <https://escholarship.org/content>

השיתוף (מתי, כיצד ועם מי ישותף), או במתן תמריצים לשיתוף פעולה מצד מפתחי המערכות והמשתמשים בהן (כגון על ידי הבטחת השמירה על מידע מסחרי רגיש של חברות המשתפות).²²⁹

5.6. אחריותות

על רקע האתגרים בהחלת ההסדרים הקיימים על השימוש במערכות מבוססות בינה מלאכותית, עיון בספרות האקדמית מעלה מספר הצעות ביחס למשטר האחריות המשפטי שראוי להחיל, הן לגבי דיני הנזיקין והן ביחס למשפט הפלילי, בכל הנוגע למקרים שבהם השימוש במערכות אלה הביא לתוצאות שליליות. להלן יוצגו מספר חלופות מרכזיות שהוצעו למשטר האחריות ולהסדרים שדרכם ניתן לבחון את השימוש במערכות, תוך עמידה על היתרונות והחסרונות העיקריים של כל אחת מהאפשרויות האלה. כך גם, תתואר האפשרות להגביר את מידת האחריות באמצעות הטמעת נורמות ארגוניות חדשות המותאמות לשימוש במערכות מבוססות בינה מלאכותית.

א. משטרי האחריות והסדרים חלופיים אפשריים

אפשרות אחת, היא להמשיך לעשות שימוש במשטרי האחריות הקיימים, ובפרט בעולות בדיני הנזיקין ובעבירות של רשלנות במשפט הפלילי, תוך ניסיון להתגבר על סוגיות ושאלות שמציב השימוש במערכות מבוססות בינה מלאכותית בפני יישום ההסדרים הקיימים. זאת, בין היתר, תוך הסתייעות בחזקות משפטיות ודוקטרינות כגון סעיף 41 לפקודת הנזיקין, הקובע כי בתנאים מסוימים,²³⁰ כאשר "הדבר מעיד על עצמו", הנטל יעבור לנתבע להוכיח שלא התרשל.

אפשרות נוספת, היא קביעת הסדרים חדשים וייעודיים שיסדירו את האפשרות להטיל אחריות על רקע השימוש במערכות מבוססות בינה מלאכותית. בפרט, אחת האפשרויות המרכזיות ליצירת הסדר חדש היא להחיל במסגרתו משטר של אחריות מוחלטת או קפידה, שאינה תלויה באשם.

אפשרות אחרת, היא כי לצד ההחלה של משטר האחריות המועדף ביחס לשימוש במערכות מבוססות בינה מלאכותית, ייקבעו הסדרים חדשים הכוללים "נמלי מבטחים" (Safe Harbors) שיגנו באופן מוחלט מפני הטלת אחריות, או לחלופין, שיובילו לכך ששאלת האחריות תיבחן על בסיס משטר אחריות אחר (למשל תחת משטר אחריות של רשלנות במקום של אחריות חמורה).

ב. הטמעת נורמות ארגוניות להגברת האחריותות

לצד עיצוב משטרי האחריות ועד לגיבוש הנורמות המשפטיות בצורה ודאית ויציבה, יש חשיבות לתהליכי הפנמה של הסיכון בידי הארגונים המפתחים והמשתמשים בבינה מלאכותית.

הטמעה של נורמות ארגוניות, בין אם באופן וולונטרי ובין אם כתוצאה מחובה רגולטורית (עוד בטרם קביעת היקף החובה המשפטית בדיעבד), תסייע לקידום ניהול הסיכונים וההגנה על זכויות יסוד ואינטרסים ציבוריים. זאת, במטרה להביא לכך שמפתחי המערכות והמשתמשים בהן יבינו כיצד מערכות אלה משליכות על זכויות יסוד ואינטרסים ציבוריים, ויוכלו לנקוט בצעדים הנדרשים על מנת למנוע או לצמצם את החשש שהשימוש בהן יגרום לתוצאות שליליות.

²²⁹ Arnold & Toner, לעיל ה"ש 189, בעמ' 19.

²³⁰ ס' 41 לפקודת הנזיקין. שלושת התנאים המנויים בסעיף הם:

(1) הנזיק לא ידע ולא היה יכול לדעת מה היו הנסיבות שהובילו לנזק;
(2) הנזק נגרם על ידי נכס שלנתבע היה שליטה מלאה עליו;
(3) נראה לבית המשפט שאירוע המקרה שגרם לנזק מתיישב יותר עם המסקנה שהנתבע לא נקט זהירות סבירה מאשר עם המסקנה שהוא נקט זהירות סבירה.

בכלל זה, לדוגמה, ניתן לעודד או לדרוש ממפתחי המערכות והמשתמשים בהן לבצע תהליכי הערכת השפעה או הערכת סיכון באופן עיתי, או למנות אחראיים ארגוניים שתפקידם להבטיח שהארגון נוקט באמצעים המתאימים כדי להתמודד עם האתגרים והסיכונים הנשקפים ממערכות אלה.

5.7. פרטיות

בהתאם לדרישות הדין הישראלי ולעקרונות רגולציה מהעולם,²³¹ ועל רקע החשש לסיכונים לפרטיות, בחלק זה יתוארו אמצעים מקובלים להקטנת החשש לפגיעה בפרטיות. כמובן שאין מדובר ברשימה ממצה, וייתכנו אמצעים ומנגנונים נוספים להפחתת החשש לפגיעה בפרטיות.

ראשית, ובאשר לצורך בגיבוש הסכמה מדעת מקום בו נעשה שימוש במידע אישי, יצוין כי עמדת הרשות להגנת הפרטיות, בהתאם להמלצות שפורסמו גופים בינלאומיים,²³² היא כי כשאיסוף המידע נעשה באמצעות מערכות אוטומטיות (כגון "בוטים"), ובכללן מערכות מבוססות בינה מלאכותית, יש לוודא כי יוצגו לנושא המידע כל הפרטים הנדרשים לפי סעיף 11 לחוק הגנת הפרטיות. עוד הביעה הרשות עמדתה, כי כאשר עיבוד המידע נערך באמצעות מערכות אוטומטיות, במסגרת הליך היידוע יש לפרט על אופן פעולת המערכות האמורות, ככל שהדבר רלוונטי לגיבוש ההסכמה וככל שפירוט זה אפשרי מבחינה משפטית, טכנולוגית, ומסחרית. כן המליצה הרשות, כי יוסבר לנושא המידע על אודות פרטי מידע בהם עשויות המערכות להשתמש במסגרת השימוש במידע הנוגע אליו, והמקור של פרטי מידע אלו.²³³ זאת ועוד, הרשות ציינה כי כשמערכת בינה מלאכותית מאפשרת להסיק על אדם מסקנות הנובעות מעיבוד מידע על אחרים בעלי מאפיינים דומים, יש ליידעו על כך.

המלצות נוספות שפורסמו על ידי גופים בינלאומיים, כוללים במסגרת חובת היידוע גם הודעה על כך שנושאי המידע מבצעים אינטראקציה עם בינה מלאכותית או מספקים מידע אישי אשר יעובד על ידי מערכת מסוג זה, ובמסגרת זו חלה החובה ליידע אותם מהן מטרות המערכת והשפעותיה האפשריות, מהו המידע המשמש את המערכת, ומה ההיגיון העומד בבסיס החלטותיה.²³⁴ נדמה כי כדי ליתן מענה לאתגר הפרטיות יהיה מקום לקחת בחשבון אף היבטים אלה, ובמסגרת זו יש חשיבות לכך שיידוע כאמור יעשה באופן פשוט, מובן ואפקטיבי.

שנית, ובהתאם לעיקרון צמצום המידע, יש לפעול על מנת שהמטרות שלשמן יעובד המידע, יוגדרו בבירור, יתועדו ויתאימו לציפייה סבירה של נושאי המידע בנוגע לשימוש במידע הנוגע להם. בהקשר זה, יש מקום להביא בחשבון גם את הצורך בצמצום המידע האישי, המשמש את תהליכי פיתוח בינה מלאכותית, כך שהוא יהיה מוגבל למידע הרלוונטי והנדרש עבור המטרות שהוגדרו. לצד זאת, יצוין כי פיתוח מערכת בינה מלאכותית נושא פעמים רבות אופי מחקרי. לא תמיד ידוע מראש אלו משתנים יסיעו בבניית מודל חיזוי איכותי, ולצורך הפיתוח מבוצע תהליך של ניסוי וטעיה. כמו כן, הליך איסוף מאגר המידע עשוי להיות מורכב, ארוך ויקר, ולפיכך יש רצון ליצור מאגרים אשר יוכלו לשמש למגוון גדול של פעילויות מחקריות בעתיד, כאשר אופי המחקר עצמו אינו ידוע במועד האיסוף. לפיכך, יישום אמצעי זה במקרים קונקרטיים איננו טריוויאלי.

²³¹ ראו למשל החלטת ה-Global Privacy Assembly, בו חברות רשויות להגנת המידע והפרטיות מרחבי העולם, בנושא בינה מלאכותית, משנת 2018 (להלן: החלטת ה-Global Privacy Assembly) https://globalprivacyassembly.org/wp-content/uploads/2018/10/20180922_ICDPPC-40th_AI-Declaration_ADOPTED.pdf.

²³² עמדת הרשות להגנת הפרטיות בעניין חובת יידוע במסגרת איסוף ושימוש במידע אישי, לעיל ה"ש 27, פס' 24.
²³⁴ ראו גם החלטת ה-Global Privacy Assembly, לעיל ה"ש 231, בנושא בינה מלאכותית משנת 2020.

שלישית, יהיה מקום להביא בחשבון, במקרים מסוימים, גם את האפשרות לשימוש במנגנונים מגבירי פרטיות, אשר יסייעו בצמצום היקפי המידע האישי, כגון התממה (אנונימזציה); שימוש במידע סינתטי (מידע אשר נוצר באופן מלאכותי וניתן להשתמש בו כתחליף למידע המקורי האמיתי); הוספת "רעש"; ²³⁵ Differential Privacy; ²³⁶ Federated learning (אימון המידע בכמה נקודות מקומיות, וחיבור התובנות למודל אחד, בלי לחלוק את המידע האישי); ²³⁷ ועוד. שיטות נוספות להגברת פרטיות הרלוונטיות לשלב הסקת המסקנות לגבי אדם ספציפי הן עיבוד המידע בצידוד הקצה של המשתמש ולא בשרת מרכזי, המרת המידע לקוד שאינו נקרא על ידי אדם ועוד. אף קביעת מדיניות לעניין מועדי מחיקת המידע המשמש לאימון המערכת, אשר תכלול קריטריונים לבחינת השאלה לשם מה נדרש המידע, האם שמירתו הכרחית, ואילו טכנולוגיות מגבירות פרטיות ניתן ליישם ביחס למידע אשר לא נמחק, תוכל לסייע בהתמודדות עם אתגר הפרטיות. יודגש כי תחום הטכנולוגיות מגבירות הפרטיות הינו תחום דינאמי אשר מתפתח במהירות, בין היתר בשל דרישות הרגולציה ברחבי העולם. משכך, פירוט הטכנולוגיות שנסקר אינו ממצה כלל וכלל.

בנוסף, יוזכר, כי בהתאם להוראת תקנה 2(ג) לתקנות הגנת הפרטיות (אבטחת מידע), התשע"ז-2017, כל בעל מאגר מחויב לבחון אחת לשנה האם הוא שומר במאגר מידע רב מן הנדרש למטרות המאגר. במסגרת בדיקה כזו, המתחייבת כאמור על פי הדין, מומלץ גם לבחון את הטכנולוגיות הקיימות בשוק באותה עת, בהן ניתן להסתייע כדי להתמים את המידע ובכך לצמצם דה-פקטו את המידע האישי שאינו נדרש עוד.

לבסוף, ולצורך יישום ההיבטים שתוארו לעיל, יש מקום, במקרים המתאימים, לקבוע נהלים ומנגנוני יישום ופיקוח פנימיים על השימוש במידע אישי. במסגרת זאת, יוכלו הרגולטורים או הרשות להגנת הפרטיות לקבוע הוראות מתאימות בדבר מינוי ממונה הגנה על פרטיות בארגון, ²³⁸ כמו גם על ביצוע הכשרות מתאימות וכן על קביעת מנגנוני בקרה ופיקוח על שימוש במידע אישי. מעבר לכך, ייתכן כי במקרים המתאימים, יהיה ניתן לעשות שימוש בתסקיר השפעה על פרטיות בארגון לשם בחינת מידת הנחיצות של השימוש באמצעים שתוארו לעיל. ²³⁹ לא כל שימוש במערכת מבוססת בינה מלאכותית מעורר את אותה מידה של חשש בנוגע לשמירה על הפרטיות של נושאי המידע; בהתאם, הרגולטורים הרלוונטיים והרשות להגנת הפרטיות יוכלו לבחון באופן מעמיק ומפורט, ובשים לב לאופי המערכת ומידת הסיכון והרגישות שיש בה, מהם האמצעים המתאימים בנסיבות העניין, לרבות אלה שתוארו לעיל, לצורך הבטחת ההגנה הנדרשת לזכות לפרטיות.

²³⁵ קרי, שינוי חלק מהמידע האישי באופן רנדומלי, תוך שימור ערכים סטטיסטיים במערך הנתונים, ²³⁶ שיתוף מידע פומבי על מערך נתונים על ידי תיאור דפוסי הקבוצות בתוך מערך הנתונים, מבלי לשתף מידע על פרטים במערך הנתונים עצמו.

²³⁷ כגון טכניקת למידת מכונה המאמנת אלגוריתם על פני מספר התקני קצה מבוזרים, או שרתים המחזיקים בדגימות נתונים מקומיות מבלי להחליף אותן.

²³⁸ הרשות להגנת הפרטיות ²³⁹ תסקיר השפעה על הפרטיות: מדריך עזר מתודולוגי (2021) [https://www.gov.il/BlobFolder/](https://www.gov.il/BlobFolder/https://www.gov.il/BlobFolder/)

6. חלק שישי: המלצות למדיניות רגולציה ואתיקה לתחום הבינה המלאכותית

פרק זה מפרט המלצות למדיניות אתיקה ורגולציה בתחום הבינה המלאכותית בישראל, תוך התחשבות בעקרונות רגולטוריים, משפטיים ואתיים מקובלים בעולם, וזאת, בהתאם להחלטת ממשלה 212. לצורך גיבוש המלצות אלה נסקרו המאפיינים הייחודיים של טכנולוגיית הבינה המלאכותית, וכן האתגרים והסוגיות המתעוררים בעקבות פיתוח ושימוש במערכות מבוססות בינה מלאכותית, והאופן בו מדינות מפותחות עוסקות בתחום זה. המדיניות המומלצת מבקשת לאזן בין הצורך בוודאות ובהירות לבין האינטרסים הציבוריים והזכויות שעל הפרק, והחשש מהתערבות שאינה מוצדקת העלולה לפגוע בחדשנות טכנולוגית.

6.1. אימוץ מדיניות רגולציה ממשלתית לתחום הבינה המלאכותית

מוצע לאמץ מדיניות רגולציה שתכלול את הרכיבים שיפורטו להלן, ותסייע לממשלה בבחינת הצורך בהתערבות רגולטורית ביחס לפיתוח ושימוש בבינה מלאכותית ואופי ההתערבות.

6.1.1. החשיבות של מדיניות רגולציה ומעמדה

נקודת המוצא להמלצה זו היא שישנה חשיבות רבה לכך שרגולציה המשפיעה על פיתוח בינה מלאכותית והשימוש בה במגזר הפרטי תקודם או תותאם לפי מדיניות רגולציה ממשלתית אחידה וקוהרנטית. זאת, על מנת להשיג את יעדי המדיניות של הממשלה, לקדם את תחום הבינה המלאכותית, להגן על זכויות היסוד של הפרט ועל אינטרסים ציבוריים, ולצמצם את החשש לפגיעה בחדשנות הטכנולוגית. בשים לב לכך שמדובר בתחום טכנולוגי וחדשני שמתפתח במהירות, שחשיבותו הכלכלית גבוהה, ושעשוי להשפיע על כלל מרחבי החיים, לקביעת מדיניות כאמור משנה חשיבות.

בהתאם לכך, מוצע לאמץ מדיניות רגולציה, ולכוון את גורמי הממשלה, ובפרט הרגולטורים הרלוונטיים, להתחשב בה בפעילותם בקביעת רגולציה ביחס לבינה מלאכותית. לדוגמה, מוצע כי רגולטור ממשלתי השוקל להחיל רגולציה על שימוש מסוים בבינה מלאכותית בתחום שעליו הוא אמון יביא בחשבון את החשיבות שבגיבוש אותה רגולציה בשים לב לנעשה במדינות המפותחות, יבצע הליך סדור של בחינת הסיכונים בשימוש שבו מדובר, יבצע תהליך שיתוף ציבור משמעותי, ויעדיף, במקרים המתאימים, להשתמש בכלי רגולציה מתקדמים ובכלי נסיינות רגולטורית (דוגמת "ארגו החול" הרגולטורי (Regulatory Sandbox)).

אשר לתוכן מדיניות הרגולציה, מוצע לאמץ את רכיבי מדיניות הרגולציה המפורטים להלן, שגובשו בין היתר בשים לב ל"עקרונות מנחים לאסדרה מיטבית" הקבועים בחוק עקרונות האסדרה, התשפ"ב-2021; למסמכי מדיניות רגולציה שפרסמו המדינות המובילות בעולם בפיתוח רגולציה בתחום הבינה המלאכותית; להמלצות שניתנו ולעבודות קודמות שנעשו בישראל במסגרת הפעילות הציבורית בתחום; ולדיוני הצוות הבין-משרדי שעסק ברגולציה ואתיקה במהלך גיבוש התכנית הלאומית לבינה מלאכותית בהתאם להחלטת ממשלה 212.

רכיבי מדיניות הרגולציה המוצעים להלן עומדים כל אחד בפני עצמו, אך מתחברים יחד לכדי מדיניות שלמה. מטבע הדברים, לא כל הרכיבים ניתנים ליישום ביחס לכל רגולציה בתחום הבינה המלאכותית, וההתחשבות בהם ראויה להיות תמיד על פי ההקשר והנסיבות ובשים לב לשיטות העבודה המקובלות, אך מוצע כי ככלל יהיה צורך להתחשב בהם ולשאוף ליישום.

6.1.2. רכיבי מדיניות הרגולציה המוצעים לאימוץ

א. רגולציה ענפית

אין בשלב זה מקום לקידום חקיקת מסגרת רוחבית על כל תחום הבינה המלאכותית, וזאת בין היתר עקב המגוון הרחב של השימושים בבינה מלאכותית בתחומים וענפים שונים בעלי מאפיינים נבדלים; הקצב הגבוה של ההתפתחויות הטכנולוגיות; והיעדר הבנה מספיקה של השלכות הטכנולוגיות ברמה המעשית והמשפטית. לכן, לעת עתה, מומלץ לפעול כך שהרגולטורים הענפיים יבחנו את הצורך בקידום הרגולציה קונקרטית בתחומם, תוך פעולה לאור מדיניות ממשלתית מתואמת ואחידה ובסיוע מוקד הידע הממשלתי כפי שיפורט בהמשך.

עם זאת, אין ספק שהשאלה אם נכון לאמץ הסדרים רוחביים, מתי, ומה תכנם, היא עניין שיש להמשיך ולתת עליו את הדעת, גם בשים להתפתחויות בעניין זה במישור הבין לאומי, ולא מן הנמנע שיהיה צורך בכך בעתיד. לכן מומלץ להמשיך ולעקוב אחר ההתפתחויות הטכנולוגיות והרגולטוריות בארץ ובעולם, ולשוב ולשקול אפשרות זו בהמשך.

במסגרת רכיב זה במדיניות הרגולציה מוצע כי בשלב זה ובהתאם למצב הדברים הקיים כיום, ככל שתקודם רגולציה שמטרתה להסדיר פיתוח בינה מלאכותית או שימוש בה, הדבר ייעשה ברמת הענף (סקטור) או התחום שבו היא נדרשת על בסיס צורך קונקרטי ובהובלת הרגולטור האמון עליו, תוך פעולה לאור מדיניות ממשלתית אחידה באמצעות תיאום. בהקשר זה מוצע גם למסד מנגנוני תיאום בין רגולטורים ועם מוקד הידע (עליו יפורט להלן), ובפרט כאשר להוראותיהם של רגולטורים עשויה להיות השפעה רוחבית. אך, כאמור, לעת הזו מוצע לא לקדם חקיקה רוחבית ייעודית להסדרת תחום הבינה המלאכותית.

תחום הבינה המלאכותית הוא רחב ומתייחס לטכנולוגיות שונות, להן שימושים מגוונים, המופיעים כמעט בכל ענף במשק, ובכלל זה בריאות, תחבורה, פיננסים, חקלאות וחינוך. נוכח הבסיס המשותף של הבינה המלאכותית מצד אחד, וגיוון השימושים מצד שני, שאלה מרכזית שמתעוררת היא אם יש להסדיר פיתוח ושימוש בבינה מלאכותית (ככל שיש צורך בכך) באמצעות רגולציה כלל-משקית אחת, או באמצעות רגולציות ספציפיות לכל ענף בנפרד.

אפשרות אחת היא רגולציה רוחבית וכללית שתשמש כמסגרת רגולטורית המסדירה פיתוח ושימוש בבינה מלאכותית בכלל הענפים, באופן הדומה לתחולת הרגולציה בתחום התחרות, הפרטיות והגנת הצרכן. אפשרות אחרת היא רגולציה ענפית, שתחול רק ביחס למגזר או לשוק מסוים שאותו היא נועדה להסדיר באופן ספציפי, כגון בריאות או בנקאות, או ביחס לחתך נושאי, כגון במסגרת דיני השוויון.²⁴⁰ להמחשה, ההבדל יכול להיות בין קידום "חוק בינה מלאכותית" שיכלול הוראות כלליות שיחולו על כל המשק, לבין קביעת כללים פרטניים על ידי רגולטור מסוים לגבי שימוש בבינה מלאכותית שבתחומו.

כמפורט בפרק "פעילות ארגונים בינלאומיים ומדינות מפותחות" לעיל, **האיחוד האירופי** מקדם טיוטה של חוק הבינה המלאכותית (Artificial intelligence Act), המציעה מסגרת רגולטורית רוחבית ראשונה מסוגה לבינה מלאכותית. הטיוטה קובעת מסגרת משפטית שתחול

²⁴⁰ יצוין שיש גם מי שהציעו רגולציה שתהיה רוחבית יותר, ותסדיר באופן כללי את התחום הטכנולוגי. כלומר, רגולציה שתסדיר חדשנות טכנולוגית, אולם על פניו נראה שמדובר ביריעה רחבה מידי.

על פיתוח ושימוש בכל המערכות המבוססות על בינה מלאכותית באיחוד האירופי (להרחבה ראו סקירת טיוטת החקיקה של האיחוד האירופי לעיל). לעומת זאת, מדינות מובילות אחרות, ובהן **ארה"ב**,²⁴¹ **יפן**²⁴² ו**סינגפור**,²⁴² נוקטות בגישה שונה, שבבסיסה איתור סיכונים והתמודדות איתם ברמת הארגון, באמצעות שפה משותפת. שפה זו מאפשרת בירור וליבון מעמיקים לגבי הצורך והאופי של מעורבות רגולטורית, אך לא מקודמת לצידה רגולציה רוחבית לתחום.

מדינה מובילה נוספת שאימצה מדיניות של רגולציה ענפית לתחום הבינה המלאכותית היא **בריטניה**. כאמור בסקירה לעיל, מדיניות הרגולציה שפרסמה ממשלת בריטניה מעדיפה בצורה ברורה שימוש ברגולציה ענפית, וזאת לאור היתרונות שראתה בכך. התפיסה הבריטית מבוססת על יצירת עקרונות משותפים לבינה מלאכותית, והבאתם לכדי יישום בהתאם לשימוש הרלוונטי ותוך הסתמכות על הידע והמומחיות של גופי הרגולציה הסקטוריאליים.

אמנם, לקידום רגולציה רוחבית יתרונות בתחום האחידות והקוהרנטיות בהתמודדות עם טכנולוגיות בינה מלאכותית. במקרים רבים האתגרים הבסיסיים הכרוכים בבינה מלאכותית זהים או דומים, ללא קשר לענף שבו היא מוטמעת. הסדרת אתגרים אלה באמצעות רגולציות שונות ומקבילות עלולה ליצור אי-אחידות, פערים, סתירות וחוסר ודאות. זאת ועוד, סטנדרטים אחידים מאפשרים שפה משותפת בין גורמי הממשלה, גורמים פרטיים, ואף התממשקות טובה יותר בין מערכות הבינה המלאכותית עצמן, בענפים שונים.

ואולם, השימושים בבינה מלאכותית מגוונים ונעשים ככלל בהקשר מסוים, כך שלא נקל ניתן לקבוע רגולציה כוללת שתתאים באופן מספק למאפיינים של כל שימוש ותטפל בכולם. **בנוסף**, נוכח ריבוי השימושים האפשריים בבינה מלאכותית, ממילא עשוי להתעורר צורך להתייחס לכל שימוש ותחום פעילות על בסיס מאפייניו ומכלול נסיבותיו במסגרת הרגולציה הרוחבית. **כמו כן**, גיבוש רגולציה רוחבית יהיה כרוך בהליך ממושך (בפרט כשיש צורך לעגן אותה בחקיקה), מה שעלול להביא לכך שבמועד סיומו התוצאה לא תענה לצרכים המקוריים, או תגרום לקיבוע תוצאות המשקפות את המציאות הטכנולוגית והחברתית הקיימת בשלב מסוים, ולהתיישן עם שינוי המציאות הטכנולוגית והחברתית בכל ענף. זאת, בשים לב לכך שההתפתחויות הטכנולוגיות הן מהירות, כאשר משמעותן והשלכותיהן המעשיות והמשפטיות טרם הובררו כל צרכן. **לבסוף**, בעת הזו נראה כי מוטב להרחיב את התשתית העובדתית והמקצועית בעניין מלאכת הרגולציה של פיתוח ושימוש בבינה מלאכותית בישראל.

גם דו"ח בינה מלאכותית בשירותים פיננסיים עסק בהבחנה בין רגולציה ענפית ורוחבית. החוקרים הביעו עמדה שלפיה "בשלב זה החסרונות של קביעת עקרונות אחידים שיחולו בתחומי משפט שונים עולים על היתרונות", ובהתאם המליצו כי "בשלב זה אין לקבוע עקרונות אחידים" ביחס לכל השימושים בבינה מלאכותית. לצד זאת, הם הצביעו על "צורך בתיאום" כדי להתמודד באופן רגולטורי מתואם עם סוגיות בעלות מאפיינים רוחביים. בנוסף, הם המליצו לבחון את הרגולציה הפיננסית – שעמדה במוקד מחקרם – כמקשה אחת, עקב הקשר

Ministry of Economy, Trade and Industry, Consumer Affairs Agency, Expert Group on how AI²⁴¹ Principles Should be Implemented, AI Governance in Japan Ver. 1.1., [https://www.meti.go.jp/](https://www.meti.go.jp/Government of Japan, Ministry of Economy, Trade and Industry, Study Group on New Governance Models in Society, Governance Innovation ver.2 - A Guide to Designing and Implementing Agile Governance, 2021, https://www.meti.go.jp/)

Personal Data Protection Commission Singapore, Singapore's Approach to AI Governance,²⁴² <https://www.pdpc.gov.sg/Help-and-Resources/2020/01/Model-AI-Governance-Framework>

בין התחומים השונים הנכללים בה (בנקאות, ניירות ערך ושוק ההון), וכן הציעו להפקיד בידי גורם ממשלתי אחד את "תיאום סוגיית האסדרה של יישומי בינה מלאכותית בתחומים שונים" (ראו גם המלצה בעניין "מיסוד מוקד ידע ותיאום ממשלתי להסדרת בינה מלאכותית").

יחד עם האמור, וכפי שיפורט בהרחבה להלן, נוכח חשיבות הנושא, ההערות השונות שהתקבלו לטיטוט מסמך זה וההתפתחות המהירה של התחום – מוצע להקים מוקד ידע ממשלתי לתחום הבינה המלאכותית. מוקד הידע הממשלתי יפתח את המומחיות המקצועית והטכנולוגית הנדרשת, ובין היתר יסייע לגורמי הממשלה והרגולטורים הרלוונטיים במיפוי השימושים בבינה מלאכותית והאתגרים הנלווים להם בענפים שעליהם הם אמונים ובפיתוח הרגולציה בתחום זה בהתאם למדיניות הממשלה, באופן קוהרנטי ומתוך מבט רוחבי. מוקד הידע גם יוכל להמשיך ולעקוב אחר ההתפתחויות במישור הרגולציה הבין לאומית, כמו גם במישור הטכנולוגי ולגבש המלצות לשינוי מדיניות הרגולציה ובכלל זאת לעניין הצורך ברגולציה רוחבית בעתיד. זאת, נוכח ההבנה כי ההתפתחויות בתחום זה עשויות להצריך שינויי מדיניות גם בהקשר זה.

ב. חשיבות לרגולציה התואמת לנעשה במדינות מפותחות ובארגונים בינלאומיים

מוצע כי הרגולציה ביחס לפיתוח ושימוש בבינה מלאכותית, ככל שתאומץ, תביא בחשבון את הרגולציה הנוהגת במדינות מובילות בתחום ושאימצה על ידי ארגונים בינלאומיים, ותתאים עצמה במידת האפשר לרגולציה רלוונטית שנקבעה במדינות ובארגונים כאמור.

פרקטיקות ופיתוחי מדיניות של מדינות וארגונים שונים בעולם עשויים להוות מקורות השראה רלוונטיים לנעשה במדינת ישראל, בין היתר ביחס להגנה על זכויות יסוד ואינטרסים ציבוריים. כמו כן, התאמת הרגולציה הישראלית לסטנדרטים גלובליים מובילים יכולה למנוע חסמים רגולטוריים ייחודיים בפני כניסת טכנולוגיה לישראל, ולתמוך ביכולתה של התעשייה הישראלית לייצא למדינות היעד העיקריות.

מדינות רבות בעולם וארגונים בינלאומיים עוסקים בקידום עקרונות משפטיים ורגולציה לתחום הבינה המלאכותית. עם המדינות המובילות נמנות לדוגמה ארה"ב, בריטניה, קנדה, יפן וסינגפור. בין הארגונים נמנים לדוגמה ה-CAI, UNESCO, OECD (ועדה ייעודית של מועצת אירופה העוסקת בהסדרת בינה מלאכותית), GPAI והאיחוד האירופי. בנוסף, נעשית עבודה נמרצת בארגוני תקינה בינלאומיים כמו ISO/IEC, IEEE ו-NIST לפיתוח תקנים ל"בינה מלאכותית אחראית". אלה, ורבים אחרים, עוסקים בהסדרת בינה מלאכותית, ויש להניח כי לתוצרי העבודה שלהם תהיה השפעה גלובלית, כולל על הנעשה בישראל. כבר היום, תוצרים של ארגונים ומדינות אלה מסייעים רבות לגיבוש המדיניות שמופיעה במסמך זה.

בשים לב לגודלם היחסי של שווקי ישראל, ולכך שטכנולוגיות שמפותחות בישראל מכוונות פעמים רבות בעיקר לשווקים הגלובליים, יש חשש כי יצירה של רגולציה ייחודית לישראל עבור תחום הבינה המלאכותית, עלולה, בין היתר, להיתרגם לחסם רגולטורי ייחודי בפני כניסת מערכות מבוססות בינה מלאכותית לישראל, ולהפסד התועלות הכלכליות והחברתיות מכך.

לפיכך, למעט במקרים חריגים, מוצע להימנע מרגולציה ייחודית בקשר לבינה מלאכותית, ולהתאים ככל האפשר את הרגולציה בישראל לזו שאימצו ויאמצו המדינות המובילות בתחום והארגונים הבינלאומיים. ניתן לצפות כי עשויות להתפתח גישות שונות, כפי שניכר כבר היום

ביחס לשאלת הצורך בחקיקת מסגרת רוחבית לתחום הבינה המלאכותית. במקרים כאלו, יהיה צורך לבחור בין גישות מקובלות שונות מהעולם, בהתחשב במכלול היתרונות והחסרונות של כל אחת. בדרך זו ניתן לוודא שהרגולציה בישראל בתחום הבינה המלאכותית תהיה תואמת, ככל הניתן, לתוכן הרגולציה המקובלת בשווקים הגלובליים (כגון התקנים המתפתחים), אף אם ההסדרים המשפטיים ואופן מימושם יהיה שונה.

הדבר עולה בקנה אחד גם עם המגמה ההולכת ומתגברת בשנים האחרונות להתאמת הרגולציה בישראל לרגולציה המקובלת בעולם. ביטוי מרכזי למגמה זו מופיע בחוק עקרונות האסדרה הקובע כי רצוי שרגולציה תיקבע ככלל על בסיס כללים ואמות מידה מקצועיים שגובשו בארגונים בינלאומיים, או שחלים במדינות מפותחות עם שווקים משמעותיים. ביטוי נוסף למגמה זו מצוי ב"רפורמת היבוא" שעברה במסגרת חוק ההסדרים בשנת 2021 ונועדה להתאים את הרגולציה על היבוא בישראל לרגולציה זרה ובכך להגביר תחרות ולהוזיל מחירים.

גם ועדת המשנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה בראשות פרופ' קרין נהון (ראו פרק "העיסוק הממשלתי בבינה מלאכותית" לעיל), המליצה על התאמת הרגולציה בתחום הבינה המלאכותית לנעשה במדינות מפותחות, ובשון דו"ח הוועדה: "הוועדה רואה חשיבות להתאים את האסדרה הנבחרת לחקיקה, למדיניות ולתקינה מקובלות במדינות המפותחות, על מנת שמדינת ישראל תוכל להישאר בחוד החנית של התחום".

ג. החלת רגולציה המבוססת על כלים ומסגרות לניהול סיכונים

רכיב זה במדיניות הרגולציה מציע כי רגולציה המסדירה פיתוח ושימוש בבינה מלאכותית, תותאם ככלל לסיכונים הנשקפים מסוג הטכנולוגיה ומהשימוש הספציפי שאותו היא נועדה להסדיר, ותהיה תוצר של ניהול סיכונים שביצע הרגולטור ביחס אליו. כך שכלל לא תחול באופן אחיד על טכנולוגיות ושימושים שרמת הסיכונים והחששות לגביהם שונה באופן ניכר.

לבינה המלאכותית מגוון רחב של שימושים, שאינם שווים זה לזה במידת הסיכון והרגישות הקשורה בהם (למשל, על פניו נראה כי יש הבדל בין שימוש בבינה מלאכותית כדי להציע לצרכן סרט או שיר לטעמו, לבין שימוש בה לאבחון רפואי או החלטת אשור). על כן, ייתכן בהחלט שגם בתוך ענף רגולטורי מסוים, יהיה שוני במידת ההצדקה להתערבות רגולטורית בין שימוש אחד לאחר, הנובע מכך ששימוש אחד מעורר חששות וסיכונים גדולים יותר – אם משום שהפגיעה האפשרית בזכויות או אינטרסים חמורה במיוחד, ואם משום שהטכנולוגיה עדיין אינה בשלה לשימוש המוצע. כך, במסגרת מדיניות של ניהול סיכונים ייתכן שיהיו שימושים בבינה מלאכותית שיימצא שאינם מחייבים אסדרה, וייתכן ששימושים אחרים יימצאו כה מסוכנים עד שאין לאפשר אותם כלל. בהתאם לכך, מוצע להבחין בין רמת הסיכון המשתנה של השימושים השונים במסגרת הרגולציה, ולגבשה באופן שיביא לידי ביטוי ביצוע הליך כלשהו לניהול סיכונים על ידי הרגולטור (ראו גם המלצה בעניין כלי לניהול סיכונים להלן).

יצוין כי באופן כללי תפיסות ניהול סיכונים בתחום הבינה המלאכותית באות לידי ביטוי בשני היבטים. היבט ראשון הוא ניהול הסיכונים על ידי גורמי הממשלה והרגולטורים הרלוונטיים כחלק ממדיניות הרגולציה. כך, לפי תפיסת ניהול סיכונים גוברת ההצדקה להתערבות רגולטורית בפיתוח ושימוש בבינה מלאכותית, מקום שבו הסיכון לפגיעה בזכויות יסוד או אינטרסים ציבוריים גבוה. בהתאמה, ככל שהסיכון נמוך יותר, ההצדקה להתערבות

מצמצמת, וכללי ההתנהגות עשויים להיות בגדר המלצה בלבד. היבט שני הוא ניהול הסיכונים ברמת הארגון המפתח או משתמש בבינה מלאכותית. תפיסת ניהול סיכונים זו מבוססת על תפיסות מקובלות לניהול סיכונים, בין היתר בהתאם לעקרונות האתיים (ראו הצעה להלן), כדי לאתר את הסיכונים הנובעים מבינה מלאכותית לזכויות יסוד ואינטרסים ציבוריים, ולנקוט באמצעים טכנולוגיים ונהליים בכדי לצמצם אותם לרמה מקובלת.

הממשק בין שני היבטים אלה, הוא שתפיסת ניהול הסיכונים של מדיניות הרגולציה מסייעת להפעלת סמכות הרגולטור בין היתר כדי לבחון אם ההתמודדות הארגונית הוולונטרית מספקת מענה הולם לסיכון לפגיעה בזכויות יסוד או באינטרסים ציבוריים.

טיוטת הרגולציה של האיחוד האירופי (גרסת הנציבות (EU Commission) משנת 2022) מציעה לקבוע סיווג למערכות בינה מלאכותית עם דרישות וחובות שונות המותאמות על פי "גישה מבוססת סיכונים". לפי גישה זו, יישומי בינה מלאכותית יוסדרו רק בהתאם לנדרש כדי לטפל בסיכון קונקרטי. הטיוטה מבחינה בין מערכות ב"סיכון בלתי מקובל", "סיכון גבוה", "סיכון מוגבל", ו"סיכון מינימלי", ומגדירה את התהליך לסיווג המערכות לתוך קטגוריות אלה. לשם ההמחשה, ב"סיכון בלתי מקובל" נכללות מערכות המשמשות רשויות ציבוריות ל"דירוג חברתי" (social scoring)²⁴³ ומערכות שמשתמשות בטכניקות תת-הכרתיות מניפולטיביות; וב"סיכון גבוה" נכללות מערכות לזיהוי ביומטרי, חינוך, תעסוקה ומערכות המשמשות לגישה לשירותים פרטיים חיוניים. בהתאם לסיווג המערכות, לפי רמות הסיכון, מוטלות חובות שונות על מפתחי המערכות והמשתמשים בהן. מאיסור על שימוש במערכת, דרך עמידה בדרישות רגולטוריות לשם שימוש במערכת, ועד שימוש (כמעט) חופשי במערכת. כפי שצוין בפרק הסקירה הבינלאומית לעיל, גרסת הפרלמנט האירופי מציעה מספר שינויים ומבקשת לכלול קטגוריה חדשה ביחס ל-foundation models, אך משמרת את המבנה העיקרי שהציעה הנציבות.

גם ועדת המשנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה הציעה חלוקה בין אזורי סיכון בינוני וגבוה (בהם נדרשת לשיטתה חקיקה); אזורי סיכון בינוני ונמוך (בהם תידרש תקינה); ואזורי סיכון נמוך (שלא דורשים מגבלה משפטית). בנוסף, הוועדה הציעה כלי שפיתחה, הכולל בין היתר קבוצה של שאלות ראשוניות לבחינת ההשפעה של מערכת מבוססת בינה מלאכותית, ובכלל זה מהי עוצמת הפגיעה הפוטנציאלית בפרט או בציבור, ומהי עוצמת הפגיעה הצפויה אם יעשו שימוש לרעה במוצר או שהוא יצא משליטה.

לבסוף יצוין, כי גם חוק עקרונות האסדרה מורה על התאמת הרגולציה לסיכונים. בהקשר זה, החוק קובע כי רגולציה מיטבית היא כזו הנקבעת במטרה להביא למירב התועלת לחברה ולמשק, תוך שקילת האינטרס המוגן וההשפעות הנובעות מקביעת הרגולציה או אי-קביעתה, עלות הציות לה, ו"ככלל, על בסיס ניהול סיכונים". בדברי ההסבר לחוק, הצורך בניהול הסיכונים הובהר כך: "מוצע כי אסדרה מיטבית תיקבע, ככלל, לפי עקרונות של ניהול סיכונים.

²⁴³ "דירוג חברתי" ("social scoring") עוסק בקביעת ציון לאדם על בסיס התנהגותו בשורה של פרמטרים ותחומים, תוך פיתוח הרעיון של "דירוג אשראי" למשל, הקובע את רמת סיכון או מהימנות האשראי של אדם, לדירוג רחב יותר; הרשות להגנת הפרטיות **דירוג חברתי בראי הזכות לפרטיות: סקירת רקע בעניין שימוש במערכות לדירוג חברתי** https://www.gov.il/he/departments/publications/reports/social_ranking (2021).

זאת, מתוך הבנה שככלל, אין זה רצוי מבחינה כלכלית וחברתית לפעול באמצעות אסדרה כדי לאיין סיכונים באופן מוחלט".

ד. פיתוח הרגולציה באופן מודולרי, תוך הסתייעות בנסיינות רגולטורית, ושמירה על ניטרליות

טכנולוגיות

מוצע כי רגולציה שנועדה להסדיר פיתוח ושימוש בבינה מלאכותית תקודם ותתפתח בשלבים ובאופן גמיש בהתאם להתפתחויות הטכנולוגיות, וכן מוצע כי ייעשה שימוש בנסיינות רגולטורית, ובכלל זה בפיילוטס וב"ארגזי חול" רגולטוריים שיאפשרו הכנסה בטוחה של מערכות בינה מלאכותית. הכל, תוך שימוש בכלים ניטרליים מבחינה טכנולוגית (כלומר, לא כללים מפורטים ונוקשים המעוצבים לפי מערכת מסוימת כך שמתאימים רק לה).

דומה כי האתגר המשמעותי ביותר הניצב בפני הרגולטורים בבואם להסדיר טכנולוגיה בכלל, ובינה מלאכותית בפרט, הוא הקצב המהיר והדינמיות המאפיינים את ההתפתחות הטכנולוגית (מה שידוע כ"pacing problem"). המתח המובנה שבין נוקשות המערכת הרגולטורית לבין דינמיות ההתפתחות הטכנולוגית מייצר אתגר דו-כיווני. מצד אחד, הרגולציה לעיתים קרובות מתקשה להדביק את הקצב המהיר של ההתפתחויות הטכנולוגיות, ולהציע הגנה מספקת מפני הסיכונים הכרוכים בהן. שינוי ועדכון של הרגולציה כרוך לעיתים קרובות בהליכים שהשלמתם נמשכת זמן לא מבוטל, בין אם המדובר בנסיבות בהן יש צורך לתקן חקיקה ראשית, ובין אם נדרש תיקון של חקיקת משנה או עדכון נהלים של הרגולטור. במצבים אלה הרגולציה "נשארת מאחור" ומותירה ואקום רגולטורי שמאפשר התפתחות טכנולוגית ללא פיקוח מספק והגנה הולמת על זכויות ואינטרסים. מצד שני, הקושי בעדכון רגולציה קיימת מביא לכך שהרגולציה מהווה לעיתים חסם בלתי מוצדק הבולם התפתחות טכנולוגית רצויה. במצבים אלו, הרגולציה – שפעמים רבות נקבעה בתקופה בה כלל לא היה ניתן להעלות על הדעת את ההתפתחות הטכנולוגית שאירעה – מונעת את הפיתוח או השימוש בטכנולוגיה ללא הצדקה.

אתגר נוסף הוא שלעיתים עולה צורך ברגולציה בשלב שבו עדיין אין ודאות בנוגע לשימושים וההשלכות של טכנולוגיית בינה מלאכותית מסוימת, מה שמשפיע על היכולת להסדיר אותה. יש לציין, כי ככלל קיים מתח מסוים (trade-off) בין הבנת השימושים וההשלכות של טכנולוגיה חדשה לבין מידת היכולת להשפיע על ההתפתחות שלה. שכן לרוב ניתן להסדיר יותר בקלות טכנולוגיה נתונה כשהיא עדיין "צעירה" ולא פופולרית, כאשר במצב זה לעיתים קרובות ההשלכות הבלתי צפויות והלא רצויות שלה עדיין אינן ניכרות בבירור; או שניתן להמתין ולזהות השלכות אלה, אבל אז ייתכן שיהיה קשה להסדיר אותה לאחר התקבעות השוק.

על רקע זה, ומסיבות נוספות לרבות פערי מידע ומומחיות, נודעת חשיבות לקיומם של מחזורי רגולציה לומדים וגמישים, הכוללים בחינה של הצורך בשינוי מדיניות לאחר שהתקבלה ויכולת התאמה מהירה שלה.²⁴⁴ כלומר, יש יתרון לקביעת כלי המדיניות ועקרונות התוכן בהסדרים המאפשרים שינוי והתאמה, למשל, עדיפות לעיגון הרגולציה בכלים גמישים יחסית כגון תקנות ונהלים על פני חקיקה ראשית, מקום בו ניתן לקבוע רגולציה במתכונת זו. בהמשך לכך, יש חשיבות לבחינה ועדכון של הסדרים אלו לעיתים תכופות בהתאם להתפתחות הטכנולוגית.

²⁴⁴ שם.

זאת ועוד, על רקע זה מתחדדת החשיבות של הטמעת כלים ויכולות שיאפשרו למערכת הרגולטורית להיות יותר לומדת, דינמית וגמישה. יש צורך בהטמעת יכולות של נסיינות רגולטורית ונכונות לקדם נסיינות ביחס לבינה מלאכותית. יכולות אלה קריטיות בעידן הנוכחי המתאפיין בהתפתחות טכנולוגית מואצת, ורלוונטיות במיוחד לגבי בינה מלאכותית. שכן נדרשת גישה זריזה ויעילה לרגולציה שתאפשר למצות את פוטנציאל התפתחויות בתחום זה, לאפשר ואף לעודד אותן, אך לעצב אותן באופן שלא פוגע בזכויות יסוד ובאינטרסים ציבוריים.

אחד הכלים המרכזיים המשמשים לנסיינות רגולטורית הוא "ארגז החול" הרגולטורי (Regulatory Sandbox) (להלן: הסנדבוקס), כלי מדיניות המאפשר לרגולטורים לייצר "סביבת ניסוי" רגולטורית. במסגרת זו, הרגולטור מבצע התאמות ברגולציה הקיימת (ככלל נותן הקלות ביחס לחלק מדרישותיה) כדי לאפשר למפוקחים לפתח או להכניס לשוק טכנולוגיה חדשה באופן מבוקר, מצומצם ולזמן מוגבל.²⁴⁵ בנוסף, הסנדבוקס מסייע לרגולטורים להתמודד עם פערי מידע ומומחיות בינם לבין השוק ועם קצב ההתפתחות המהיר של הטכנולוגיה, שכן רגולציה זמנית, מדודה וגמישה מאפשרת תהליך למידה "תוך כדי תנועה" ומקנה כלים להסרה מבוקרת וזהירה של חסמים ממוקדים.²⁴⁶ המפוקחים מצדם מרוויחים מההזדמנות לפעול בסביבה "ידידותית", ובלי לשאת בנטל רגולטורי שלא חיוני לפעילותם.²⁴⁷

הגם שהשימוש הרווח בסנדבוקס הוא כאשר נדרשת הקלה או הגמשה ביחס לרגולציה הקיימת בשוק מוסדר, ניתן לעשות בו שימוש לעיתים גם במצב ההפוך שבו על פניו נראה כי נדרשת תוספת או הקשחה ביחס לרגולציה הקיימת בשוק שאינו מוסדר. מצב זה, שניתן לכנות "סנדבוקס הופכי" (Reverse Sandbox), הוא שימוש בסנדבוקס לשם הוספת תנאי המגביל פיתוח או שימוש בטכנולוגיה, כתחליף לחסימתה או להגבלתה באופן נוקשה וקבוע יותר. למשל, לעיתים מבוקש להכניס לשוק טכנולוגיה חדשה שיוצרת סיכונים שטרם הוסדרו במישור, אך לא מן הנמנע שיוסדרו בעתיד. במצב זה, הרגולטור יכול (בכפוף לסמכות ולהתנהלות המתאימה על פי כללי המשפט המנהלי) להציע למשתתפים להשתמש בסנדבוקס לצורך למידת הטכנולוגיה והשלכותיה, ואולי אף לקבוע תנאים ייעודיים שימנעו בטווח הקצר פגיעה באינטרסים פיקוחיים שבאחריותו. בתמורה לכך, הוא יוכל להימנע בשלב זה מהסדרה נוקשה ומהירה שתגביל את הפיתוח והשימוש בטכנולוגיה. הסנדבוקס ההפוך עשוי להתגלות כשימושי גם ביחס לבינה מלאכותית, שכן במקרים רבים הסיכונים לזכויות יסוד ואינטרסים ציבוריים הכרוכים בה אינם מוסדרים באמצעות רגולציה קודמת. הוא אף עשוי להיות כלי מעודד הטמעת טכנולוגיה, בעקבות הגברת אמון הציבור בה, הנובעת ממעורבות הרגולטור.

גם ועדת המשנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה של בינה מלאכותית הכירה באתגר הקשור לדינמיות ההתפתחות הטכנולוגית בתחום הבינה המלאכותית. בין היתר על רקע זה, המליצה הוועדה על בחינה קבועה ובפרקי זמן קצרים (בהשוואה לרגולציה בתחומים מסורתיים) של המדיניות הרגולטורית, וכן המליצה לשלב תהליכים של ניסוי מבוקר למדיניות, בפרט באמצעות סנדבוקס. יצוין, כי גם בטיטוט הרגולציה

²⁴⁵ משרד המשפטים, דו"ח הצוות הבין-משרדי לרגולציה חכמה: תכנית לאומית למדיניות רגולציה, בעמ' 129 (2021) (להלן: דו"ח הצוות הבין-משרדי לרגולציה חכמה). <https://www.gov.il/BlobFolder>

²⁴⁶ שם, בעמ' 130.

²⁴⁷ אחיעז, חמדני, עמירם וקסטיאל, לעיל ה"ש 25, בעמ' 21; Lawrence G. Baxter, *Adaptive Financial Regulation*; and RegTech: A Concept Article on Realistic Protection for Victims of Bank Failures, 66 DUKE L.J. 567, 600-01 (2016).

האירופית, תחת פרק שעוסק ב"אמצעים לתמיכה בחדשנות", מוצע להשתמש בסנדבוקס כדי לבחון עמידה של מערכות מבוססות בינה מלאכותית בדרישות שקובעת הטיטה.²⁴⁸

בנוסף, ישנה חשיבות לשימוש בכלים רגולטוריים וברגולציה שיהיו ניטרליים מבחינה טכנולוגית. כלומר, לא כללים משפטיים ורגולטוריים מפורטים המתאימים לסוג טכנולוגיה מסויים שעלול להשתנות במהרה, ולהפוך ללא אקטואלי, או אף להוביל להתקבעות השוק על פתרון טכנולוגי מסוים. זאת, בדומה לעקרון בעינין "גמישות" המוכר במדיניות הרגולציה בארה"ב כפי שהיא מופיעה בחוזר ה-OMB (ראו לעיל פרק "מדיניות הרגולציה בארה"ב").

ה. בחינת השימוש בכלי רגולציה מתקדמים

לאור השלב הראשוני שבו מצויות הטכנולוגיה והרגולציה בתחום הבינה המלאכותית, ועל מנת לעודד את הפיתוח והשימוש בבינה מלאכותית ולהפיק את התועלות החברתיות והכלכליות הנובעות מכך, מוצע כי במקרים המתאימים ייעשה שימוש ברגולציה המאפשרת השגת מטרות אלה. בכלל זה, מוצע לבחון שימוש בכלי רגולציה מתקדמים, כגון עקרונות אתיים לתחום הבינה המלאכותית, תקינה, המלצות רגולטור לאימוץ וולונטרי ורגולציה עצמית (מפוקחת או בלתי מפוקחת).

ככלל, כאשר רגולטור ניצב בפני טכנולוגיה חדשה המשנה מהמציאות הרגולטורית שנהגה בשוק לפני כניסתה, ובענייננו בינה מלאכותית, עומדות בפניו מספר גישות אפשריות מרכזיות לתגובה.²⁴⁹ אפשרות אחת היא חסימה של הטכנולוגיה החדשה, באמצעות קביעת כללים משפטיים האוסרים על כניסת הטכנולוגיה החדשה, או באמצעות מתן פרשנות לכללים משפטיים קיימים האוסרת את השימוש בה. אפשרות שנייה והפוכה, היא אי-התערבות בכניסת הטכנולוגיה החדשה, לאחר שנמצא שאין קושי או אינטרס פיקוחי המצדיק התערבות. אפשרות שלישית, היא יישום רגולציה קיימת, דהיינו לאפשר את כניסת הטכנולוגיה החדשה, וליישם לגביה את הכללים המשפטיים הקיימים כך שיסדירו אותה בדומה למצב שלפני כניסתה. אפשרות רביעית, היא קביעת רגולציה חדשה, כלומר פיתוח מסגרת רגולטורית חדשה ומעטפת משפטית ייעודית שתעגן אותה, כדי להסדיר את כניסת הטכנולוגיה החדשה.

בעשורים האחרונים, ועם התפתחות עולם מדיניות הרגולציה באופן כללי וביחס לרגולציה של טכנולוגיה בפרט, התפתחו שיטות רגולציה המשלבות בין אפשרויות אלה, ובדרך זו מאפשרות ומעודדות חדשנות טכנולוגית. בכלל זה ניתן למנות שימוש בעקרונות אתיים לא מחייבים המוכוונים בין היתר למגזר הפרטי כגון אלה שיוצעו להלן; בתקינה; בהמלצות של הרגולטור הניתנות לאימוץ וולונטרי ואינן מחייבות; וברגולציה עצמית. הרעיון הבסיסי בכל אלה, הוא לעודד חדשנות תחת נכונות וולנטרית לפעילות אחראית ושימת עין מצד הרגולטור והממשלה. עניין זה משמעותי במיוחד לבינה מלאכותית, עקב הצורך לנהל את הסיכונים הכרוכים באימוץ האלגוריתמי וביכולת המשתנה לצפות את תוצאות פעולתו. המגמה המסתמנת בקרב מדינות העולם היא של קידום מהיר של תקינה וקביעת כללים במקביל לפיתוח הטכנולוגי, ומגמה זו מתורגמת לדרישה מהרגולטורים לקיים תהליך בחינה, מעקב ותגובה הנדרש לצורך אימוץ

²⁴⁸ ס' 53 לטיטת הרגולציה האירופית. מעניין לציין כי הטיטה נותנת עדיפות לספקים קטנים וחברות הזנק בגישה לסנדבוקס, בסעיף 55 לטיטת הרגולציה האירופית (להלן: **טיטת הרגולציה האירופית**) <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=celex%3A52021PC0206>.

²⁴⁹ Eric Biber, Sarah E. Light, J.B. Ruhl & James Salzman, *Regulating Business Innovation as Policy* *Disruption: From the Model T to Airbnb*, 70 VANDERBLIT L. REV. 1561, 1605 (2019).

בינה מלאכותית בצורה מהירה יחסית. כתוצאה מכך, יש לבחון העשרת סל הכלים הרגולטורי בכלים המאפשרים למידה, בחינה והתנסות בשימוש בטכנולוגיה, בניהול הסיכונים הקשור בה, ובהתאמה של סל הכלים הרגולטורי אליה.

לשם ההמחשה, דרך אפשרית במקרים המתאימים לכך היא רגולציה עצמית (Self-Regulation), שמשמעותה העברת האחריות לפיתוח כללי ההתנהגות והוכחה על עמידה בהם למפתחי המערכת או לגורמים אחרים בתעשייה, לעיתים תוך שמירה על פיקוח שלטוני ברמות משתנות (Enforced Self-Regulation). רגולציה עצמית יכולה להיות מופקחת (כשאחת החלופות היא רגולטור שאוכף על התעשייה את הכללים שנקבעו על ידה עבור עצמה) ואז היא דומה יותר לרגולציה המסורתית; והיא יכולה להיות בלתי מופקחת ולמעשה "להיאכף" על ידי הגוף שאימץ אותה ביחס לעצמו באופן וולונטרי. ככלל, נראה שיהיה מתאים יותר לעשות שימוש ברגולציה עצמית כשיש אמון בין הרגולטור למפוקחים הפוטנציאליים; כשהרגולציה מותאמת באופן פרטני למפוקח מסוים (tailor made); ובעיקר כשהסיכונים הגלומים בטכנולוגיה בהקשר הקונקרטי הם מצומצמים.

יתרון מרכזי לרגולציה עצמית, הוא שחלק מפערי המידע והמומחיות בין הרגולטור למפוקחים נפתרים, עקב העובדה שהתעשייה נושאת בחלק מהאחריות להסדרת הפעילות בתחום. בנוסף, עצם הידיעה כי ככל שהרגולציה העצמית לא תיתן מענה הולם לסיכונים הרגולטור יצטרך לקבוע רגולציה שלטונית, מהווה תמריץ ליישום הרגולציה העצמית. כמו כן, יש לזכור כי לגבי חלק ניכר מהסיכונים המתעוררים בהקשרי בינה מלאכותית, יש לתעשייה תמריץ משמעותי להתמודד עמם, באופן שעשוי לייתר במקרים המתאימים רגולציה שלטונית כופה (למשל, ניתן להניח שגם יצרני כלי רכב עצמאיים (אוטונומיים) דואגים לבטיחות השימוש בהם, וגם מפתחי מערכות לקבלת החלטות אשראי מעוניינים לוודא שהן לא יפלו כדי שאכן ישתמשו בהן).

מצד שני, החיסרון העיקרי לרגולציה עצמית קשור בחשש שהכללים שייקבעו במסגרתה לא יספקו את ההגנה הנדרשת על זכויות יסוד ואינטרסים ציבוריים, כפי שהיה ראוי וניתן להגן עליהם באמצעות רגולציה שלטונית. בנוסף, ככלל יש רצון ליצור כללים שיהיו ניתנים לאכיפה, כשבמצב של רגולציה עצמית שאינה מופקחת לא תהיה אכיפה שלטונית, וגם במצב של רגולציה עצמית מופקחת רמת האכיפה שונה מחמת ההבדל במעמד הכללים ובסנקציה שלצדם. כמו כן, קיים חשש מהקניית יתרון לגופים פרטיים גדולים שיגבשו את הרגולציה העצמית, ויקבעו את ה"סטנדרט", באופן שעשוי לשמר את כוחם בשוק ולהקשות על מתחרים קטנים וחדשים.

כמצוין לעיל, קיימים גם אמצעי רגולציה מתקדמים נוספים. כך, לעניין קביעת כללים אתיים וולונטריים על ידי הרגולטור, שאף שאינם מחייבים מבחינה משפטית, פיתוח ושימוש בבינה מלאכותית תוך התחשבות בהם, תהא עדיפה משמעותית ותוכל לתת מענה לקשיים שונים שפורטו במסמך זה. כך, ניתן יהיה לפעול בכדי לעודד עמידה בתקנים הבינלאומיים המתהווים גם אם אלה עדיין אינם מחייבים.

בין היתר, ניתן יהיה לשקול קיומו של "תו תקן", שלא יהיה בהכרח מחייב, אך מי שיעמוד בו יזכה לאישור כי הוא פועל בהתאם לכללים אתיים או להנחיות וולונטריות המלמדות כי מודל הבינה המלאכותית פותח תוך התחשבות בסיכונים השונים הנובעים ממנו (זאת, למשל בדומה

למדבקות המופיעות על גבי מוצרי חשמל ומלמדות על היקף הצריכה האנרגטית או לסימונים על מוצרי מזון בדבר היותם רוויי שומן או סוכר).

עקב כך יש חשיבות לעודד פיתוח ושימוש בכלי זה כדי להשיג את יעדי ההסדרה, אולם יש לעשות בו שימוש זהיר במקרים המתאימים בשים לב לחסרונותיו, וכן לעקוב אחר יישומו כדי לבחון את הצורך בהתערבות במידה שלא ניתנת הגנה נאותה לזכויות יסוד ואינטרסים ציבוריים.²⁵⁰ יצוין כי ישנו כמובן קשר בין שימוש בכלי רגולציה מתקדמים המוצעים כאן לבין גישת ניהול הסיכונים שמוצעת לעיל. בפרט, ככל שהסיכונים הכרוכים בפיתוח ושימוש בבינה מלאכותית גדולים יותר, יהיה מקום לנקוט משנה זהירות בשימוש בכלי רגולציה מתקדמים ורכים יותר בשל החשש להתממשות סיכונים אלה, ולהיפך.

יצוין כי ועדת הבדיקה לבחינת הצורך בהתערבות ממשלתית לשם האצת התפתחות תחום הבינה המלאכותית ומדע הנתונים שהקים פורום תל"ם, בראשות ד"ר ארנה ברי, המליצה על פיתוח תקינה לבינה המלאכותית "הן מבחינת אלגוריתמיקה ומודלים והן מבחינת הנתונים". זאת, על מנת להבטיח פיתוח ושימוש אתי ובטוח בבינה מלאכותית ו"יכולת בקרה ומדידה של התוצרים ועמידתם בתקנים השונים". עבודה נמרצת בכיוון זה נעשית כיום בעולם, בין היתר על ידי ארגוני תקינה בינלאומיים כגון IEEE, ISO ו-NIST, ויש להניח כי תשמש כלי משמעותי בהסדרת התחום.

עוד יצוין, כי במסגרת קבלת הערות הציבור לטיטוט מסמך זה, עלו הצעות שונות לגבי מתכונת הרגולציה – היו שטענו שגם השימוש בכלי רגולציה מתקדמים עלול לפגוע בקידום החדשנות הטכנולוגית, והיו שטענו שלנוכח האתגרים המשמעותיים והפגיעה האפשרית בזכויות ואינטרסים אין די ב"רגולציה רכה". לנוכח ההמלצה להתבסס בשלב זה על הרגולטורים הענפיים, בשים לב למאפיינים הייחודיים של כל תחום פעילות, אין מקום להכריע בעניין זה באופן רוחבי כעת, ויש להותירו להכרעה פרטנית של הרגולטורים בתחומם.

ו. עיצוב הרגולציה תוך שיתוף הציבור ובכלל זה התעשייה, האקדמיה וארגוני חברה אזרחית

מוצע כי רגולציה המסדירה פיתוח ושימוש בבינה מלאכותית, תפותח ותקודם תוך שיתוף בעלי הידע והמומחיות בתחום, וכן בעלי העניין שצפויים להיות מושפעים ממנה. זאת, במידה הנדרשת בנסיבות העניין, ועל מנת שהרגולציה תהיה מבוססת על תשתית מקצועית וטכנולוגית איכותית ותבטא איזון בין הזכויות והאינטרסים השונים.

הטכנולוגיה משתנה תכופות, על ידי חברות בעלות אמצעים ומומחיות, המעצבות את הטכנולוגיה ובקיאות בה. ההתקדמות הטכנולוגית נעשית ככלל ללא מעורבות הרגולטורים, וחברות הטכנולוגיה נהנות מהגנה על הסודות המסחריים שלהן כלפי מתחרים. כך שהמידע בנוגע לטכנולוגיות בינה מלאכותית פעמים רבות לא מצוי באופן מלא בידי הרגולטורים, ולמעשה תחום הטכנולוגיה בכלל וטכנולוגיות הבינה המלאכותית בפרט מתאפיין במידה רבה באי-סימטריה של מידע. אף במקרים שבהם מסוגלים הרגולטורים להשיג מידע מחברות הטכנולוגיה, עלול להתעורר קושי להבין את המידע הטכני, ולחזות את השפעות הטכנולוגיות.

²⁵⁰ C(2021)99/ADD1 7-8 Practical Guidance on Agile Regulatory Governance to Harness Innovation

כך שמעבר לסוגיית קבלת המידע, המורכבות הטכנולוגית מחייבת בקיאות ומומחיות טכנית מצד הרגולטורים, והקדשת משאבים רבים לשם כך, שאינם תמיד בנמצא.

ככלל, יש צורך לעצב רגולציה על בסיס תשתית עובדתית איכותית, בשיתוף בעלי העניין, המידע והמומחיות ותוך שמיעת הציבור הרחב. מעבר לחובות הבסיסיות בעניין זה מתחום המשפט המנהלי, חוק עקרונות האסדרה קובע כי "אסדרה מיטבית" היא כזו שתהליך גיבושה וקביעתה מתבססים על נתונים הנוגעים לעניין ונעשים בהתאם לעקרון השקיפות ותוך שיתוף הציבור במידה הנדרשת בנסיבות העניין. צורך זה בא לידי ביטוי במיוחד כשמדובר ברגולציה של טכנולוגיה ובפרט של בינה מלאכותית, על רקע פערי המידע והמומחיות הנסקרים לעיל.

בהתאם לכך, רכיב זה מדגיש את הצורך בשיתוף ציבור כחלק מההתמודדות עם פערי המידע והמומחיות, ובמסגרת קיום השיח הדינמי על התפיסות הערכיות ביחס לשימושים בטכנולוגיה חדשה. מטבע הדברים, הכוונה היא לשיתוף הציבור במידת האפשר ובשים לב למכלול הנסיבות של הציבור הרלוונטי ובעלי העניין, אך ראוי לשים דגש גם על חברות קטנות ובפרט חברות ההזנק (סטארט-אפ), ארגוני חברה אזרחית והציבור בכללותו, כדי להגביר את השקיפות, אמון הציבור, והמומחיות בתחום. לצד זאת, יש לפעול לחיזוק היכולות הטכנולוגיות והמומחיות אצל גורמי הממשלה והרגולטורים, על מנת שיוכלו לגבש עמדה מקצועית עצמאית.

6.2. אימוץ עקרונות אתיים לתחום הבינה המלאכותית

מוצע לאמץ עקרונות אתיים משותפים עבור תחום הבינה המלאכותית, על בסיס עקרונות ה-OECD, כדי לסייע לרגולטורים ולארגונים בתחום זה. מוצע כי העקרונות לא יהיו בעלי מעמד משפטי מחייב, אך ישמשו בסיס לשיח משותף ותיאום לפי ההקשר הנורמטיבי והרגולטורי.

6.2.1. החשיבות של אימוץ עקרונות אתיים ומעמדם

ראשית, העקרונות מאפשרים שפה משותפת, הן בין ענפי המשק השונים, הן מול הנעשה במדינות העולם המפותחות לגבי אופן ההתמודדות האחראי עם חששות וסיכונים מוכרים הכרוכים במערכות מבוססות בינה מלאכותית, לרבות האתגרים שנסקרו לעיל במסמך זה.

שנית, העקרונות מאפשרים לגשר בין ערכים, יישומים טכנולוגיים והיבטים משפטיים, תוך ממשק לתפיסה הבינלאומית המתהווה של "בינה מלאכותית מהימנה". אימוץ ופרסום העקרונות במסגרת המדיניות מאפשר שקיפות ושיח ציבורי לגבי העקרונות כנקודת מוצא למדיניות רגולציה לתחום הבינה המלאכותית. בנוסף, לעקרונות חשיבות כמכנה משותף לשיח בינלאומי בנושא זה.

שלישית, העקרונות מסייעים לקידום מדיניות רגולציה קוהרנטית בתחום הבינה המלאכותית, בכך שהם מכוונים את הרגולטורים, את גורמי הממשלה ואף את גורמי התעשייה השונים להיבטים

המרכזיים שיש להידרש אליהם, ומגבירים את הוודאות בקרב המגזר הציבורי והפרטי בעניין המדיניות הממשלתית לגבי בינה מלאכותית.

אין כוונה להקנות לעקרונות אלה מעמד משפטי, ובהתאם לכך הם לא מחייבים גורמים פרטיים או ציבוריים לפעול או שלא לפעול בדרך כלשהי, ואינם מחליפים את המסגרת המשפטית החלה או מהווים כלי לפרשנות משפטית (עם זאת, עקרונות אלה משקפים היבטים בהם ראוי להתחשב במסגרת פיתוח ושימוש בבינה מלאכותית ובעת קביעת רגולציה בתחום זה).

6.2.2. התבססות על עקרונות ה-OECD

בהמשך להחלטת ממשלה 212, העקרונות שמוצע לאמץ מבוססים בעיקר על העקרונות המוצעים בהמלצות ה-OECD לקידום בינה מלאכותית מהימנה,²⁵¹ תוך התאמה שלהם ליעדי המדיניות הישראלית ובהתבסס על הנעשה במדינות מפותחות, לרבות התאמות שבוצעו במדינות אלה.

עד כה נכתבו בעולם מסמכים שונים הכוללים עקרונות אתיים או עקרונות משפטיים שהם רלוונטיים לתחומי העיסוק בבינה מלאכותית. מחקרים מונים עשרות מסמכים ובהם עקרונות לבינה מלאכותית מהימנה או אתית שנערכו על ידי ארגוני חברה אזרחית, חברות טכנולוגיה, ארגונים פרטיים, ממשלות, וארגונים בינלאומיים.²⁵² ועדת המשנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה גיבשה אף היא רשימת עקרונות אתיים שיש להתחשב בהם.²⁵³

לצד זאת, לעקרונות ה-OECD שנסקרו לעיל (ראו פרק "המלצות ה-OECD בתחום הבינה המלאכותית") חשיבות ייחודית, משום שגובשו בתיאום בין המדינות המפותחות מתוך שותפות ורצון לעשות שימוש מועיל בבינה מלאכותית באופן שמכבד את עיקרון שלטון החוק, ערכים דמוקרטיים ואינטרסים ציבוריים. מדינות אלה הסכימו על עקרונות אלו כקווים מנחים בתחום זה, תוך גישור על השוני בהיבטים המשפטיים, הרגולטוריים והכלכליים ביניהן. עקב כך העקרונות משקפים שיח מפותח שיש לו ערך לגיבוש מדיניות פנים מדינתית. בנוסף, יש לו חשיבות במרחב הבינלאומי כמכנה משותף מוסכם בין המדינות המפותחות, לרבות במסגרת שיתוף הפעולה בין האיחוד האירופי וארה"ב כמתואר לעיל (ראו פרק בעניין "המלצות ה-OECD").

6.2.3. העקרונות האתיים המוצעים לאימוץ

מוצע לאמץ את ששת העקרונות האתיים המשותפים לתחום הבינה המלאכותית המפורטים להלן. לאור השונות באופני השימוש בבינה מלאכותית, מרכיב מרכזי באופן התחולה של העקרונות הוא יישום לפי הקשר ובהתאם לתפיסות מקובלות של ניהול סיכונים.

א. בינה מלאכותית תשמש לקידום צמיחה, פיתוח בר-קיימא ומובילות ישראלית בחדשנות

²⁵¹ ראו המלצות ה-OECD, לעיל ה"ש 8.

²⁵² מחקר של אוניברסיטת הרווארד משנת 2020 סקר מסמכים אתיים רבים שהופקו בידי ארגונים שונים, את העקרונות הכלולים בהם, ואת אופן פרשנותם. ראו Jessica Fjeld, Nele Achten, Hannah Hilligoss, Adam Nagy & Madhulika Srikumar, *Principled Artificial Intelligence: Mapping Consensus in Ethical and Rights-Based Approaches to Principles for AI*, 2020-1 Berkman Klein Center Research Publication. 1 (2020); <https://papers.ssrn.com/sol3>

לסקירה משנת 2022 ראו Nicholas Kluge Corrêa, Camila Galvão, James William Santos, Carolina Del Pino, Edson Pontes Pinto, Camila Barbosa, Diogo Massmann, Rodrigo Mambrini, Luiza Galvão, Edmund Terem & Nythamar de Oliveira, *Worldwide AI Ethics: a Review of 200 Guidelines and Recommendations for AI Governance*, CORNELL UNIVERSITY (Jun. 23 2022), <https://arxiv.org/abs/2206.11922>.

²⁵³ דוח ועדת המשנה לעיל ה"ש 21.

שימוש אחראי בבינה מלאכותית מהימנה הוא אמצעי לעודד צמיחה, פיתוח בר-קיימא, רווחה חברתית וקידום המובילות הישראלית בתחום החדשנות.

עיקרון זה מצביע על החשיבות של שימוש אחראי בבינה המלאכותית להשגת יעדי המדיניות שמוגדרים בהחלטת ממשלה 212 ויעדי מדיניות מרכזיים נוספים.

גישה זו גם מכירה בכך שחדשנות טכנולוגית כגון טכנולוגיית הבינה המלאכותית יכולה לתרום משמעותית לרווחת הפרט, ולתחומים חשובים כגון בריאות הציבור, או צמיחה בת-קיימא. עקב כך, יש צורך במבט מאוזן המביא בחשבון הן את הסיכונים והן את התועלות, ולבחון את אלה למול נקודת המוצא החברתית-טכנולוגית הקיימת.

בעולם של שווקים טכנולוגיים גלובליים, ולנוכח השלב בו נמצאת הטכנולוגיה, יש לתת את הדעת לחשיבות של שימור עוצמה טכנולוגית ויכולת טכנולוגית. זאת, תוך פעולה במסגרת ערכית ואתית. על רקע המאפיינים של הפיתוח והשימוש בטכנולוגיה בשווקים גלובליים, ועל רקע מקומה של ישראל בשווקים אלה, יש לקחת בחשבון את הקשר בין עיקרון זה לבין העקרונות הנוספים.

אופיו של עקרון זה חורג מהשימוש הרגיל במונח "אתיקה", והנושא עלה גם במסגרת הערות הציבור לטיוטת המסמך. עם זאת, לנוכח חשיבותו של העיקרון נראה שנכון למנות אותו באותה רשימה עם יתר העקרונות, ולהימנע מאבחנה מסורבלת בין סוגים שונים של עקרונות מנחים.

ב. האדם במרכז: כיבוד זכויות יסוד ואינטרסים ציבוריים

פיתוח בינה מלאכותית, או שימוש בה, יש לעשות תוך כיבוד שלטון החוק, זכויות יסוד ואינטרסים ציבוריים ובפרט תוך שמירה על כבוד האדם ועל הזכות לפרטיות.

עיקרון זה מציין את נקודות המוצא הנורמטיביות בישראל ובמדינות המפותחות הכוללות את שלטון החוק והמשטר הדמוקרטי, וכן הגנה על זכויות יסוד ואינטרסים ציבוריים.²⁵⁴ הביטוי "האדם במרכז" ("Human-centered AI") מהווה מחולל מרכזי לצורך פירוש ופיתוח העקרונות המשפטיים והרגולטוריים בתחום הבינה המלאכותית.²⁵⁵ תפיסה זו נועדה להזכיר כי בעת פיתוח ושימוש בטכנולוגיות חדשניות, האדם, ויחסי הגומלין שלו (במגוון הקשרים) עם החברה בה הוא פועל, מהווים נקודת מכוון נורמטיבית.

תפיסת "האדם במרכז" נועדה להבטיח שבעת פיתוח ושימוש בבינה מלאכותית, ובפרט הישענות על חישובים אלגוריתמיים מורכבים המייצרים החלטות או תחזיות, תינתן תשומת הלב הנאותה לערכים הרלוונטיים להקשר שבו נעשה השימוש.

עקרון זה נועד להעצים, לפי ההקשר והסיכון, את כיבוד הערכים הקיימים אגב פיתוח בינה מלאכותית ושימוש בה; בין היתר מתוך הכרה בכך ששימוש בכלי בינה מלאכותית יכול לשמש לקידום ההגנה על זכויות האדם.

לצד הערכים הכלליים, מצוינים בסעיף במפורש גם "כבוד האדם" ו"הזכות לפרטיות". לכבוד האדם מעמד חשוב ויסודי כעוגן לזכויות יסוד,²⁵⁶ והוא מופיע בעקרון זה כביטוי לתפיסת

²⁵⁴ ראו גם ס' 1 לחוק עקרונות האסדרה, התשפ"ב-2021.

²⁵⁵ ראו את האמור במסמך הלוואי להמלצות ה-OECD, הכולל התייחסות לשיקולים השונים שהוצגו בדיונים המכינים ובניסוח ההמלצות, AI in Society, לעיל ה"ש 7, פרק 4 <https://www.oecd-ilibrary.org/sites/>.

²⁵⁶ ראו חוק-יסוד: כבוד האדם וחירותו.

"האדם במרכז" ועל מנת לבטא את האוטונומיה והרצון של הפרט.²⁵⁷ הזכות לפרטיות מוזכרת באופן ספציפי לאור היותה זכות משמעותית בפעולות הקשורות ב"נתוני עתק" (איסוף ועיבוד מידע) ובעת קבלת החלטות על סמך מידע אישי. עקב כך, להגנה על הפרטיות רלוונטיות ניכרת לעקרונות האתיים המקובלים ביחס לבינה מלאכותית. עקרונות אלה מבקשים לצמצם איסוף מידע מזוהה, להבחין בין הגנה על מידע רגיל ומידע רגיש ולוודא כי מידע נאסף בהסכמה ולמטרה מוגדרת.²⁵⁸

ג. שוויון ומניעת אפליה פסולה

בפיתוח ושימוש בבינה מלאכותית יש להביא בחשבון את הצורך בשוויון וגיוון, את הפוטנציאל הגלום בבינה מלאכותית לקידום שוויון מצד אחד, ואת החשש להטיה במערכות בינה מלאכותית והסיכון לאפליה פסולה כנגד יחידים או קבוצות מצד שני.

עיקרון זה משקף את התפיסה הבסיסית של הוגנות והתבונה המצטברת בדבר החששות להתרחשות אפליה פסולה בעקבות השימוש בבינה מלאכותית, שלעיתים אף תהיה בלתי חוקית. בכל הקשור לבינה מלאכותית, נעשה שימוש בהוגנות (fairness) כעיקרון הנוגע למניעת אפליה, ואשר נועד לבטא את הצורך בבחינה של מכלול הנסיבות הצריכות לעניין, על מנת לחשוף חשש להטיה אסורה, וזאת לנוכח מגוון התרחישים הטכנולוגיים והחברתיים הרלוונטיים.²⁵⁹ בין היתר, כמפורט בפרק העוסק באתגר האפליה לעיל, חששות אלה כוללים את האפשרות להטיה בנתונים המשמשים לאימון המערכת או שימוש במאפיין רגיש באחד המשתנים הנאספים באופן שעלול להביא לתוצאות מפלות, וכן את האפשרות להטיות מושרשות בתחום המומחיות המשמש לאימון הבינה המלאכותית.²⁶⁰ הזכות לשוויון מוזכרת בין השאר עקב החשש כי חישובים אלגוריתמיים עלולים לשקף הטיות בחברה.²⁶¹

מדובר בעיקרון שנועד לבטא את הצורך להביא בחשבון את ההיבטים הנוגעים לשוויון ולמניעת אפליה במסגרת פיתוח ושימוש בבינה מלאכותית. בהקשר הנוכחי, הכוונה אינה להגדיר

²⁵⁷ דוח ועדת המשנה, לעיל ה"ש 21, בעמ' 13:

"כיבוד זכויות אדם והגנה עליהן - כיבוד כלל זכויות האדם ומודעות לצורך הגובר בהגנה על זכויות מסויימות שעלולות להיפגע יותר בעידן הבינה המלאכותית וביניהן:

א. שמירה על שלמות הגוף - הגנה מפני פגיעה בחייו או בגופו של אדם.
ב. פרטיות - לרבות מניעת פגיעה בפרטיות באיסוף המידע, ניתוח ועיבוד המידע, שיתוף המידע ובשימושים חדשים בו.
ג. השמירה על האוטונומיה - שמירה על חירות האדם. בין השאר, מניעת השפעה שאינה הוגנת על התנהגות היחיד, המרחב הפרטי שלו וזכותו "להיעזב לנפשו".
ד. זכויות אזרחיות ופוליטיות - לרבות חופש הביטוי וחופש הדת והמצפון".

²⁵⁸ AI in Society, לעיל ה"ש 7, בעמ' 87-89.

²⁵⁹ שם, בעמ' 90; ראו גם בטיטת ההנחיה הבריטית:

Embed considerations of fairness into AI

In many contexts, the outcomes of the use of AI can have a significant impact on people's lives - such as insurance, credit scoring or job applications. Such high-impact outcomes - and the data points used to reach them - should be justifiable and not arbitrary.

In order to ensure proportionate and pro-innovation regulation, it will be important to let regulators continue to define fairness. However, in any sector or domain we would expect regulators to:

- interpret and articulate 'fairness' as relevant to their sector or domain, decide in which contexts and specific instances fairness is important and relevant (which it may not always be), and design, implement and enforce appropriate governance requirements for 'fairness' as applicable to the entities that they regulate

²⁶⁰ מאפיין רגיש הינו או מאפיין כגון דת, גזע, מין, וכו' או נתון אחר שיש לו זיקה לאותו מאפיין (פרוקסי).

²⁶¹ להרחבה ראו אתגר האפליה לעיל וכן: EPRS Study 2022, לעיל ה"ש 74, בעמ' 1.

במדויק את אופן תחולת הזכות לשוויון ואיסור להפלות, כי אם להצביע על הסיכון המוגבר לכך עקב מאפייני הטכנולוגיה, ולהגביר את מודעות בעלי העניין לכך.²⁶²

היקף השתרעות הזכות לשוויון או הגדרת אפליה פסולה, נובע, בין היתר, מאופי השימוש בבינה המלאכותית ומידת רגישותו.²⁶³ כך שבמקרים המתאימים, יהיה צורך לנקוט אמצעים, בהתאם להקשר ולרמת הסיכון, להתמודדות עם חששות אלה. למשל, ייתכן שיהיה מקום לוודא שאין הטיה גלויה בנתונים, וכן לבצע בקורות ובדיקות כדי לבדוק את הסיכון להטיה סמויה. מכל מקום, חשוב להדגיש, כי עקרון השוויון ראוי לבוא לידי ביטוי גם בתוצאה – כך שהתוצאה לא תהיה מפלה, ולא רק הדרך אליה – בשלבי הפיתוח ואיסוף המידע.

התמודדות עם הסיכון כראוי, כמו גם המודעות לו, יכולה להביא לשיפור קבלת ההחלטות הכוללת, תוך מיקסום התועלות מבינה מלאכותית והפחתת הסיכון להטיה.

ברמה המעשית, לצורך התמודדות עם החשש להטיה, יש חשיבות לניהול סיכונים לאורך כל "מחזור החיים" של מערכת בינה מלאכותית. זאת, בין היתר, באמצעות התמודדות עם הטיה באיסוף ועיבוד מקדים (pre-processing) של הנתונים; בשלב בניית המודל, האימון והערכה; ובשלב היישום ובדיקת ההשפעה על יחידים וקבוצות בעת ההפעלה.²⁶⁴

ד. שקיפות והסברות

בפיתוח בינה מלאכותית ושימוש בה, יש להביא בחשבון, בשים לב למגוון השיקולים הנוגעים לעניין, את הצורך ליידע מי שבא במגע עם בינה מלאכותית או מושפע מפעילותה ואת הצורך במתן הסבר להחלטתה או לאופן שבו פעלה, בין היתר בהתחשב במידת השפעתה, השלכותיה על מי שמושפע ממנה והאפשרויות הטכנולוגיות הזמינות.

עיקרון השקיפות וההסברות נועד להתמודד עם האפשרות שאדם יקיים אינטראקציה עם בינה מלאכותית ולא ידע על כך, וכן עם הסיכונים האינהרנטיים הנובעים מהמפגש בין יכולת ההבנה האנושית וקבלת ההחלטות האלגוריתמית.²⁶⁵ אמנם בני האדם הם אלה שמייצרים את הבינה המלאכותית, אולם היא אינה מובנת לרוב הציבור, וגם למומחים לבינה מלאכותית ייתכן שיהיה קושי להסביר אופן קבלת החלטה קונקרטית או תוצאותיה. שכן כמפורט לעיל, בעולמות הבינה המלאכותית נעשה שימוש תדיר בדימוי של קבלת ההחלטה האלגוריתמית כ"קופסא שחורה", שבה לא ניתן להבין את אופן קבלת ההחלטה או "לנהל שיח" עם מקבל ההחלטה על מנת לעמוד על התשתית העובדתית ושיקול הדעת המקצועי שהפעיל.

היבט זה מקבל משנה חשיבות כשבינה מלאכותית מעורבת בהחלטה שיש בה פגיעה בזכויות של אדם. במקרים אלה גוברת החשיבות לקיומו של מידע על עצם השימוש בבינה מלאכותית,

²⁶² שם.

²⁶³ הדיון על אודות עיקרון השוויון ואיסור האפליה במשפט הישראלי הינו רחב ומורכב. ראו למשל במשפט החוקתי והמנהלי: ברק ארז, **המשפט המינהלי**, ב' 673-721 (2010); לסקירה על המשפט הפרטי וביקורת עליה: אברהם, לעיל ה"ש 39; דיני העבודה: שרון רבין מרגליות, "שלושה דורות של אפליה תעסוקתית: הישגים ומגבלות של המאבק לקיום שוויון תעסוקתי", **עבודה חברה ומשפט** טו 7, 7 (2018).

²⁶⁴ ראו EPRS Study 2022, לעיל ה"ש 74, בעמ' 13, 6, 1.

²⁶⁵ דוח ועדת המשנה, לעיל ה"ש 21, הפרד בין שני עקרונות אלה, כמפורט להלן: א. **שקיפות** – הנגשה, בהתאם להקשר ולנסיבות, של מידע על התהליך עצמו ועל דרך קבלת ההחלטה. לכל הפחות יש לאפשר קבלת מידע על התהליך, על מנת לקבוע אם הוא עונה על העקרונות האתיים.

ב. **הסברות** – הסבר הפעולה וההחלטה – היכולת להסביר את תהליך קבלת ההחלטה של המערכת (ברמת המשתמשים כפרטים, גם ברמת הכלל אם המערכת משפיעה על קבוצות, וגם עבור מפעילי המערכת עצמם); בדו"ח פורום תל"ס, לעיל ה"ש 22, גם הופרדו שני עקרונות אלה בעמ' 62.

וכן על הפרמטרים והלוגיקה ששימשו בסיס להחלטה או הפעולה של הבינה המלאכותית. בנוסף, הסיכון והקושי בהקשר זה גוברים בנסיבות שבהן עקב ההתייעלות הכרוכה במיחשוב ושימוש במערכת מבוססת בינה מלאכותית, יש המרה של תהליך שבו מעורב אדם לתהליך טכנולוגי בלבד ללא מעורבות אנושית (כגון רכישת מוצר או שירות, קבלת ייעוץ או המלצות).

על כן, בהתאם לעיקרון זה, ראוי כי ככלל מפתח בינה מלאכותית או משתמש בה יגלה למי שמקיים עמה אינטראקציה או מושפע מפעילותה את דבר השימוש בבינה מלאכותית, וכן ינגיש, במידת האפשר והנדרש בנסיבות העניין, מידע אודות מאפייני ההחלטה של המערכת. כן ראוי, כי יינקטו אמצעים סבירים, בהתאם למקובל בתחום, כדי לאפשר להתחקות אחר טעמי ההחלטה. זאת, בשים לב לאילוצים הטכנולוגיים, ותוך איזון בינם לבין התכלית הניצבת ביסוד הרצון לספק הסבר על אופן פעילות המערכת.

בהתאם למסמך המלווה את המלצות ה-OECD, עיקרון ההסברות מיועד להתמודד עם האתגר האמור ועוסק בהבנה של תוצאה ספציפית של החלטה או תחזית של מערכת בינה מלאכותית.²⁶⁶ המסמך המלווה מציין כי להסברות יש מחירים ועלויות ולעיתים מערכות בעלות תוצאות יותר מדויקות הן מורכבות יותר להסבר. בהתאם לכך, יש לבחון את היחס שבין העלות לתועלת לפי כל הנסיבות, כגון למשל רגישות ההחלטה והשימוש.²⁶⁷

יצוין כי גם בהתאם לדיני הפרטיות בישראל ובעולם, כאשר פעילות הבינה המלאכותית מבוססת על איסוף או עיבוד מידע אישי, ייתכנו חובות ספציפיות בנושא זה.²⁶⁸

עיקרון השקיפות וההסברות רלוונטי גם ליחסים שבין יצרן בינה מלאכותית לבין המשתמש בה, בהם עשוי להיות פער מידע אודות מאפיינים אלה. סגירת פער מידע זה בהתאם לעיקרון השקיפות וההסברות מאפשר לארגון המשתמש לעמוד מצדו בחובות החלות עליו כלפי מי שמושפע מהחלטותיו. בנוסף, קיום עיקרון השקיפות וההסברות מעצים את המשתמש בשיח עם היצרן, ומאפשר לו לנהל שיח עשיר יותר לגבי מאפייניה ואופן השימוש הנכון בה, חלוקת הסיכונים ביחס אליה ומעקב אחר השימוש בה. השקיפות וההסברות מאפשרות לארגון המשתמש גם להתאים את אופן השימוש והבקרה שלו על השימוש, ובכלל זה עמידה שלו בעיקרון השקיפות וההסברות כלפי מי שמושפע מהשימוש שלו בטכנולוגיה.

ה. אמינות, עמידות, אבטחה ובטיחות

בפיתוח ושימוש במערכות מבוססות בינה מלאכותית יש להביא בחשבון את הצורך בכך שהן תהיינה אמינות, מאובטחות ובטיחותיות, כך שבתנאי שימוש צפוי או שימוש שגוי או בתנאים מסוכנים אחרים, הן יפעלו כראוי ולא יהוו סיכון בטיחותי בלתי סביר. לשם כך יש לנקוט אמצעים סבירים בהתאם לתפיסות מקצועיות מקובלות של ניהול סיכונים, על מנת לצמצם סיכוני בטיחות ואבטחת מידע לכל אורך "מחזור החיים" של המערכות.

עיקרון זה משקף את הסיכונים הנובעים מטכנולוגיות המידע. כמפורט לעיל בפרק "סוגיות ואתגרים המתעוררים בקשר לבינה המלאכותית", הניסיון המצטבר בתחום אבטחת המידע והגנת הסייבר מלמד על חשיבות ההתמודדות עם סיכונים אינהרנטיים הנובעים משימוש

²⁶⁶ AI in Society, לעיל ה"ש 7, בעמ' 93.

²⁶⁷ שם, בעמ' 95.

²⁶⁸ ראו עמדת הרשות להגנת הפרטיות בעניין חובת יידוע במסגרת איסוף ושימוש במידע אישי, לעיל ה"ש 27; כן ראו אחיעז, חמדני, עמירם וקסטיאל, לעיל ה"ש 25, בעמ' 42-43 על סקירה על הוראות החקיקה האירופית.

במחשבים ותקשורת, ועם החשש לשימוש לרעה בידי גורמים זדוניים המנצלים "חולשות" בטכנולוגיה. מאפיינים אלה רלוונטיים גם לתחום הפיתוח והשימוש בבינה מלאכותית.

בנוסף לסיכונים אלה, הבינה המלאכותית חשופה לסיכונים ייחודיים הקשורים באופן אימון המערכת, החשש מפני מניפולציה שתבוצע כלפיה, וסיכון הנובע מאופן השימוש בה. היקף השימוש הצפוי בבינה מלאכותית, במוצרים שכיום אין להם יכולת חישובית (כגון מכונות ומוצרים מסוגים שונים), מרחיב את היקף הרלוונטיות של סוגיות בטיחות הנובעות מסיכונים ממוחשבים. כך, כל עוד מדובר במוצרים שיש בהם רכיבים פיזיים בלבד, הסיכון הבטיחותי מהם נובע מאופן פעולתם הפיזי. שילוב רכיבים ממוחשבים מייצר סיכון מסוג חדש. זאת, בפרט שבינה מלאכותית עשויה ללמוד ולהשתנות תוך כדי ההפעלה שלה לאחר שלב הייצור.

על מנת לקדם את הוודאות והיכולת לנהל את סיכוני הבטיחות, יש חשיבות בהעברת המידע הנדרש על מאפייני הבטיחות בין הגורמים השונים המשתמשים בבינה מלאכותית, על מנת לאפשר ניהול סיכונים אחיד. בנוסף, בדומה לפיתוח תוכנה, יש מקום לכך שהליך עיצוב ופיתוח בינה מלאכותית יאפשר בחינה ובדיקה לאחר, באמצעות תיעוד.

לצד חשיבותו העצמאית שתוארה לעיל, עקרון זה תומך ביכולת לבחון את קיום העקרונות הנוספים. זאת, משום שקיום תהליכי פיתוח ובדיקה מהימנים של טכנולוגיה מאפשרים גם מעקב טוב יותר אחרי אופן קיום העקרונות הנוספים המפורטים בפרק זה.

1. אחריותיות

מפתחי בינה מלאכותית, מפעיליה או המשתמשים בה יגלו אחריות לתפקודה התקין, ולקיום העקרונות האתיים האחרים בפעילותם, בין היתר בשים לב לתפיסות מקובלות של ניהול סיכונים ולאפשרויות הטכנולוגיות הזמינות.

עיקרון האחריותיות נועד לבסס את קיומו של "גורם אחראי" לבינה מלאכותית, שעשוי להיות למשל המפתח או המפעיל של מערכת מבוססת בינה מלאכותית. העיקרון המוצע מצביע על כך שמערכות מבוססות בינה מלאכותית נוצרות בידי אדם, ולכן כברירת מחדל אדם הוא שאחראי להן.²⁶⁹ בכך העיקרון מבטא גם את חשיבות המעורבות האנושית בפעילות בינה מלאכותית.

אחריות זו משמעה בין היתר הצורך לנקוט אמצעים לתפקודה התקין, לפעול לצמצום החשש לפגיעה בזכויות יסוד ובאינטרסים ציבוריים, ולשאוף לקיום העקרונות שתוארו לעיל. אופן ההתמודדות עם הסיכונים ויישום העקרונות מבוסס על אמצעים סבירים ויכולת צפיות סבירה, לפי מידת ההתפתחות הטכנולוגית. העיקרון משרת גם את היכולת לבצע בקרה בידי גורם בלתי תלוי, במקרים הרלוונטיים, בנוגע לאופן הפיתוח של בינה מלאכותית והשימוש בה.

עיקרון זה מבטא אמון בגורמים האמורים, ומאפשר להם לפעול בתנאי שניהגו כמי שאחראיים לבינה המלאכותית. הוא מאפשר לגשר בין העקרונות המפורטים לעיל, לבין החדשנות המבוצעת בפועל בידי ארגונים המפתחים או משתמשים בבינה מלאכותית.

Joanna J. Bryson, *The Artificial Intelligence of the Ethics of Artificial Intelligence: An Introductory*²⁶⁹
Overview for Law and Regulation, in: THE OXFORD HANDBOOK OF ETHICS OF AI 3, 4-5 (Markus D.
Dubber, Frank Pasquale & Sunit Das eds., 2020)

עיקרון האחראיות משקף את התפיסה הבסיסית שלפיה ראוי כי מי שיוצר את הסיכון יתמודד איתו ויהיה אחראי לו, בלי להחצין את הסיכון לצדדים אחרים. בנוסף, עיקרון זה ותפיסה הוליסטית לגבי ניהול הסיכונים בארגון תורמת להעלאת החוסן הארגוני והיכולת להשתמש במלוא התועלות הטכנולוגיות.²⁷⁰ לעיקרון זה זיקה הדוקה לתפיסות מקובלות של ניהול סיכונים כמפורט בהמלצה העוסקת בפיתוח והנגשה של כלי כזה להלן.

6.3. מיסוד מוקד ידע ותיאום ממשלתי להסדרת בינה מלאכותית

בהחלטת ממשלה 173 נקבע כי על משרד החדשנות, המדע והטכנולוגיה להקים מוקד ידע ותיאום ממשלתי בתחומי הבינה המלאכותית אשר יעסוק בנושאי רגולציה, מדיניות מידע ונתונים, אתיקה, שיתוף פעולה בינלאומי אזרחי והטמעה במגזר הציבורי האזרחי; וזאת בשיתוף מחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים (להלן: "מוקד הידע").

מוצע לפעול בהקדם להקמת ומיסוד מוקד הידע האמור, לנוכח ההתפתחויות בתחום הבינה המלאכותית ופוטנציאל ההשפעה הרב של מערכות מבוססות בינה מלאכותית, ובשים לב לחסרונות הכרוכים בהחלטה שלא לקדם בשלב זה חקיקת מסגרת רגולטורית רוחבית (בפרט החשש מפערים, חפיפות וסתירות ברגולציה של הרגולטורים הענפיים שיסדירו בינה מלאכותית בתחומם).

בהקשרים בהם עוסק מסמך זה – רגולציה ואתיקה – מוצע כי מוקד הידע ישמש כגוף פנים-ממשלתי מקצועי ובעל מומחיות, אשר ימלא, בין היתר, את התפקידים הבאים:

- א. ייעוץ לרגולטורים הענפיים בבחינת הצורך ברגולציה בתחום הבינה המלאכותית בענף עליו הם מופקדים, ובאפיון וקידום רגולציה כזו ככל שהיא נחוצה;
- ב. קידום תכנון, תיאום ושיתוף-פעולה בין רגולטורים בתחום הבינה המלאכותית, כדי לצמצם חפיפות וסתירות בין רגולציה שנקבעת על ידי רגולטורים שונים, ולהגביר ודאות ואחידות;
- ג. הובלת מהלכים מתואמים ורוחביים למימוש מדיניות הממשלה בתחום הבינה המלאכותית, ובפרט המדיניות המוצעת במסמך זה;
- ד. ייעוץ לממשלה בענייני רגולציה ומדיניות רגולציה בתחום הבינה המלאכותית, ובכלל זאת בעניינים הנוגעים לעיצוב, עדכון ופיתוח המדיניות שבמסמך זה, וכן מעקב אחר יישומה;
- ה. הובלת הייצוג והמעורבות של ישראל בפיתוח רגולציה ותקינה בפורומים בינלאומיים;
- ו. הנגשה ופיתוח של מידע וכלים לשימוש אחראי בבינה מלאכותית עבור הרגולטורים והציבור.
- ז. הקמת פורומים רלוונטיים לקיום שיח ושיתוף ידע עם גורמי תעשייה, אקדמיה, מגזר שלישי וממשל, כמפורט בהמשך ההמלצות.

²⁷⁰ ראו ממצאים מסקר גרטנר בנושא זה, <https://blogs.gartner.com/avivah-litan/2022/08/05/ai-models-under-attack-conventional-controls-are-not-enough>.

לנוכח תפקידיו המשמעותיים, מוקד הידע נדרש להיות מורכב מעובדי מדינה בעלי מומחיות בתחומי מדיניות רגולציה, טכנולוגיה ומשפטים. כמו כן, עליו להחזיק ביכולת המעשית, בהיבטים של תקנון ותקצוב, לתת מענה אפקטיבי לאתגרים השונים בתחומי הרגולציה של בינה מלאכותית. זאת, על מנת שיוכל להיות בעל ראייה רוחבית ולהחזיק באחריות לעידוד פיתוח ושימוש במערכות מבוססות בינה מלאכותית תוך שמירה על זכויות ואינטרסים, והשגת היעדים הממשלתיים.

על מנת לכוון את מוקד הידע בהתאם למכלול הזכויות והאינטרסים הנוגעים לעניין, להקנות לו ראיית רוחב על הנעשה בממשלה ובמשק הישראלי, לחזק את מומחיותו ויכולותיו, ולהעצים את השפעתו – מוצע להקים למוקד הידע ועדת היגוי, שבה יכהנו הנציגים הבאים, או מי מטעמם: מנכ"ל משרד החדשנות, המדע והטכנולוגיה (שיכהן כיו"ר הוועדה); מנהל מוקד הידע; המשנה ליועץ המשפטי לממשלה (משפט כלכלי) במשרד המשפטים; הממונה על התקציבים במשרד האוצר; ראש רשות הרגולציה; ראש רשות הגנת הפרטיות; ראש מערך הדיגיטל הלאומי; ומנכ"ל רשות החדשנות.

תפקידי ועדת ההיגוי יהיו, בין היתר, להמליץ לשר החדשנות, המדע והטכנולוגיה על אסטרטגיה לאומית בתחום הבינה המלאכותית; לבחון את הצורך בשינוי מדיניות הרגולציה בתחום הבינה המלאכותית; לייצץ לעניין תכנית העבודה השנתית של מוקד הידע; להמליץ על יעדי מוקד הידע לטווח הזמן הבינוני.

השימושים בבינה מלאכותית מגוונים והיא מוטמעת כמעט בכל ענף במשק, ובכלל זה בריאות, תחבורה, פיננסים, חקלאות וחינוך. גיוון השימושים הוא השיקול המכריע שבגינו מומלץ במסגרת מדיניות רגולציה, כפי שפורט לעיל, שלא לאמץ חקיקת מסגרת רוחבית שתסדיר את כל השימושים בבינה מלאכותית, ולקבוע הסדרים רגולטוריים לפי הנדרש על בסיס הצורך הקונקרטי בכל ענף ובהובלת הרגולטור האמון עליו (ראו לעיל בעניין מדיניות הרגולציה המוצעת).

במקביל, לנוכח היותה של הבינה המלאכותית תחום טכנולוגי חדשני אשר מתפתח במהירות ומחייב ידע המקצועי בתחומים טכנולוגיים ורגולטוריים, וכן לנוכח חשיבות התאימות עם הנעשה בתחום בעולם, אשר מחייבת היכרות עם מדיניות הרגולציה במדינות מובילות ובארגונים בינלאומיים, יש טעם במיסוד גורם ממשלתי אשר יסייע לרגולטורים הענפיים בבחינת הצורך ברגולציה בתחום הבינה המלאכותית בענף עליו הם מופקדים ובקידום רגולציה כזו ככל שהיא נחוצה.

כמו כן, לנוכח הבסיס המשותף של האתגרים הרגולטוריים הרוחביים בתחומי הבינה המלאכותית כפי שנמנו בחלקם לעיל, כגון אחריותיות, אפליה, מעורבות אנושית או הסברתיות, יש צורך לתאם בין הרגולטורים הענפיים ולהתאים את הרגולציה הענפית למדיניות ממשלתית אחידה. תיאום והתאמה אלו יאפשרו, בין היתר, לוודא שאכן נעשית רגולציה ענפית בתחומים בהם היא נחוצה ולסייע לרגולטורים לשם כך; לשמור על הקוהרנטיות של הרגולציה הענפית; למנוע חפיפות בין רגולטורים או סתירות בין כללי רגולציה; לתאם היכן שקביעת רגולציה בתחום או ענף אחד עשויה להשפיע על תחומים נוספים, ולהגביר את הוודאות והאחידות; וכן להוביל מהלכים מתואמים ורחבים להתמודדות עם אתגרים ושימושי בינה מלאכותית חדשניים חוצי-ענפים.

תיאום מובנה זה נועד להתגבר על האתגר הטמון בכך שלעיתים קרובות טכנולוגיות בינה מלאכותית מעוררות סוגיות המטופלות על ידי כמה רגולטורים שונים, מה שעלול להקשות על הרגולטורים להגיב לטכנולוגיות באופן המצריך תיאום עם רגולטורים אחרים.²⁷¹

גם ועדת המשנה של המיזם הלאומי למערכות נבונות בנושא אתיקה ורגולציה של בינה מלאכותית המליצה, במילותיה: "להקים מנגנון תיאום פנים ממשלתי על-משרדי. כל זאת על מנת ליצור מדיניות אחידה, ברורה וקוהרנטית, בין כלל משרדי הממשלה".

המלצה דומה נכללה בטיוטת מסמך זה, אשר לאחר פרסומה להערות הציבור ובין היתר בעקבות זאת, הוחלט בהחלטת ממשלה 173 על הקמת מוקד הידע במשרד החדשנות, המדע והטכנולוגיה, ובשיתוף מחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים.

מעבר למשימות הייעוץ לרגולטורים והתיאום ביניהם, יש צורך בגורם שיהיה מופקד על יישום כלל המלצות המדיניות המפורטות במסמך זה, מעקב אחר ביצוען וגיבוש המלצות נוספות.

לאור האמור לעיל, מומלץ כי מוקד הידע ימלא, בין היתר, את התפקידים הבאים:

- א. **ילווה את גורמי הממשלה והרגולטורים וייעץ להם** בזיהוי ומיפוי השימושים בבינה מלאכותית בתחומם, הסיכונים הכרוכים בכך והמענים האפשריים, לרבות רגולציה. ובהמשך לכך, באפיון וקידום רגולציה כזו ברמה הענפית, ככל שהיא נחוצה.
- ב. **יקים ויוביל של פורום רגולטורים ופורום לשיתוף ציבור בתחום הבינה המלאכותית.** זאת, וצעדים נוספים, כדי לקדם תכנון, תיאום ושיתוף-פעולה בין רגולטורים בתחום הבינה המלאכותית, וכך לצמצם חפיפות וסתירות ברגולציה, ולהגביר ודאות ואחידות.
- ג. **יתמוך ביישום מדיניות הרגולציה והעקרונות האתיים**, ויבחן מעת לעת את הצורך בהתאמתם, בין היתר לאור המהלכים שיתגבשו במדינות מובילות ובארגונים בינלאומיים.
- ד. **ייעץ לממשלה ביחס לרגולציה בתחום הבינה המלאכותית ויבחן מעת לעת אם נדרשות התאמות במדיניות הרגולציה**, בפרט לעניין ההמלצה להימנע בשלב זה מרגולציה רוחבית. זאת, בין היתר לאור ההתפתחויות הטכנולוגיות, לרבות בתחום הבינה המלאכותית היוצרת על ענפיה, ומידת האפקטיביות של הרגולציה הענפית בהתמודדות עם האתגרים השונים.

ה. **יוביל את הייצוג והמעורבות של ישראל בפיתוח רגולציה ותקינה בפורומים בינלאומיים.**

- ו. **ינגיש מידע וכלים לשימוש אחראי בבינה מלאכותית עבור הרגולטורים והציבור**, ויוביל את העבודה לפיתוח מסגרת מומלצת (וולונטרית) לניהול סיכונים.

לנוכח תפקידיו המשמעותיים, מוקד הידע נדרש להיות מורכב מעובדי מדינה בעלי מומחיות מקצועית בתחומי מדיניות רגולציה וכן הבנה ומומחיות בתחומי הטכנולוגיה והמשפטים, כך שיוכל להיות בעל ראייה רוחבית ולהחזיק באחריות לקידום האינטרס של עידוד פיתוח ושימוש במערכות מבוססות בינה מלאכותית תוך שמירה על זכויות ואינטרסים, והשגת היעדים הממשלתיים.

²⁷¹ [C\(2021\)99/ADD1](#) 7-8; Practical Guidance on Agile Regulatory Governance to Harness Innovation
באופן כללי, לגבי אתגר התיאום בין רגולטורים, ראו יובל רויטמן "הרפורמה הרגולטורית: בין הגלוי לסמוי" משפט חברה ותרבות, מסדירים רגולציה – משפט ומדיניות 425, 429 (ישי בלנק, דוד לוי-פאור ורועי קרייטנר, עורכים, 2016).

כמו כן, על מוקד הידע להחזיק ביכולת המעשית, בהיבטים של תקנון ותקצוב, לתת מענה אפקטיבי לאתגרים השונים בתחומי הרגולציה בתחום הבינה המלאכותית – באופן שיבטיח מענה רגולטורי יעיל וקוהרנטי בתחומים בהם הוא נחוץ לקידום השוק ולהגנה על זכויות ואינטרסים.

קשה להפריז בחשיבות פעילותו של מוקד ידע מקצועי אשר יחזיק באמצעים המתאימים כדי לקדם את הטיפול בסוגיות הנדונות במסמך זה, בקצב שתואם להתפתחויות הטכנולוגיות והרגולטוריות, ונושא זה עלה שוב ושוב גם בהערות הציבור לטיוטת מסמך זה.

נוכח התפקידים המשמעותיים של מוקד הידע, ובמיוחד בשים לב להמלצה להימנע מחקיקה רוחבית בעניין בינה מלאכותית, יודגש כי מוקד הידע יוכל לבצע את מכלול התפקידים המוטלים עליו רק אם יועמדו לרשותו האמצעים התקציביים המתאימים לכך. ויובהר, בין היתר, תפקידיו של מוקד הידע כמצוין לעיל, יהיו פיתוח תורת רגולציה לבינה מלאכותית, פיתוח מודל לאיתור וניהול סיכונים, ליווי הרגולטורים הענפיים ותיאום ביניהם, מעקב ומעורבות בפיתוח החקיקה והתקינה הבין לאומית, קיום קשר מתמיד עם גורמי תעשייה, אקדמיה, מגזר שלישי וממשל והמשך מעקב אחר מדיניות הרגולציה וגיבוש המלצות לעדכונה. זאת, בסביבה בין לאומית וטכנולוגית שההתפתחויות בה מהירות ביותר, וכאשר השפעת תחום הבינה המלאכותית על הכלכלה היא דרמטית. על כן, קיימת חשיבות בל תתואר בהקצאת המשאבים המספקים למוקד הידע, כך שזה יוכל לבצע את מכלול התפקידים המוטלים עליו בצורה ראויה.

עוד מוצע, כי מוקד הידע יפעל במשרד החדשנות, המדע והטכנולוגיה, ותוקם עבורו ועדת היגוי שבה יכהנו נציגי הגופים הממשלתיים והציבוריים הרלוונטיים. זאת, על מנת לכוון את מוקד הידע בהתאם למכלול הזכויות והאינטרסים הנוגעים לעניין; להקנות לו ראיית רוחב על הנעשה בממשלה ובמשק הישראלי; לחזק את מומחיותו ויכולותיו; ולהעצים את השפעתו.

כן מוצע, שתפקידיו ועדת ההיגוי יהיו, בין היתר:

- א. להמליץ לשר החדשנות, המדע והטכנולוגיה על אסטרטגיה לאומית בתחום הבינה המלאכותית – כגוף בעל ראייה רחבה של הפעילות הממשלתית, ועדת ההיגוי היא גורם מתאים לגיבוש מסמך מדיניות לאומית בתחום הבינה המלאכותית, כפי שנעשה במדינות שונות, ועדכונו מעת לעת. מסמך המדיניות הלאומית יתייחס לתחומי הפעילות של מוקד הידע, אך לא יוגבל אליהם.
- ב. לבחון מעת לעת את הצורך בשינוי מדיניות הרגולציה בתחום הבינה המלאכותית – כמפורט לעיל, אחד מתפקידיו של מוקד הידע יהיה לבחון את הצורך בעדכון מדיניות הרגולציה כפי שהוצעה במסמך זה. ראוי כי נושא משמעותי זה יובל על ידי ועדת ההיגוי אשר בה חברים שחקני מפתח מגוונים מהגופים הציבוריים הרלוונטיים.
- ג. לייעץ לעניין תכנית העבודה השנתית של מוקד הידע – לנוכח חשיבות פעילותו של מוקד הידע, וריבוי תחומי המומחיות והאינטרסים המשולבים בפעילותו, יש מקום לשלב את כלל הגורמים המיוצגים בוועדת ההיגוי בהשפעה על תכנית העבודה השנתית שלו, ובהערכת הצלחתה. מעורבות כזו תוכל לתרום להצלחת פעילות מוקד הידע ולהשגת יעדיו.
- ד. להמליץ על יעדי מוקד הידע לטווח הזמן הבינוני – לנוכח היתרונות המוסדיים של ועדת ההיגוי, כמפורט לעיל, יש מקום להסתייע בה לשם הצבת יעדים רב-שנתיים בפני מוקד הידע מתוך ראייה אסטרטגית, אשר חורגת מעבר לפעילות השוטפת.

לאור ייעודה ותפקידה, מוצע כי הרכב ועדת ההיגוי יכול את הנציגים הבאים (או מי מטעמם):

- א. המנהל הכללי של משרד החדשנות, המדע והטכנולוגיה – יו"ר הוועדה.
- ב. מנהל מוקד הידע.
- ג. המשנה ליועמ"ש לממשלה (משפט כלכלי).
- ד. הממונה על התקציבים במשרד האוצר.
- ה. ראש רשות הרגולציה.
- ו. ראש מערך הדיגיטל הלאומי.
- ז. ראש רשות הגנת הפרטיות.
- ח. המנהל הכללי של רשות החדשנות.

6.4. מיפוי השימושים בבינה מלאכותית והאתגרים הנלווים להם בענפים מוסדרים

מוצע כי גורמי הממשלה והרגולטורים הרלוונטיים יפעלו בהקדם לזיהוי ומיפוי של השימושים הקונקרטיים במערכות מבוססות בינה מלאכותית הנעשים על ידי המפוקחים או השחקנים בענף המוסדר שעליו הם אמונים; האתגרים, החששות והסיכונים הכרוכים בכך; והמענים האפשריים. מוקד הידע והתיאום הממשלתי יסייע לגורמי הממשלה והרגולטורים במלאכה זו וירכז אותה.

כאמור, בשנים האחרונות חל גידול משמעותי בהיקף השימושים במערכות מבוססות בינה מלאכותית והן מוטמעות כמעט בכל ענף במשק, כשההתפתחויות ביחס לבינה המלאכותית היוצרת מאיצות תהליך זה. חלק ניכר מהענפים שבהם נעשה שימוש במערכות מבוססות בינה מלאכותית מוסדר באמצעות רגולציה ענפית (כגון רפואה ובנקאות) או רוחבית (כגון תחרות ופרטיות) וכפוף לרגולטור ענפי (כדוגמת משרד הבריאות והפיקוח על הבנקים בבנק ישראל) או רוחבי (כדוגמת רשות התחרות והרשות להגנת הפרטיות). לעיתים הרגולציה הנוהגת כוללת אף בחינה מראש של הרגולטור הרלוונטי או דרישה לקבלת רישיון מהרגולטור כתנאי לפעילות.

עם זאת, בשיח ראשוני שקיים הצוות הבין-משרדי שעסק ברגולציה ואתיקה במהלך גיבוש התכנית הלאומית לבינה מלאכותית עם רגולטורים שונים, עלה כי יש מקום לעריכת מיפוי של שימושים משמעותיים במערכות מבוססות בינה מלאכותית על ידי גורמים מפקחים או שחקנים שאופי פעילותם לא בהכרח הכפיף אותם עד כה לפיקוח, בתחום שעליו הם אמונים ובהתאם לסמכותם לפי דין. בהמשך לכך, יהיה צורך לבחון את הסוגיות, החששות והסיכונים המתעוררים לגבי שימושים אלה, בין היתר לאור האתגרים שנמנו במסמך זה (אפליה, מעורבות אנושית, הסברתיות, גילוי, אמינות עמידות, אבטחה ובטיחות, אחריות ופרטיות), ואת המענים האפשריים.

מומלץ כי מוקד הידע והתיאום הממשלתי ילווה את מלאכת המיפוי והבחינה שיבצעו הרגולטורים הענפיים; יסייע לגורמי הממשלה והרגולטורים העוסקים בה; ויקדם אותה באופן אקטיבי. כך, בין היתר, מוקד הידע יוכל לסייע לרגולטורים למקד את השאלות ולייצר מהממצאים תובנות משמעותיות. בנוסף, על רקע ממצאי המיפוי יוכל מוקד הידע והתיאום הממשלתי לייצר תמונת מצב רחבה עבור הממשלה, על אודות היקף השימוש בבינה מלאכותית ואופן ההתמודדות עם סיכונים לזכויות יסוד ואינטרסים ציבוריים, כמו גם על חסמים לקידום בינה מלאכותית.

על מנת שיוכל לבצע תפקידים אלה, מוצע כי הרגולטורים וגופי הממשלה יעדכנו את מוקד הידע בנוגע לרגולציה העוסקת בתחום הבינה מלאכותית שהם מקדמים, וכי ימסרו למוקד הידע מעת לעת סטאטוס עדכני בנוגע לנעשה בתחומם, ועל ההתפתחויות הטכנולוגיות והרגולטוריות.

יודגש כי אין בהצעה זו כדי לקרוא להתערבות רגולטורית מידית בפועל שתשפיע על היכולת לפתח או להשתמש בבינה מלאכותית כשהדבר אינו נדרש, ומטרתה לסייע בהבנת תמונת המצב במבט צופה פני עתיד, ובגיבוש מענים רגולטוריים במידת הצורך ובהתאם לסיכונים קונקרטיים.

6.5. הקמת פורום רגולטורים ופורום לשיתוף ציבור לתחום הבינה המלאכותית

מוצע להקים פורום מקצועי פנים-ממשלתי הכולל נציגי רגולטורים ומומחים לטכנולוגיה, מדיניות ומשפט, על מנת לקדם את הרגולציה הענפית בתחום הבינה המלאכותית באופן מתואם וקוהרנטי, ותוך שיתוף פעולה ולמידה הדדית. בנוסף, מוצע להקים פורום לשיתוף ציבור הכולל נציגי תעשייה, אקדמיה, ארגוני חברה אזרחית והציבור הרחב, כדי לדון בסוגיות של הסדרת בינה מלאכותית.

בחלק הראשון, מוצע לקבוע "שולחנות עגולים" שיתקיימו באופן עיתי או לפי הצורך בין רגולטורים וגורמי ממשלה רלוונטיים, בדגש על תחומים עתירי "חיכוך", במטרה לדון בנושאים שעל סדר היום ולבחון את הצורך בעדכון מדיניות הרגולציה לתחום הבינה המלאכותית ממבט מאקרו. במסגרת זאת, ניתן יהיה לקיים למידת עמיתים בנוגע להתפתחויות הטכנולוגיות בתחום, לגישות חדשניות לרגולציה ולתובנות מהשטח. בדרך זו, הפורום יסייע באחידות וביישום מדיניות רגולציה ואתיקה זו. בנוסף, ההתפתחויות המהירות בתחום הטכנולוגי מחייבות בחינה עיתית של המלצות המדיניות, ולכן מוצע שמטרה נוספת של הפורום האמור תהיה לדון בהיבטים אלה באופן קבוע.

בחלק השני, מוצע לקבוע פורום נוסף לשיתוף ציבור, עם נציגי תעשייה, אקדמיה וחברה אזרחית, וכן בעלי עניין חוץ-ממשלתיים רלוונטיים נוספים, במטרה לאתר את התחומים העיקריים שבהם יש צורך בהתאמה של הרגולציה או מדיניות הרגולציה בתחום הבינה המלאכותית, כדי להסיר חסמים לפעילות כלכלית או חברתית רציפה ויעילה או לגבש מענים מותאמים לצרכים.

תחום הבינה המלאכותית מובל בידי התעשייה והאקדמיה ולכן יש חשיבות לעדכון ושיח רציף בינם לבין גורמי הממשלה. במסגרת פורום זה, ניתן יהיה לדון בהתפתחויות הטכנולוגיות, בשאלה אם כלי רגולציה מסוימים נותנים מענה לאתגרים השונים כמו גם לעקרונות שפורטו לעיל ובשאלות נוספות. כמו כן, ניתן יהיה לקיים שמיעת ציבור רציפה שתסייע לממשלה ולרגולטורים השונים לגבש מדיניות ולבחון אותה מעת לעת. יתרון נוסף של פורום זה ושל קיום שיח פתוח ולומד, הוא בהיותו מגביר שקיפות וככזה מסייע אף להגברת אמון הציבור והוודאות הרגולטורית.

יצוין כי מיסוד פורום קבוע תואם גם את המלצות ה-OECD לגבי רגולציה גמישה, וכי פורמט שולחנות עגולים הוצע ככלי כללי לתיאום בין רגולטורים על ידי צוות רגולציה חכמה.²⁷²

מומלץ להפקיד את מלאכת הריכוז של פורומים אלה בידי מוקד הידע והתיאום הממשלתי שיוקם.

6.6. תיאום המעורבות בפיתוח הרגולציה והתקינה בפורומים בינלאומיים

מוצע כי גורמי הממשלה הפועלים במסגרת פורומים וארגונים בינלאומיים, או עומדים בקשרי עבודה עם מדינות מפותחות אחרות בתחום זה יפעלו יחד לקידום מדיניות רגולציה ואתיקה מאוזנת ותואמת למדיניות רגולציה זו וליעדים הממשלתיים, בהובלת מוקד הידע. זאת, בשים לב להחלטת ממשלה 212 והחלטת ממשלה 173, שהטילו על שר החדשנות, המדע והטכנולוגיה להוביל את שיתופי הפעולה הבינלאומיים האזרחיים בתחום הבינה המלאכותית, לייצג את ישראל בפורומים בינלאומיים אזרחיים, ולפעול לקידום שיתופי פעולה בינלאומיים אזרחיים בתחומי הבינה

²⁷² דוח צוות הבין-משרדי לרגולציה חכמה, לעיל ה"ש 245, בעמ' 118.

המלאכותיות בשיתוף עם עוסקים בכך ברשות החדשנות, במשרד החוץ, במחלקת ייעוץ וחקיקה (משפט בינלאומי) ומחלקת ייעוץ וחקיקה (משפט כלכלי) ובגופים רלוונטיים נוספים.

ארגונים בינלאומיים ומדינות רבות בעולם עוסקים במדיניות ביחס לבינה מלאכותית. בין הארגונים נמנים לדוגמה ה-OECD, UNESCO, CAI (ועדה של מועצת אירופה), GPAI (Global Partnership on Artificial Intelligence) – ארגון המאגד נציגי ממשלות, מדענים, חברות טכנולוגיה, ארגוני חברה אזרחית וארגונים בינ"ל, שמטרתו לקדם את שיתוף הפעולה הבינלאומי בנושא ולהוות מרכז ידע לסוגיות מובילות של בינה מלאכותית, במטרה לקדם מערכות בינה מלאכותית אמינות²⁷³ והאיחוד האירופי.²⁷⁴ עם המדינות המובילות נמנות למשל ארה"ב, בריטניה, קנדה, יפן וסינגפור. בנוסף, נעשית עבודה נמרצת בארגוני תקינה בינלאומית כמו ISO ו-IEEE לפיתוח תקנים ל"בינה מלאכותית אחראית".²⁷⁵ אלה, ורבים אחרים, עוסקים בהסדרת בינה מלאכותית, ויש להניח כי לתוצרי העבודה שלהם תהיה השפעה גלובלית, כולל על הנעשה בישראל. בנוסף, ישראל משתתפת בחלק מהפורומים המערבים ארגונים בינלאומיים ומדינות אלה ומקיימת שיח עם.

על רקע זה, ובשים לב להחלטת ממשלה 212 המצוטטת לעיל, מוצע כי גורמי הממשלה הרלוונטיים ימשיכו לפעול באופן אקטיבי במסגרת פורומים בינלאומיים שונים לקידום מדיניות רגולציה ואתיקה מאוזנת ומתואמת בתחום הבינה המלאכותית לאור האינטרסים הישראליים, בהתאם למדיניות הרגולציה המוצעת במסמך זה וליעדים הממשלתיים; כאשר מוקד הידע הממשלתי יוכל להיות הפלטפורמה לסנכרון כלל הפעילויות הממשלתיות בהקשר זה. בד בבד, מוצע כי ההשתתפות בפורומים הבינלאומיים השונים תשמש גם ככלי מרכזי ללמידה מניסיוןן של המדינות השונות, להעמקת הידע בישראל בתחומים אלה, ולהבנת המגמות הבינלאומיות המסתמנות.

6.7. הנגשת מידע וכלים לשימוש אחראי בבינה מלאכותית, ובכלל זה כלי לניהול סיכונים

מוצע להנגיש מידע וכלים לשימוש אחראי בבינה מלאכותית, ובכלל זה לשקול פיתוח או אימוץ של כלי אחיד לניהול סיכונים ביחס לשימוש בבינה מלאכותית שיאפשר שפה משותפת בין גורמי הממשלה והרגולטורים ובינם לבין גורמים פרטיים. השפה האחידה תסייע למגזר הפרטי להעריך את הסיכונים הנלווים לשימוש מסוים בבינה מלאכותית, וכן לרגולטורים לבחון את הסיכונים הכרוכים בשימוש בה בתחום שעליו הם אמונים ואת הצורך בהתערבות.

באופן כללי, תפיסות ניהול הסיכונים נועדו למנוע סיכונים לזכויות יסוד ואינטרסים ציבוריים. האתגרים שנסקרו לעיל והעקרונות האתיים שהוצעו לעיל, מסייעים בהכוונת ארגונים לגבי סוגי הסיכונים הנפוצים ב"מחזור החיים" של הפיתוח והשימוש בבינה מלאכותית.

תפיסות ניהול סיכונים מהוות מכנה משותף למדיניות בינה מלאכותית במדינות המפותחות, ונכללות בין היתר בבסיס המלצות ה-OECD משנת 2019, טיוטת החקיקה האירופית, עקרונות תפיסת המדיניות האמריקאית, טיוטת האמנה המקודמת בוועדת CAI של מועצת אירופה ועוד.

²⁷³ ראו <https://gpai.ai/>.

²⁷⁴ ארגון המאגד נציגי ממשלות, מדענים, חברות טכנולוגיה, ארגוני חברה אזרחית וארגונים בינ"ל, שמטרתו לקדם את שיתוף הפעולה הבינ"ל בנושא ולהוות מרכז ידע לסוגיות מובילות של בינה מלאכותית, במטרה לקדם מערכות בינה מלאכותית אמינות. ראו <https://gpai.ai/>.

²⁷⁵ לסקירה על התקנה בארגון ISO ובגופי התקנה האירופיים ראו Stefano Nativi & Sarah de Nigris, *AI Watch: EC Proposal for an AI Standardization Landscape State of Play and Link to the AI Regulatory Framework*, EUROPEAN COMMISSION (2021); לסקירה על קידום התקנה בארה"ב ראו NIST (2023) *AI Risk Management Framework*, <https://www.nist.gov/itl/ai-risk-management-framework>.

תפיסת ניהול הסיכונים בבינה מלאכותית נשענת על תפיסות ניהול סיכונים מקובלות (שנועדו להתמודד עם אי ודאות שנוצרת כתוצאה מפעילות כלשהי),²⁷⁶ על תפיסות ניהול סיכונים לניהול טכנולוגיות המידע,²⁷⁷ ועל ההקשר המשפטי והחברתי הקונקרטי בו מופעלת הבינה המלאכותית.

המחקר המלווה להמלצות ה-OECD מסכם באופן כללי את עקרונות-העל של ניהול סיכונים הנדרש בכל "מחזור החיים" של בינה מלאכותית:²⁷⁸ (1) הגדרת יעדים, פונקציות או תכונות של מערכת הבינה המלאכותית; (2) זיהוי בעלי עניין וגורמים מעורבים; (3) הערכה של ההשפעות הפוטנציאליות, הן התועלות והן הסיכונים לבעלי העניין והגורמים המעורבים; (4) התאמה של אסטרטגיות צמצום סיכונים לסוג הסיכונים; (5) יישום האמצעים לצמצום סיכונים; (6) בקרה, הערכה ומשוב.²⁷⁹

חלק משמעותי בניהול הסיכונים הוא קיום תהליכים סדורים ומתועדים, המאפשרים לבחון את אופן יישום העקרונות בסיטואציה שבה נעשה שימוש במערכות מבוססות בינה מלאכותית. בנוסף, תהליכים אלה מאפשרים להתחקות אחר תהליך הפיתוח והשימוש, ולבקר היבטים אלה במידת הצורך כדי להפיק לקחים, לתקן החלטות שגויות, ולשפר את הליך קבלת ההחלטות, באופן רציף. עקב כך הם מדגישים תיעוד, עבודה בהתאם לכללים מקובלים בתחום וקיום נהלים.

עקרונות אלה הם כלליים ומהווים מכנה משותף רחב בתחום ניהול הסיכונים. לצד זאת ניתן לראות במדינות העולם ובגופי התקינה תפיסות המרחיבות היבטים אלה. כך לדוגמה, טיוטת החקיקה האירופית כוללת שילוב של דרישות ניהול סיכונים שמוצע לקבוע בחוק, והשלמה של דרישות אלה באמצעות תקינה מתאימה. סעיף 9 לטיוטת החקיקה של האיחוד האירופי כולל פירוט של הרכיבים הנדרשים למנגנון ניהול סיכונים עבור מערכות בינה מלאכותית המסווגות ב"סיכון גבוה".²⁸⁰ לצד זאת, מוצע בטיטוט החקיקה, לגבי שימושים מסוימים, כי ארגון שעומד בתקינה שנקבעה בהתאם לחקיקה יהיה בעל "חזקת" עמידה בהוראות החוק.²⁸¹ כך שהנחת העבודה היא שהתקינה תהיה

²⁷⁶ המחלקה העוסקת בבינה מלאכותית ב-OECD מבססת את עבודתה על תפיסות ניהול הסיכונים המקובלות בסדרת תקני ISO בנושא ניהול סיכונים:

ISO, ISO 31000, *Risk management – Guidelines*, <https://www.iso.org/iso-31000-risk-management.html>
"ISO 31000, Risk management – Guidelines, provides principles, a framework and a process for managing risk. It can be used by any organization regardless of its size, activity or sector. Using ISO 31000 can help organizations increase the likelihood of achieving objectives, improve the identification of opportunities and threats and effectively allocate and use resources for risk treatment. However, ISO 31000 cannot be used for certification purposes, but does provide guidance for internal or external audit programmes. Organizations using it can compare their risk management practices with an internationally recognized benchmark, providing sound principles for effective management and corporate governance."

ISO 31000 - Risk management, <https://www.iso.org/publication/PUB100426.html>
"ISO 31000 provides direction on how companies can integrate risk-based decision making into an It is an open, . reporting, policies, values and culture, organization's governance, planning, management meaning it enables organizations to apply the principles in the standard to, principles-based system ; the organizational context."

התקן אומץ כתקן ישראלי: מכון התקנים הישראלי, ת"י 31000 חלק 1, ניהול סיכונים: הנחיות למימוש התקן הישראלי ת"י 31000-ניהול סיכונים-עקרונות וקווים מנחים <https://www.sii.org.il/he/>. תקנים אלה שימשו את OECD גם בגיבוש המלצות לניהול סיכונים בתחום הגנת הסייבר, OECD, *Digital Security Risk Management* (2015) <https://www.oecd.org/sti/ieconomy/digital-security-risk-management.htm>

²⁷⁷ ראו למשל את סדרת תקני ISO בתחום ניהול סיכונים סייבר: ISO/IEC 27001 Information security management, <https://www.iso.org/standard/27001>. הדגש של תקנים אלה הוא במערכת ניהול טכנולוגית המידע כדי לנהל את סיכוני האבטחה.

התקן אומץ כתקן ישראלי: מכון התקנים הישראלי, ת"י - ISO/IEC 27001 מערכת ניהול אבטחת מידע <https://www.sii.org.il/he/iso-27001>
²⁷⁸ AI in Society, לעיל ה"ש 7, בעמ' 96-95.

²⁷⁹ שם.
²⁸⁰ טיוטת הרגולציה האירופית, לעיל ה"ש 248, ס' 9 "Risk Management System".
²⁸¹ שם, בס' 42 "Presumption of conformity with certain requirements".

מפורטת יותר, ותאפשר מצד אחד ודאות לגבי אופן עמידה בה, ומצד שני לארגונים שיבחרו לעמוד בדרישות ניהול הסיכונים החקוקות בדרך אחרת, לעשות זאת.²⁸² טיוטת החקיקה כוללת גם דרישה לביצוע "בדיקת התאמה" ("conformity assessment") לחקיקה למערכות ב"סיכון גבוה".²⁸³

תפיסות ניהול סיכונים באות לידי ביטוי בשני מובנים מרכזיים. מובן ראשון עוסק **בנקודת המבט של הרגולטור**. תפיסות ניהול סיכונים משמשות גורמי ממשלה ורגולטורים לגיבוש רגולציה ומדיניות רגולציה ביחס לבינה מלאכותית. מובן זה מקושר גם לחוק עקרונות האסדרה, בכך שהוא משמש כלי עבודה לבחינת ההצדקה להתערבות בפעילות המגזר הפרטי. כך, בהתאם לתפיסת ניהול סיכונים, ההצדקה לקביעת כללי התנהגות לפיתוח ושימוש בבינה מלאכותית גוברת מקום שבו הסיכון לזכויות יסוד או אינטרס ציבורי הוא גבוה. בהתאמה, ככל שהסיכון נמוך יותר, ההצדקה להתערבות מצטמצמת, וכללי ההתנהגות מהווים רק המלצה לאופן ההתמודדות הארגוני עם הסיכון. מובן שני עוסק בנקודת המבט של גורמים פרטיים המפתחים או משתמשים בבינה מלאכותית. מובן זה עוסק בתפיסות מקובלות לניהול סיכונים (כמפורט להלן), בהתאם לעקרונות האתיים, כדי לאתר את הסיכונים הנובעים מבינה מלאכותית לזכויות יסוד ואינטרסים ציבוריים, ולנקוט באמצעים טכנולוגיים ונהליים לצמצם אותם לרמה מקובלת. ביצוע האמור בצורה שיטתית גם מאפשר להפיק מידע בצורה אחידה על הבינה המלאכותית ואופן ניהול הסיכונים שלה. מידע זה רלוונטי לגורמים נוספים הקשורים בשימוש בבינה המלאכותית, כגון שותפים עסקיים או לקוחות, הנדרשים לקבל החלטות ולנהל את הסיכונים הקשורים בפעילות שלהם בנושא זה.

הממשק בין שני מובנים אלה, הוא שתפיסת ניהול הסיכונים מסייעת להפעלת סמכות הרגולטור בין היתר כדי לבחון האם ההתמודדות הארגונית הוולונטרית בהתאם ליישום תפיסת ניהול הסיכונים על ידו מספקת מענה הולם לסיכון לפגיעה בזכויות יסוד או באינטרסים ציבוריים. בנוסף, תפיסת ניהול הסיכונים מהווה שפה משותפת לשיח בין הרגולטור לבין הארגון.

כמצוין לעיל, יש תועלת בקביעת שפה משותפת וכלי עבודה דומים לצורך קידום המדיניות הרגולטורית באופן אחיד, ויישום העקרונות על ידי גורמים פרטיים בענפי המשק השונים. כלי אחד ליצירת השפה המשותפת הוא העקרונות האתיים המוצעים לעיל (לפיו למשל הרגולטור והמפוקח מכירים שניהם את החשש לאפליה ויודעים מהי הסברתיות). כלי נוסף להתוויית שפה משותפת, שמוצע כאן, הוא תפיסת ניהול הסיכונים (שלפיה הם יוכלו למשל לבחון באופן דומה מה נחשב לשימוש בסיכון גבוה או נמוך, ובאיזה סוג של שימוש ניכרת שכיחות גבוהה של סיכונים).

לקיומה של שפה משותפת גם במסגרת ניהול הסיכונים מספר יתרונות, ובכלל זה, קידום שפה אחידה להתמודדות עם טכנולוגיה מסוימת; קידום תפיסה קוהרנטית לגבי פעילות שחוצה תחומי סמכות רגולטוריים; שמירה על זיקה בין תפיסות ניהול סיכונים בתחום הבינה המלאכותית להוראות אחרות בתחום ניהול הסיכונים, כדי להקל על ארגונים בנטל ניהול הסיכונים החל עליהם; אפשרות פיתוח תחום בודקי הסיכונים באופן אחיד; סיוע לרגולטורים בבחינת התחומים הדורשים

²⁸² לפרשנויות על אודות תהליכי העבודה המפורטים הנדרשים כדי לעמוד בחקיקה האירופית ראו Luciano Floridi, Matthias Holweg, Mariarosaria Taddeo, Javier Amaya Silva, Jakob Mökander & Yuni Wen, *CapAI - A Procedure for Conducting Conformity Assessment of AI Systems in Line With the EU Artificial Intelligence Act* SSRN (Mar. 23, 2022), <https://papers.ssrn.com/> or <https://papers.ssrn.com/>.
²⁸³ טיוטת הרגולציה האירופית, לעיל ה"ש 248, בס' 43; ליחס שבין "בדיקת התאמה" לתפיסה המוכרת בחקיקת הפרטיות של "תסקיר השפעה על הפרטיות", ראו Katerina Demetrou, *Introduction to the Conformity Assessment Under the Draft EU AI Act, and How it Compares to DPLAS, FUTURE OF PRIVACY FORUM* (Aug. 12, 2022), <https://fpf.org/blog/i>.

את התערבותם; והתאמה לכללים המתקדמים בשווקים של המדינות המפותחות. שפה משותפת תאפשר גם לארגון להעביר את המידע הנדרש לרגולטור לצורך בחינת הצורך בהתערבות.

כפי שתואר לעיל, טיוטת הרגולציה של האיחוד האירופי קובעת במפורש את רמות הסיכון השונות. היא מבחינה בין מערכות ב"סיכון בלתי מקובל", "סיכון גבוה", "סיכון מוגבל", ו"סיכון מינימלי", ומגדירה את התהליך לסיווג המערכות לתוך קטגוריות אלה. גם ה-OECD פעל לייצר תבחינים לסיווג מערכות בינה מלאכותית, אשר נועדו למפות את הרכיבים השונים במערכות בינה מלאכותית שעשויים להצביע על מידת רגישותן. ואולם, בשונה מהשימוש בתבחינים באיחוד האירופי, אלה של ה-OECD אינם כוללים הכרעה ערכית לגבי מהי רמת סיכון גבוהה.

על רקע זה, מוצע כי במסגרת תפיסה אחידה לניהול סיכונים, תישקל קביעת שיטה הכוללת תבחינים אחידים למיון מערכות בינה מלאכותית. עם זאת, מוצע כי הקביעה הנורמטיבית לגבי שילוב התבחינים העוברים את הסף שעשוי להצדיק התערבות רגולטורית, תיעשה על ידי כל רגולטור ברמה הענפית ובהתאם לסביבה המשפטית והרגולטורית החלה. בהקשר זה יודגש כי מוקד הבחינה אינו בהכרח הטכנולוגיה ככזו, אלא בעיקר סוג השימוש בה (כלומר, בחינה של הרובד האפליקטיבי והסדרה ברמת האפליקציה). זאת, משום שמדובר בטכנולוגיה הניתנת ליישום בהקשרים שונים.

גופי התקינה בעולם כבר עמלים על תקנים הנוגעים למערכות בינה מלאכותית והתאמתם לנורמות המתהוות בשיח הבינלאומי – כך ב-ISC/IEC,²⁸⁴ וכך גם בגוף התקינה בארה"ב NIST.²⁸⁵ יצוין כי גם במסגרת חוק הבינה המלאכותית האירופי עתידים להתפרסם כלים דומים על ידי גוף התקינה האירופי CEN/CENELEC,²⁸⁶ וכי כלי נוסף לניהול סיכונים מפותח לצורך אמנת מועצת אירופה.²⁸⁷ בתחומים אחרים של ניהול סיכונים טכנולוגיים מידע, כגון בהגנת סייבר ובהגנה על הפרטיות, כבר התפתחו תקנים ושיטות עבודה שמטרתם ניהול שיטתי ורציף של הסיכון ומניעת התממשותו.

לאור האמור לעיל, מוצע לשקול פיתוח או אימוץ של כלי אחיד לניהול סיכונים שיהיה נגיש בשפה העברית וביחס לנעשה בישראל, וישמש הן את גורמי הממשלה והרגולטורים והן גורמים פרטיים, לשם הערכת הסיכון ביחס לשימושים מסוימים במערכות מבוססות בינה מלאכותית.

מומלץ כי משימה זו תבוצע בהובלת מוקד הידע והתיאום הממשלתי, וכי מעבר לכלי לניהול סיכונים, מוקד הידע יעסוק בפיתוח והנגשה של כלים נוספים עבור גורמי הממשלה והרגולטורים, וכן עבור גורמים פרטיים – על מנת לקדם פיתוח ושימוש אחראי בבינה מלאכותית.

יצוין כי במסגרת הליך שיתוף הציבור לטיטוט מסמך זה, התקבלו מספר הערות המתייחסות לשלבים מתקדמים יותר בגיבוש הכלי לניהול סיכונים, וביניהן: הצורך לבחון את הסיכונים בכל שלב בדרך שבה מערכת בינה מלאכותית מפותחת ומוטמעת, ובכלל זה בשלב יצירת מאגר הנתונים ללמידת המערכת, בשלבי הלימוד והאימון, ובשלב הטמעת המערכת במוצר הלקוח; הרעיון ליצור

ISO/IEC 23894:2023, Information Technology — Artificial Intelligence — Guidance on Risk Management ²⁸⁴ <https://www.iso.org/standard/77304.html>

NIST AI 100-1, Artificial Intelligence Risk Management Framework (AI RMF 1.0) ²⁸⁵ <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>

Ensuring a Harmonised AI and ICT Professionalism Standardization: Read the New CEN White Paper, CEN-CENELEC (Jun. 30, 2023), <https://www.cenelec.eu/> ²⁸⁶

Committee On Artificial Intelligence (CAI): Outline of Huderia Risk and Impact Assessment ²⁸⁷ Methodology, COUNCIL OF EUROPE (2022) <https://rm.coe.int/>

סרגל "תועלות" לצד סרגל ה"סיכונים" שייבחן במקביל על מנת לאזן בין הסיכון והתועלת; האפשרות להטיל חובה על חלק מהגורמים הפרטיים העוסקים בפיתוח ושימוש בבינה מלאכותית לנהל תהליך ניהול סיכונים פרטני ושקוף ועוד. בהמשך העבודה על גיבוש הכלי לניהול סיכונים, שלפי ההמלצה כאמור תבוצע בהובלת מוקד הידע הממשלתי, ראוי יהיה לבחון גם הערות אלה.

הערות הציבור

טיוטה של מסמך זה הופצה להערות הציבור ביום 30.10.2022, ובעקבותיה התקיים תהליך שימוע ציבורי נרחב. במסגרת זאת התקבלו התייחסויות רבות ותגובות בכתב מאת גורמים שונים מהמגזר הציבורי הרחב, ומאת נציגי אקדמיה, ארגוני חברה אזרחית ותעשייה. בנוסף, טיוטת המסמך עמדה במוקד יום עיון במתכונת "שולחנות עגולים" בארגון משרד החדשנות, המדע והטכנולוגיה וייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים, וכן במוקד יום עיון נוסף שארגן מרכז הנשיא מאיר שמגר למשפט דיגיטלי וחדשנות באוניברסיטת תל אביב בהשתתפות המעורבים בגיבוש המסמך. בהמשך אף נערכו פגישות עם מספר גופים ששלחו הערות בכתב וחפצו בשיח משלים בעל-פה.

ההערות הציבורי שהתקבלו הציפו את מורכבות הנושאים בהם עוסק המסמך, ואת המתח בין הצורך בקביעת רגולציה בתחום הבינה המלאכותית לבין החשש מפני רגולציה עודפת שתפגע בחדשנות הטכנולוגית. מרבית ההערות נגעו להמלצות המפורטות במסמך, ולהלן עיקריהן:

ביחס להמלצה להימנע בשלב זה מחקיקה רוחבית ולפעול להסדרת תחום הבינה המלאכותית באמצעות הרגולטוריים הענפיים, הערות רבות התייחסו לחשש מחפיפה בין הרגולטורים; מאי-בהירות ומחוסר קוהרנטיות בין הדרישות בתחומים שונים, אפילו ברמת המינוחים וההגדרות; ומחוסר יכולתם של הרגולטורים הענפיים, בהיעדר ידע, כלים ותמריצים מתאימים, לקבוע רגולציה עילה בתחומים בהם היא נדרשת. עוד נטען, כי בכל אופן ההתפתחויות בתחום הבינה המלאכותית מחייבות תיקוני חקיקה בתחומים כגון זכויות יוצרים והגנת הפרטיות.

ביחס למדיניות הרגולציה המוצעת, היו שהעירו שהסתפקות בעקרונות אתיים שאינם בעלי מעמד מחייב תאפשר לגופים שלא לפעול לפי שיקולים אתיים מטעמי רווח, תוך עיכוב הדיון על רגולציה אפקטיבית ("טיוח אתי"); וכן שרגולציה רכה עלולה להיפך ל"עלה תאנה". עוד נטען, כי לכל הפחות ביחס לשימושים מסוימים נדרשים איסורים קטגוריים בחקיקה, ובהיעדרה מדינת ישראל עלולה לפגור אחר הנעשה בעולם ולהפוך ל"חצר אחורית" בהגנה על זכויות יסוד ואינטרסים ציבוריים. מנגד, היו מעירים שחששו שאימוץ רגולציה בישראל בחקיקה ברוח הטיוטה האירופית עלול לפגוע בתעשייה ובחדשנות, וטענו כי ניתן להסתפק בכללי המשפט הקיימים ללא רגולציה ייעודית.

ביחס לעקרונות האתיים שמוצע לאמץ צוין, בין היתר, שהעיקרון "בינה מלאכותית תשמש לקידום צמיחה, פיתוח בר קיימא והמובילות הישראלית בתחום החדשנות" הוא עקרון מנחה ולא עיקרון "אתי" של ממש; שהעיקרון "האדם במרכז: כיבוד זכויות יסוד ואינטרסים ציבוריים" אינו מבטא רק הכרח אשר עומד במתח עם הרצון לעודד חדשנות, אלא גם אמצעי לעודד צמיחה, משום שההגנה על זכויות היסוד מבטיחה ודאות לשחקנים השונים בשוק ומונעת פיתוח יישומי בינה מלאכותית מזיקים. כמו כן הוצע להתייחס לזכויות אדם נוספות כגון חופש הביטוי, ולאינטרסים כגון הגנה על הדמוקרטיה; שהעיקרון "שוויון ומניעת אפליה פסולה" מעורר קשיים שונים, ובהם הקושי להגדיר מהי בדיוק אפליה פסולה תוך הבחנה בינה ובין התאמה אישית של שירותים, והקושי שעשוי לעמוד בפני התעשייה כאשר תתנסה במבחני שוויון תוצאתיים (שלא תמיד ברור מי יבצע ומתי) או הליכיים (כגון הקושי להשיג נתונים שיכילו ייצוג אחיד לכל קבוצות האוכלוסיה); שהעיקרון "שקיפות והסבריות" קשה להגדרה לאור הפערים בין אופן החשיבה האנושי ואופן הפעולה של מערכות בינה מלאכותית; שהעיקרון "אמינות, עמידות אבטחה ובטיחות" עשוי להקשות על פעילות מסחרית בשל היבטים הנוגעים לקניין רוחני, וכן עלתה הצעה לבצע מעקב רוחבי אחר תקלות

ואירועים בעייתיים; ושהעיקרון "אחריותיות" מעורר קושי בזיהוי הגורם האחראי לנזק במקרים בהם מעורבים בפיתוח ובשימוש במערכת בינה מלאכותית שורת גורמים נפרדים אשר היחסים ביניהם ובינם לבין הנזק מורכבים, וכן מציף את החשש שמשטר אחריות מחמיר ישמש כגורם מצנן בפני פיתוחים חדשניים.

לבסוף, ביחס להמלצה להקים מוקד ידע ממשלתי, רבים עמדו על היתרון במרכז מידע ותיאום ממשלתי שיפתח מומחיות מקצועית וטכנולוגית, בפרט לנוכח העדפת רגולציה ענפית על פני רגולציה רוחבית להסדרת תחום הבינה המלאכותית בשלב זה, אך התקבלו גם הערות שהביעו חשש מכך שהשפעתו, במבנה ובאופן שתואר בטיטוט המסמך שהופצה לציבור, תהיה מוגבלת למדי, בין היתר בהיותו נעדר מעמד סטטוטורי ויכולת להנחות את הרגולטורים הענפיים.

הערות נוספות נגעו לנושאים שאינם מטופלים במסמך זה אלא במסגרות אחרות, או שייעשו רלוונטיים בעתיד. כך ביחס לשימוש בכלי בינה מלאכותית על ידי רשויות ציבוריות – בין היתר, עלה הצורך לתמרץ את המגזר הציבורי להשתמש בכלי בינה מלאכותית, באופן שיועיל למשק ויקדם בהן את אמון הציבור; וכך ביחס לתוכן ההסדרים הרגולטוריים, נושא אשר לפי ההמלצות יידון בשלב מאוחר יותר על ידי הרגולטורים הענפיים תוך סיוע וייעוץ של מוקד הידע הממשלתי.

לסיום פרק זה יצוין פעם נוספת כי המסמך לא מתיימר להציע פתרונות מפורטים לאתגרים הנובעים מהשימוש בטכנולוגיית הבינה המלאכותית, לקבוע הסדרים מפורטים או מתודולוגיה סדורה לעיצוב הסדרים שכאלו, או להכריע בשאלות כבדות המשקל הנוגעות לתוכן הרגולציה בענפים השונים או באופן רוחבי. עקרונות המדיניות הרגולציה והאתיקה, וההמלצות ליישומן המוצעות במסמך זה, נועדו לקדם תהליך בתוך הממשלה ויחד עם הרגולטורים הענפיים אשר בסיומו יהיה ניתן להתייחס אף לנושאים אלו.

לאחר ליבון הערות הציבור, נבחנו בשנית מספר עניינים שהוארו במסגרת השיח, ונעשו תיקונים וחידודים בתוכן המסמך ובמדיניות המוצעת (כגון בעניין המבנה, התפקידים והמעמד של מוקד הידע הממשלתי); ביחס לחלק מהעניינים הומלץ על המשך מעקב ובחינה בעתיד (כגון בשאלת הצורך בחקיקה רוחבית לתחום הבינה המלאכותית); ובנוגע לחלק אחר של העניינים יהיה בהחלט מקום להתחשב בהמשך הדרך, אף אם הם אינם בהכרח רלוונטיים למסגרת של מסמך זה (כגון הערות שעסקו בתוכן הרגולציה הענפית או בכלי לניהול סיכונים שהומלץ לפתח או לאמץ בעתיד).

תודות

מסמך זה נערך במשותף על ידי גורמי המקצוע במשרד החדשנות, המדע והטכנולוגיה בראשות מנכ"ל המשרד מר גדי אריאלי, ובמחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים בראשות המשנה ליועצת המשפטית לממשלה (משפט כלכלי) עו"ד מאיר לוי.

ריכזו את העבודה על המסמך, בגרסתו הסופית, עו"ד משה אוסטר מהלשכה המשפטית של משרד החדשנות, המדע והטכנולוגיה, ועו"ד יוסף גדליהו ממחלקת ייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים; בליווי עו"ד דני חורין, היועץ המשפטי למשרד החדשנות, המדע והטכנולוגיה ולמשרד התרבות והספורט, ועו"ד ד"ר יובל רויטמן, ראש אשכול רגולציה בייעוץ וחקיקה (משפט כלכלי). נטלו חלק בכתיבת המסמך עו"ד עמית אשכנזי, לשעבר היועץ המשפטי למערך הסייבר הלאומי. כן סייעו בהכנת המסמך עו"ד מיכל בן סימון וגב' נועה אלקין.

מלבד המעורבים בגיבוש המסמך בגרסתו הסופית, נבקש לציין כי טיוטת מסמך זה החל בתקופת שרת החדשנות, המדע והטכנולוגיה, ח"כ אורית פרקש-הכהן, ובהכוונת מנכ"לית משרד החדשנות, המדע והטכנולוגיה באותה העת, גב' הילה חדד-חמלניק. עוד היו מעורבים בהכנת טיוטת המסמך מר תום דן, גב' גאיה הררי הייט, וגב' שירן מימון.

לעבודה ולגיבוש ההמלצות שותפים רבים, וביניהם:

- **במשרד החדשנות, המדע והטכנולוגיה:** גב' הדסה גטשטיין, מר איתן טי, מר אליעז לוף, גב' רינת שפרן, מר ברק גטניו, וגב' אורלי ג'קסון.
- **במחלקת ייעוץ וחקיקה במשרד המשפטים:** עו"ד שרית פלבר, עו"ד סדריק (יהודה) צבע, עו"ד טל וורנר-קלינג, עו"ד אביטל שטרנברג, עו"ד ד"ר עמרי בן-צבי, עו"ד שרה גולד, עו"ד ד"ר ליטל הלמן, עו"ד לירון מאוטנר, עו"ד רני נויבואר, עו"ד ד"ר תמר קלהורה, עו"ד איל זנדברג, עו"ד לירון מאוטנר לוגסי, עו"ד מיכל אברהם, מר גל קורן, מר אלון פרבר וגב' נועם אמיר.
- **ברשות החדשנות:** ד"ר זיו קציר, עו"ד צפריר נוימן ועו"ד לי קפון.
- **ברשות להגנת הפרטיות:** עו"ד ראובן אידלמן, עו"ד שרון שמש עזריה, עו"ד ניר גרסון.

תודתנו נתונה לכל מי שעמלו על גיבוש התייחסות והערות לטיוטת המסמך, ובכלל זה, לפרופ' ניבה אלקין-קורן, פרופ' אסף חמדני וד"ר קרני שגל-פפקורן (מרכז מאיר שמגר למשפט דיגיטלי וחדשנות באוניברסיטת תל-אביב); לד"ר תהילה שוורץ אלטשולר, עו"ד עמיר כהנא, עו"ד גדי פרל וד"ר אורי פיירמן (המכון הישראלי לדמוקרטיה); לגב' הילה הובש (חברת מייקרוסופט); לגב' גאיה הררי הייט (מכון המחקר (SNPI) Start-Up Nation Policy Institute); לעו"ד אריאל יוספי, עו"ד און דבורי ועו"ד עודד קרמר (מחלקת הרגולציה של טכנולוגיות ומסחר אלקטרוני במשרד עו"ד הרצוג, פוקס, נאמן); לד"ר טל מימון וד"ר אלעד גיל ("יתכלית" המכון למדיניות ישראלית ומרכז פדרמן לחקר הסייבר באוניברסיטה העברית); לד"ר גלית ולנר (המכון הטכנולוגי חולון ואוניברסיטת תל אביב); לד"ר הדר יוענה ז'בוטינסקי וד"ר מיכל לביא (מרכז הדר ז'בוטינסקי לחקר בין-תחומי של שווקי הון, משברים וטכנולוגיה); לד"ר דניסה קרה רשף, מר אופיר צ'רניאק ומר עידו חפץ (אוניברסיטת בר אילן); לד"ר רפאל ויונטי (חברת מטא והאוניברסיטה העברית); לד"ר שרון ילוב-הנדזל (מכללת אפקה להנדסה); לפרופ' רן גלעד-בכר וד"ר דלית קן-דרור פלדמן (הקליניקה למשפט, לטכנולוגיה

ולסייבר באוניברסיטת חיפה, ומעבדת MLwell באוניברסיטת תל-אביב; למר משה גרינהוט (מכון משפטי ארץ); ולחברת AWS – amazon web services. ההערות שהתקבלו, בכתב ובע"פ, היו מעמיקות ומועילות, ורבות מהן הוטמעו במסמך זה והביאו לשיפור המדיניות המותווית בו.

תודה נוספת לכל מי שהשתתפו בכנס ה"שולחנות העגולים" בארגון משרד החדשנות, המדע והטכנולוגיה וייעוץ וחקיקה (משפט כלכלי) במשרד המשפטים, וביום העיון בארגון מרכז הנשיא מאיר שמגר למשפט דיגיטלי וחדשנות באוניברסיטת תל אביב, וכן לחברי קהילת "סינרגIA" – קהילת החוקרים שמוביל משרד החדשנות, המדע והטכנולוגיה בתחום הבינה המלאכותית.

2023

