

סיכום דיון

מסמך זה מהווה את תמצית הדיונים של "פורום שווי הוגן", מתוך מטרה לשתף את הקהל הרחב בעיקרי הדברים. בקריאת המסמך יש להביא בחשבון שמדובר בתמצית הדיון ולא בפרוטוקול מלא. בהתאם למדיניות הפורום הדברים מובאים שלא בהכרח בציון שמות כל הדוברים. יודגש כי הדברים מייצגים את עמדותיהם האישיות והמקצועיות של חברי הפורום, ואינם מייצגים בהכרח את העמדות הרשמיות של הגופים אליהם הם משתייכים.

תאריך המפגש: 05/01/2026

נושא המפגש: השלכות ה-AI על הסקטור הפיננסי על רקע דוח הרגולטורים הפיננסיים: כיצד יראו השירותים בעתיד ומה תהיה ההשפעה על השחקנים השונים?

חומר הרקע למפגש: [חומר רקע](#); [מצגת רגולטורים](#); [מצגת שלי תשובה](#)

לרשימת המשתתפים: [לחצו כאן](#)

צוות מקצועי: דניאל אלביליה, טוהר זיני, ניר פלד, דותן פרידמן

שלומי שוב

בוקר טוב, אנחנו עוסקים הבוקר בהשלכות של הבינה המלאכותית על הסקטור הפיננסי מהכיוון האסטרטגי – מה יהיו המוצרים והשירותים, כיצד יתנהגו הצרכנים, כיצד תושפע כמות עובדים ומהות התפקידים שלהם, כיצד תשנה מפת התחרות וכו'.

בחודש שעבר עסקנו אמנם גם בבינה המלאכותית אבל מהכיוון של שוק ההון. לא זוכר שב-15 שנות פעילות של הפורום היו לנו נושא קשור ברצף – ובואו, כנראה שעוד נעסוק בו עוד רבות בפורום, הרבה מעבר לשני המפגשים האלה. אנחנו עומדים בפני רעידת אדמה שכבר החלה אבל היא מתמשכת ומתגברת כך שאנחנו ממש בתחילתה ולא יודעים עדיין מה יקרה בסופה. דבר אחד בטוח – זה לא יישאר מה שהיה. יש כאן מהפכה אמיתית שמאתגרת את האנושות וכנראה לא ראינו כדוגמתה, כך שבאופן טבעי אנחנו עוד לא מודעים לגמרי להשלכות שלה בכל תחום.

ההשפעות על הגופים פיננסיים בטווח הזמן הקצר שכבר מתרחש בפועל הן בעיקר על היעילות התפעולית, אבל כבר בטווח הזמן הבינוני יהיה מדובר בשינוי בתמהיל כוח העבודה – שמשלב עובדים אנושיים עם סוכני AI ושיפור הפיריון במקביל כנראה לרמות שגרורות יותר. אבל זה רק בדרך וטווחי הזמן האלה הולכים ומתקצרים. בטווח הזמן הארוך מדובר כבר בשינוי המודלים העסקיים, וכן שינוי בטעמי הציבור ובמודלי ההפצה. מתי זה יגיע בדיוק אנחנו לא יודעים, אבל זה ברור שזה כבר צריך להיכנס כעת לחשיבה האסטרטגית של הגופים..

הדיון אצלנו היום מתבצע גם על רקע הדוח הסופי של הרגולטורים בנושא הסדרת השימוש בבינה מלאכותית בסקטור הפיננסי ועוד מעט הרגולטורים יציגו את המסקנות שלהם. בשורה התחתונה נבחרה גישה שיותר דומה לאמריקאית (שהיא יותר מאפשרת מאשר הגישה שמסתמנת באירופה של חקיקה מיוחדת רוחבית). כלומר, הקונספט הוא שהפיקוח יהיה על בסיס הרגולציה הקיימת ולא באמצעות חקיקה חדשה, וכל רגולטור יפעל כפי שימצא לנכון תחת תחום סמכותו. תכף ננסה להבין מה זה אומר והאם ניתן לפרש זאת כרגולציה מאפשרת, מתוך מטרה שלא להפריע לפיתוח החדשנות מצידם של השחקנים השונים בסקטור העסקי.

אנחנו יודעים שכל הגופים הפיננסיים, והנציגים של כמעט כולם יושבים כעת סביב השולחן, נמצאים כיום במרוץ סביב השילוב של אלמנטים שונים של ה-AI כדי לשמור על הרלבנטיות שלהם. לפי הדו"ח של מקינזי אליו התייחסנו בחומר הרקע שקיבלתם, בינה מלאכותית צפויה לחולל מהפכה בפעילות הבנקאית ולהפחית את עלויות התפעול בעשרות אחוזים, אך יתרון זה מטבע הדברים צפוי להישחק בטווח הארוך עקב תחרות והעברת החיסכון התפעולי ללקוחות. זה אומר כמובן שבנקים חלוצים ירוויחו משמעותית ואלו שיאחרו יעמדו בפני שחיקה.

בשנים האחרונות היו לא מעט רפורמות בסקטור הפיננסי – חיבור לתשתיות תשלומים, בנקאות פתוחה, מאגרי נתוני אשראי ועוד – שעד כה לא שינו דרמטית את מצב הדברים בפועל, אבל התחושה שלי היא שיחד עם המהפכה העצומה הזאת – יהיה כאן "גיים צייג'ר". זה בדיוק מתחבר לתפיסה של הצעירים שחושבים אחרת מאיתנו – הם ממש נולדו לתוך זה ולא מדובר מבחינתם בשינוי. הטעמים שלהם – החיפוש אחר הזול והמהיר ועם פחות רגשי נאמנות וציבות מאלו שלנו, יובילו להערכתי לשינויים במפת התחרות, במיוחד כשזה בא יחד עם נשמת אפה של המהפכה זאת – הורדת חסמים. תראו למשל מה קורה היום בתחום אפליקציות המסחר וההשקעות עם הצעירים.. כלומר, עם הדיסטראפן הטכנולוגי העצום הזה, ה"פריקות" הזאת שדיברנו עליה בלא מעט דיונים שהיו לנו בעבר בפורום על היבטי התחרות – באמת תקרום עור וגידים.

שימו לב, לא מדובר רק על שירות הלקוחות ששם זה בד"כ "רץ" קודם (וראינו את הנתון שבארה"ב כמעט מחצית ממבקשי השירות בבנקים נתקלו בצ'ט-בוט), אלא זה נוגע לכל הפעילויות לרבות תמחור וחיתום אשראי וביטוח, ניהול תיקי השקעות, שיווק שמותאם להתנהגות הלקוח... כלומר התהליכים של הבק-אופיס שיכולים גם לייצר יתרון תחרותי משמעותי מעבר רק ליעילות.

תחשבו על בניה וניהול מדויקים של תיקי השקעות בהתאמה אישית עם הוזלה משמעותית של העלויות והורדה דרמטית של רף הכניסה דה פקטו לשימוש בשירותי יועץ. זה ממש שינוי מפת התחרות. תחשבו גם שכל בנק או גוף השקעות יצטרך לייצר ערך מוסף כדי לשמר את הבידול שלו והיתרון שלו מול המתחרים.

אם אנחנו מסתכלים על תחום האשראי בבנקים שיש בו מרכיב דומיננטי של ניתוח וניהול סיכונים, אז הבינה המלאכותית יכולה להוביל בקלות לרמות דיוק ויכולות עיבוד גבוהות יותר מהאנושיות. רק תחשבו איזה שינוי דרמטי זה גם מבחינת שיפור החיתום וגם מבחינת החסכון בכוח אדם.

שימו לב שבביטוח זה מורכב ורגיש יותר כי זה מתחבר לאנשים – בין אם מדובר בבניית מודלים אקטוארים המתעדכנים בזמן אמת ובין אם מדובר ביכולת חיתום פרטנית המבוססת על מקורות מידע אלטרנטיביים וכל הרגישויות הנגזרות מכך.. זה הרי ברור שכלל שהרזולוציה של החיתום הביטוחי יורדת התמחור הכלכלי הופך להיות מדויק וכלכלי יותר לכל מבוטח. אבל גם בביטוח, יש הרבה מאוד בק אופיס טכני. לא מעט עובדים שמתעסקים בניידים, בקשות של לקוחות, אישורים, כל מיני סריקות ותהליכים תפעוליים אחרים. גם הטיפול בתביעות. אין סיבה שכל זה לא יתייעל – ולמעשה זה כבר קורה בימים אלה.

יש כאן שלל סיכונים התנהגות וסודיות וגם צורך בהסברתיות ושקיפות כך שאני צופה שמנהלי הסיכונים וקציני הציות יהיו מאד מאותגרים – תחשבו אפילו על דברים שלא נהוג לחשוב עליהם, כמו סיכון ההוגנות שלא תהיה אפליה למשל בשירות בין מגזרים שונים. האם נצטרך ללמד גם את הרובוטים להיות "פוליקטלי קורקט"? האם הרגולציה תיאלץ להתערב בה אם יתגלו מקרים של אפליה לכאורה? יש גם היבטים משפטיים, למרות שלא ניכנס אליהם היום, של אחריות – הרי כולנו יודעים שעם טעות של מכונה שצודקת ב- 98 אחוז מהפעמים קשה יותר להתמודד מאשר עם טעות של עובד אנושי שצודק רק ב- 90 אחוז.. בדיוק כמו המקרה של המוניות האוטומטיות..

גם מבחינת כמות עובדים אני מעריך שאנחנו נראה שינויים דרמטיים. במחקר של מכון טאוב שפורסם לפני שנה על שוק העבודה הישראלי הסקטור הפיננסי נמצא כבעל הפוטנציאל הגדול ביותר להתייעלות תפעולית בישראל – כלומר החלפת עובדים רבים בבינה מלאכותית והגדלת הפריון לעובדים הקיימים בתפקידים שיוותרו ויעברו טרנספורמציה של שילוב עם הבינה המלאכותית.

לא נראה לי שיש מישהו שחושב שהבנקים למשל ימשיכו להעסיק כמות כל כך גדולה של עובדים.. אז אם למשל שני הבנקים הגדולים מעסיקים כיום משהו כמו 8000 עד 9000 עובדים אני לוקח סיכון ומעריך שזה ירד בטווח הארוך ואפילו הבינוני בחצי לפחות - וכמובן שיש לזה המון השלכות....

ובכלל כל הקונספט של מתווכים, סוכנים ומפיצים יעבור זעזוע לדעתי בעולם החדש. השאלה עד כמה הוא ימשיך להתקיים – מניח שחלק לפחות מהתיווך יחתך ואלה שישארו יעברו טרנספורמציה. תחשבו למשל איך יראה התחום של סוכני הביטוח בעתיד – קשה להתעלם מכך שהשיווק הישיר באמצעות ה-AI צפוי להקטין את התלות בסוכן, אבל מצד שני יתכן והתפקיד שלו ישתנה והוא יהיה סוג של זרוע אנושית של חברות הביטוח לטיפול בתביעות, כלומר- תהיה כאן חלוקת עבודה.. סוכני ביטוח שיעשו התאמות ייהנו מאפקט ה-AI בעולם התפעול והתביעות וכנראה נראה המשך של קונסולידציה בתחום. בהקשר זה גם ראינו שהפניקס עשתה לאחרונה רכישה של סוכנות ביטוח מבוססת AI והמנכל שלה, אורן כהן, נמצא כאן איתנו ותכף נשמע ממנו.

כשמסתכלים על זה בראייה אסטרטגית כוללת של הטווח הארוך כמו שאמרתי, אז הבינה המלאכותית יוצרת הזדמנות תחרותית אדירה לגופים קטנים – פינטקים יכולים להיכנס בצורה אגרסיבית לנקודת רווח ספציפיות. אפרופו מה שאמרתי מקודם על נושא הפריקות של שרשרת הערך, שהרי לכאורה התשתית הרגולטורית לכך כבר הוסדרה. זה יגרום ללחץ על מחירים ומרווחים והנקודה המרכזית היא שאנחנו נכנסים לעולם שבו השוואת המחירים והמעבר בין גופים נעשים קלים יותר ויותר. כלומר, היתרון התחרותי ייגזר פחות מגודל ונאמנות ללקוחות ויותר מהיכולת להתייעל מהר, לשלוט בחוויית הלקוח ולהציע ערך מוסף. תראו למשל את המקרה של Revolut שבהחלט מסמנת דיסטראפשיין פוטנציאלי.

מצד שני יש אפקטים מקזזים לכאורה - לגופים הגדולים יש יתרון ביציבות ובשמרנות שלהם - מבחינת רמת אמינות ואיכות השירות וכמובן באלמנטים של סודיות ואבטחת מידע. אבל מה שכן לא בטוח כמה האפקט המקזז הזה אפקטיבי כאשר מדובר בצעירים.. מה שבכל זאת יהיה יתרון משמעותי לגופים הגדולים הוא כמובן המידע שלהם על הלקוחות ודפוסי ההתנהגות שלהם – אם זה דפוס הפעילות בעו"ש בבנקים או בביטוחי חיים בחברות הביטוח.

נמצאים איתנו היום כל הגורמים הרלבנטיים – הרגולטורים הפיננסיים והגופים עצמם - ולהערכתי הולך להיות לנו אחד הדיונים החשובים שהיו לנו בראיה עתידית. רק לפני שנמשיך אציין שבמסגרת מכון הפניקס לחקר שוק ההון שתחתיו נמצא כיום פורום שווי הוגן, הקמנו לאחרונה מרכז מצויינות ל-AI בפיננסים ויש לנו שם שיתוף פעולה חזק עם בית הספר למדעי המחשב וכיתות הדאטה האחרות באוניברסיטה, כולל סגל וכולל סטודנטים מצטיינים בתארים השניים, והתחלנו כבר מספר פרויקטים משותפים של מרכז המצוינות עם גופים פיננסיים שונים בדיוק על הדברים האלה. ואם יש גופים נוספים שנמצאים כאן ועדיין לא הצטרפו אז נשמח.

אמיר ברנע

רק כמה הערות קצרות לפני שהדיון יעבור לרגולטורים ולמומחי AI בסקטור הפיננסי. אני קצת פחות מתלהב משלומי לגבי ההפיכה שבדרך. כנראה שהסיבה היא הפרש הגילים בינינו, ואולי גם על איזושהו פער מתודולוגי.

למדתי לפני שנים, וגם לימדתי לפני שנים, את מה שקוראים כלכלה פוזיטיבית. כלכלה פוזיטיבית היא טכניקת חשיבה שבה באופן תיאורטי נגזרת היפותזה ממודל התנהגות של אנשים, שאחרי פיתוחה התיאורטי היא נבחנת בשטח. האם היא מנבאת במבחן רובסטי התנהגות של האנשים ורק כשמקבל אישור אמפירי, ואישור אולי גם מחוקרים אחרים, היא מצטרפת למסה הזאת של מתודולוגיה מחקרית. כלומר, קודם כל דיון תיאורטי, שגוזר היפותזה, אחרי זה הבחינה האמפירית, ורק אחרי זה הצטרפות למדע.

ב-AI מתחילים בנתונים וגוזרים כללי התנהגות מתוך התנהגות הנתונים. כלומר, אפשר להגיע באמת לאבסורדים, מתוך הנתונים שהם הבסיס. המודל הזה לא יכול להתייחס למשל למרכז תורת המימון

שהוא, קשר בין תשואה וסיכון, דיון עומק על סיכון שיטתי, מזכיר לכם דברים שלמדתם, אולי לחלק מהוותיקים בתורת המימון. כאן מתוך הדטא גוזרים כללי התנהגות. עכשיו, להתחיל מתוך הדטא ולגזור כלי התנהגות, קראנו לאלה שעוסקים בזה בשם המעליב צ'רטיסטים. אלה שהיו מסתכלים על הצ'ארט, ומנבאים מה יקרה לשוק ההון מתוך הצ'ארט. אנחנו פסלנו אותם כבר במבחן החלש להזכירכם, מבחן הכי חלש ליעילות שוק ההון, שאנחנו עוברים בקלות, אנחנו מוחקים אותם. אני פותח היום את העיתון כלכליסט, מאמר של הכלכלן הראשי של מיטב. הוא בא ואומר לנו, אם אני קורא נכון, "ההיסטוריה באופן מובהק, לא מנבאת הצלחה למניות בוול סטריט בשנה הקרובה... כי אם שלוש שנים זה עלה ב-60%, אז בשנה הבאה זה במינוס 2%". זאת אומרת, ברור, זה AI, זה מודל AI, מה שהוא משתמש. אני מניח שהוא הפעיל מודל AI. מודל AI אמר לו, אם שלוש שנים זה עולה ב-X, בשנה הרביעית זה Y. צ'ארטריסט.

שלומי שוב

אמיר, נראה לי שאתה קצת מזלזל ב-AI...

אמיר ברנע

אני שואל את עצמי את השאלה הפשוטה הבאה, אם אדם מבקש אשראי מהבנק, ונניח ששם המשפחה שלו מתחיל באות ב', ואם יתברר שיש שכיחות גבוהה לאנשים באות ב' שלא החזירו אשראי והבקשה נפסלת, איך הוא יסביר לי במסגרת ההסברתיות שאתה כותב עליה, איך יסבירו למה הוא לא יכול לקבל אשראי כי במקרה השם מתחיל באות ב'.

זאת אומרת, אין כאן איזה תיאוריה, למה אדם כמו שאמרתי ששמו מתחיל בב', הוא כאילו נחות יותר, מבחינת יכולת החזר שלו. החסר בתיאוריה, בעיניי אולי בגלל הוותיקות, הוא בעייתי. לדעתי זה לא נושא שולי כי אני מתייחס למה ששלומי דיבר על הטווח הארוך, AI הוא מחליט וההחלטה מתבססת רק על דאטה ללא תיאוריה.

כלומר, אין ויכוח בינינו ש-AI יוצר התייעלות, אין ויכוח בינינו שהוא מכשיר עוזר להחלטות, בעיקר ברמה הטכנית. אין ספק ש-AI מסייע לנו ועוזר לנו בהחלטות. אבל השאלה האם הוא מחליף אותנו בהחלטות. לא בטווח הרחוק, כבר עכשיו, AI מחליט לגבי איומים אולי יותר קטנים, וזה רק ההתחלה, אתה אומר, נדבר על מהפכה ההחלטות הן היעד.

למשל בחומר הרקע יש שלוש החלטות, יש החלטת השקעה, החלטת אשראי והחלטת ביטוח. אני שואל את השאלה, איך מערכת AI תעזור בהחלטת השקעה בשוק ההון. האם, אתה אומר הדור הצעיר, האם מערכת ה-AI תבוא ותגיד לאדם במה להשקיע? ואיך לפעול? אני לא אומר התאמת התיק לצרכים של הלקוח, נניח. כאן מדובר על בחירה סלקטיבית. האם הוא הולך להגיד לנו, מניה א' כן, מניה ג' לא?

שלומי שוב

יועץ השקעות גם לא אומר את זה.

אמיר ברנע

קרן נאמנות שאני משקיע בה?

שלומי שוב

זה משהו אחר. קרן נאמנות, אתה משקיע בקרן נאמנות, היא משקיעה עבורך. יועץ השקעות הוא נותן לך את כל הכלים כדי שתקבל את ההחלטה, ההחלטה נשארת אצלך.

אמיר ברנע

תראו, זה שהעלויות תרדנה אני מקבל, ואתה שלומי הדגשת זאת. האם הסחירות לא תיפגע? הרי השימוש של מודל ה-AI הוא סטנדרטי? או לחילופין לכל אחד יהיה סט נתונים אחר שממנו יבנה המודל? האם נשים את מבטחנו בריבוי חברות AI, כמו ChatGPT ואחרים, Gemini וכו', האם פה נוצר הפער? האם לא תיפגע הסחירות? כאשר המודלים הופכים להיות סטנדרטים. זו שאלה שאנחנו שואלים ואין לנו תשובה.

אשאל שלוש שאלות ואסיים. שאלה ראשונה, מודל AI אחד או ריבוי מודלים שמבוססים על ריבויים של סט נתונים, כל אחד עם סט נתונים אחר. זה מה שהולך להיות בעתיד? או יהיה סט נתונים גלובלי שתהיה גישה אליו, בתוקף הרגולציה או אחרת, ואז הסטנדרטיזציה של ההחלטה?

אני שואל, מודל ה-AI עובד עם המגבלות ששלומי הזכיר, מה שאתה מכנה פוליטי קורקט, האם עצם הכנסת מגבלות למודל AI לא פוגעת בכל האופטימיזציה שלו? נניח צבע עור, מין, ג'נדר כמגבלות. אסור לך להפלות על בסיס כזה או אחר, מה זה עושה למודל?

שאלה שלישית, האם הרגולציה מעורבת במודל עצמו או רק בשאלת ההסברתיות? האם היא נכנסת למודל עצמו? ואתה אומר לי שהגישה בארץ, בהבחנה מהגישה האירופאית, אם הבנתי נכון, היא בעצם אתה ממשיך לפעול, ואתה לא נכנס למודל עצמו, אלא רק דן בתוצאות שלו במקרים הפרטניים.

שלומי שוב

תכף נשמע את הרגולטורים בזה

אמיר ברנע

חברים, מצפה לדיון מעניין.

שלומי שוב

רגע לפני שאנחנו שומעים את הרגולטורים, אני רוצה שנשמע את דני ימין יו"ר בנק דיסקונט. דני, לך יש זווית מאוד ייחודית ומעניינת, כי אתה איש טכנולוגיה ובין היתר היית בעבר מנכ"ל מייקרוסופט וכעת יו"ר של גוף פיננסי גדול. החיבור הזה מאד מעניין ונשמח לשמוע איך אתה רואה את הדברים.

דני ימין

פעם ראשונה שאני פה, אז אני רוצה להודות לשלומי, לא הכרתי את הפורום. אז באמת מרשים מאוד. ראיתי את הקמפוס הבוקר וליבי התרחב מהקמפוס היפהפה הזה.

אני מסתכל על הנושא של AI בשני מימדים, יש את המימד שאני חושב שכולנו מכירים אותו, שהוא גם בטווח הקצר והבינוני המימד שקוראים לו המהפכה התעשייתית. האנלוגיה הזו שהמנועים החליפו סוסים, מכוניות ורכבות החליפו כרכרות וכן הלאה, ובעקבות זה העולם די השתנה. ואכן, האנלוגיה הזו היא נכונה, בבנקים כמו שהוזכר לוקחים תהליכים תפעוליים ומשתמשים בAI לשנות אותם, ולוקחים תהליכי קבלת החלטות כמו איך לתת אשראי, או למי לתת אשראי, או מודלים של חיתום, ומנסים להטמיע AI. ויש גם דוגמאות מעולמות אחרים, עם צילומים של MRI וST, במקום שיש איזה מקצוע שנקרא רדיולוג. שאני סיימתי את הלימודים ועבדתי בחברה בשם 'אלסינט', ואז מקצוע רדיולוג זה היה הדבר הכי חשוב שיש, אותו אדם שפענחו את הצילומים, אז המקצוע הזה ילך ויעלם כי הAI יודע לפענח את הצילומים יותר טוב ולהגיד אם יש גידול או אין גידול.

אז יש את המימד הזה, של המהפכה התעשייתית כולנו מנסים בכלים שאנחנו מכירים, ללמוד את הטכנולוגיות החדשות ולהטמיע אותן כל אחד בארגון שלו. אני רוצה להציע מימד נוסף, מימד אחר, שהוא גם לטווח היותר ארוך שלדעתי עוד אין לנו את הטרמינולוגיה והכלים להתמודד איתו. ואני קורא לזה גילוי של ציוויליזציה חדשה במאדים.

האנושות מחפשת הרבה זמן ציוויליזציה בכוכב אחר, ותארו לכם שהתעוררנו בבוקר וגילינו ציוויליזציה חדשה במאדים. יש שם מיליארדים של גמדים קטנים, ולא גמדים, בצבע מסוים, שנראים בצורה מסוימת, ויש להם סוג של אינטליגנציה, בחלק מהמקרים יותר גבוהה משלנו, בחלק מהמקרים אולי פחותה. יש להם סוג של רגש, יש להם סוג של קשרי גומלין בינם לבין עצמם, יש להם אורח חיים כזה או אחר. יש ציוויליזציה חדשה במאדים, ואנחנו גילינו אותה. ועכשיו נשאל השאלה מה עושים? אני לא חושב שיש בחברה האנושית יכולת לדמות דבר כזה, חוץ מאשר בספרי מדע בדיוני.

ואני חושב שלא ירחק היום, וה AI יהיה בעצם כמו ציוויליזציה חדשה רק שהיא לא תהיה במאדים אלא פה. עכשיו, מה יהיה האתגר המרכזי שלנו כאנושות סביב זה? אני לא יודע להגיד. איך נתמודד עם זה? אני לא יודע להגיד. מה בכלל תהיה התוצאה הסופית של הסינתזה הזו בינינו לבינם? אני לא יודע. אני כן יודע דבר אחד, אנחנו נהיה מוקפים במיליארדי ישויות עם אינטליגנציה, הרבה פעמים, יותר גבוהה משלנו, ואנחנו נצטרך לדעת להסתדר איתם.

וזה יוביל להרבה שאלות מוסריות, ביולוגיות, פילוסופיות, כמו למשל האם לישויות האלה יש זכות בחירה? כיצד הם מתרבים? מה מערכת יחסי הגומלין בינינו לבינם? וכו'.

זו מחשבה שלא מובילה כמובן את היום יום שלי, כי היום יום שלי הוא הרבה יותר פרקטי ומעשי, אבל אני חושב שכשדנים בהסתכלות כזו, וקראתי גם את הדו"ח שצורף, אני חושב שחשוב להסתכל על שני מימדים האלה, המימד האחד הפרקטי, המעשי, מה שאני יכול לומר, המהפכה התעשייתית, והמימד של המדע בדיוני, אבל הוא כבר לא כל כך מדע בדיוני. הוא קרוב, הוא יגיע. אם זה יהיה עוד 10 שנים, 20 שנה, 30 שנה, אני יודע. אם הוא לא יגיע אחרי כולנו, זאת אומרת, בימי חיינו נראה את זה. זהו, זה מה שרציתי לחלוק. אני רוצה שוב להודות לשלומי.

אמיר ברנע

אבל, סליחה. אתה יושב ראש של בנק, נכון?

דני ימין

אני יושב ראש של בנק.

אמיר ברנע

איך אתה רואה, בתקופה הנראית לעין, את השילוב של AI במסגרת הממשק בין הבנק ללקוח?

דני ימין

אנחנו מטמיעים AI בתהליכים תפעוליים רבים, כמו למשל שירות לקוחות או דברים מהסוג הזה. אנחנו גם מנסים להאניש בא' את הAI, דרך אגב, היה לא מזמן מאמר בניו יורק טיימס שאמר על בסיס מחקר כלשהו שזו טעות להאניש בא' את הAI, ודווקא להפך. כלומר, לתת לו את הזהות הAIית שלו. זאת אומרת, אם אנחנו לסוכנת AI שלנו, של הצ'אטבוט קוראים עדי, לא בטוח שזה הרעיון הכי נכון לנסות לתת לה פנים אנושיות. אבל זו דוגמה למה שאנחנו עושים.

אני אציין עוד תופעה מעניינת שקורית לפחות אצלנו, אני מסתכל על כל הפרויקטים של הAI שאנחנו עושים. להערכתי, ואני מאוד מעורב בזה גם באופן אישי, בערך 50 אחוז צומחים מלמטה. זאת אומרת, בשונה מזה שהדברים היו קורים מלמעלה, ופרויקטים ומגדירים וכו', אתה רואה דברים שקורים מלמטה, על ידי אנשים שבאמת לא חשבת שהם יבואו עם הדברים האלה. זה מסיבה נורא פשוטה כנראה, כי כולנו משתמשים היום בAI לצרכים אישיים רבים, בונים מסלול של טיול, שואלים אותו שאלות לגבי בריאות וכולי.

אנחנו רואה בהחלט שימוש הולך וגובר ונכנסת שפה, נקרא לזה שפה AIית לארגון. פשוט שפה AIית לארגון. הרבה מדברים על זה ש AI ייתר תפקידים בארגון, בהחלט יכול להיות, אבל עבורינו זה לא לב העניין וזה אפילו מרדד את הדיון, זה יהיה אולי נגזרת, תהיה תוצאה, אבל זה לא מטרה, אם למשל אנחנו מסתכלים על מודלים של נגיד לצורך העניין של תמחור פקדונות או מודלים של אשראי,

המטרה שלנו לעשות את זה הרבה הרבה יותר מדויק, הרבה יותר נכון, הרבה יותר מתאים, הרבה יותר אפקטיבי באמצעות AI.

אמיר ברנע

אי אפשר לרמות את השיטה הזאת? הלקוח עצמו שידע שזה מהלך קבלת ההחלטות.

דני ימין

הלקוח ידע שזה AI כמו שהוא יודע היום שמי שמטיס את המטוס זה המחשב ולא הטייס. פסיכולוגיה הוא יודע שיש טייס בתא וזה חשוב אבל המטוס טס אוטומטית, והלקוחות ידעו את זה, וישלימו עם זה, ואני אומר שוב, הם ידעו שזה AI, ובהרבה פעמים זה ייתן להם יותר ביטחון.

שלומי שוב

בואו נשמע כעת את הרגולטורים.

אמיר וסרמן

שמי אמיר וסרמן, אני היועץ המשפטי של רשות ניירות ערך ועמדתי יחד עם שרית פלבר ממשד המשפטים בראש הצוות הבין משרדי לבניה מלאכותית בפיננסים. רציתי בפתח הדברים להודות לשלומי, על היוזמה לזמן את הפורום בנושא חשוב זה, וגם על ההכנה היסודית למפגש. בזכותה אנו מרוויחים עוד אפילו לפני שהמפגש מתחיל, ובוודאי שנלמד גם ממנו.

בצוות היו חברים משרדי המשפטים והאוצר, שלוש הרגולטורים הפיננסים ורשות התחרות, במטרה להציג ראייה רחבה ואחידה. נסקור בכמה שקפים בודדים כמה מתכני הדוח שפרסמנו לאחרונה.

בדוח ניסינו להתמודד עם היבטים רבים שקשורים לכניסת בינה מלאכותית לסקטור הפיננסי. הדילמה בעיסוק בנושא נובעת מכך שמצד אחד יש מגוון רחב מאוד של שימושים אפשריים לבניה מלאכותית, וגם ההשפעה הפוטנציאלית של הטכנולוגיה היא גדולה מאוד. מצד שני, בשלב הזה, השילוב של הטכנולוגיה בפועל, נעשה בקצב איטי ומבוקר. זה מאפיין גם גופים פיננסיים בחו"ל – אפשר לראות שהטמעת הטכנולוגיה היא איטית יחסית - יש עניין רב מאוד, השקעה של תשומות לא קטנות, אבל היישומים בפועל עדיין אינם רבים. קושי נוסף בעיסוק בנושא הוא שהטכנולוגיה עצמה מתפתחת מאוד מהר. הצוות שלנו התחיל לפעול זמן לא רב אחרי ההשקה של ChatGPT, אשר שינה את התפישה בדבר יכולותיה של הבינה המלאכותית והתחיל בעקבותיו שיח נרחב בנושא מודלים של LLM. בשנה האחרונה יש כבר עיסוק ב-Agents. כלומר, בפרק זמן קצר יחסית יש התפתחות של המודלים, של הטכנולוגיה, ורגולטורים הנדרשים לעצב מדיניות, מדובר באתגר מורכב. לכן ראינו בדוח נקודת פתיחה לעיסוק בנושא, והוא כולל המלצות שאפשר ליישם מיידית, אבל גם לא מעט המלצות שקוראות להמשיך לבחון את המצב ולעקוב אחר התפתחויות.

מה כולל הדו"ח? הדוח נפתח בסקירה של ההתפתחויות הרגולטוריות בעולם ובישראל. הסקירה עוסקת ברגולציה על בינה מלאכותית באופן כללי, ואחר כך ממשיכה לרגולציה על בינה מלאכותית בסקטור הפיננסי. עבודת הצוות התחילה אחרי שהתפרסם מסמך מדיניות של המדינה, שקבע כי בשלב הנוכחי לא תהיה רגולציה רוחבית על בינה מלאכותית. כלומר, התחלנו מנקודת מוצא לפיה הרגולציה בישראל תהיה ענפית, ובכל תחום שקיים בו פיקוח, הרגולטורים יתמודדו עם הסוגיה של בינה מלאכותית. זו גישה שונה מהגישה האירופאית, והיא דומה יותר לגישה האמריקאית והבריטית.

בשלב הראשון, ניסינו להסביר מהם העקרונות הראויים לדעתנו להתמודדות עם בינה מלאכותית, ועסקנו גם בשאלת הגדרת בינה מלאכותית. מדובר בעניין חמקמק ומורכב, והצענו הגדרה לצרכים של אסדרה פיננסית.

הפרק הבא בדוח עוסק בדיון והמלצות לגבי סוגיות רחב של הסברתיות, מעורבות אנושית, אחריות, ועוד. אלה בעצם סוגיות שיעסיקו כל גוף פיננסי השוקל להטמיע מערכת בינה מלאכותית, וכפועל יוצא מכך יעסיקו גם את הרגולטורים. אפשר כבר לתרגם את ההמלצות בעניינים אלה לחיי המעשה. למשל, גוף ששוקל להכניס מערכת בינה מלאכותית ישאל את עצמו - האם אני יכול לעבוד עם מערכת שאין בה הסברתיות? באיזה אופן? האם נדרשת מעורבות אנושית? על מי חלה האחריות? הדוח מבקש לענות על שאלות אלה ולספק הכוונה לגופים ולרגולטורים.

החלק הבא עוסק בהשפעות שיכולות להיות לבינה מלאכותית על השוק בכללותו, בהיבטים כמו תחרות, יציבות פיננסית, והשפעות דיסאינפורמציה. אחר כך דנו בשלושה ענפי פעילות ספציפיים שהצוות התבקש לבחון אותם, תוך דיון בהשפעות האפשרויות של בינה מלאכותית, הרגולציה הקיימת ושינויים שנחוצים בה. הדוח מסתיים בפרק שעוסק בצד המדינה - מה צריכה המדינה לעשות כדי לעודד כניסה של הטכנולוגיה, כיצד יכולים הרגולטורים עצמם להיעזר בטכנולוגיית AI, למשל לצרכים של פיקוח.

הגישה הכללית שהצגנו בדוח היא שיש מקום לעודד כניסה של בינה מלאכותית, ואפשר וצריך להתמודד עם הסיכונים הגלומים בה. הבינה המלאכותית מעוררת הרבה חששות ושאלות, למעשה רוב הדוח שלנו עוסק בסיכונים האלה, אבל תחום הפיננסים נתון לפיקוח הדוק יחסית. הרגולציה בתחום הבינה המלאכותית היא המפתח להתמודד עם הסיכונים. בהשוואה לתחומים אחרים, כמו תחומי היצירה או תחומים אחרים בהם אין פיקוח, יש בידינו כלים רגולטוריים שמאפשרים להסדיר את השוק ולווסת אותו, לראות שהוא מתנהל באופן נכון, תוך הגנת האינטרס הציבורי. בגלל מערכת הכללים הרגולטורית, שאפשר לערוך בה את ההתאמות הנדרשות, הרגשנו שאנחנו על קרקע בטוחה יחסית להגיד "כן" בשאלה אם צריך לחתור לאימוץ הטכנולוגיה. אנו סבורים שצריך לעודד אימוץ של בינה מלאכותית בסקטור הפיננסי, ובין העקרונות שהצגנו, דיברנו על חשיבות התאמת האסדרה למקובל בעולם, על גישה של ניטרליות טכנולוגית לפיה בינה מלאכותית בפני עצמה לא מחייבת בהכרח שינוי רגולטורי אלא צריך לבדוק את ההשלכות שלה על השירות או המוצר הפיננסי, ולתת מענה בהתאם לבעיות שמתעוררות. דגש נוסף בעקרונות שהצגנו הוא נקיטה בגישה מבוססת סיכונים,

שמתמקדת בעיקר במערכות בעלות סיכון גבוה. בדוח הסברנו גם איך לאמוד את הסיכון - הוא כולל הן סיכון ללקוחות או למשקיעים, הן סיכונים לגוף הפיננסי עצמו, וצריך להביא בחשבון את שני הממדים. כאשר גוף פיננסי מאמץ מערכת של בינה מלאכותית, עליו לאמוד את הסיכון הגלום בה, ולהפעיל כלים משמעותיים יותר אם מדובר במערכת בסיכון גבוה. אעביר כעת לשרית שתציג שתי סוגיות שדנו בהן בדוח.

שרית פלבר

שלום אמיר, שלום לכולם, תודה על ההזמנה. כפי שאמיר אמר, זה דו"ח ארוך. אנחנו קראנו אותו וקראנו לפני שפרסמנו אותו, ומאיר שיושב כאן לשמאלי, אמר לי שרית, הפרק שתעשי לו הגהה הכי טובה זה פרק התודות.

בשולחן נמצאים הרבה מאוד אנשים שתרגמו לדו"ח, תרגמו להקמת הצוות אז זו הזדמנות להגיד באמת תודה גדולה למי שדחף אותנו לצוות, לחברים לצוות, אז תודה גדולה. כנראה שלא לקחתי את העצה של מאיר ולא חשבתי על זה מראש, אז אני מתנצלת אם שכחתי משהו, אבל זו הזדמנות טובה לומר תודה רבה על התקופה וגם על התרומה הזאת. אז כמו שדובר כאן כבר בהצגה של הדברים, אחת הסוגיות המרכזיות שמתעוררות כמין כזה שמש אסוציאציות שחושבים על בינה מלאכותית זו הקופסה השחורה. ואומנם שלומי אמר לי לא להיות משפטנית מדי, אבל אני משפטנית, ואין דבר שמפחיד יותר משפטנים מהקופסה השחורה הזאת. לא לדעת למה המערכת קיבלה את ההחלטה שלה. תחשבו כמה המערכות הרגולציה שלנו בנויות סביב היכולת להסביר פעולה או החלטה. אם יש טענות להטיה זה עובר דרך ההסברתיות. לראות אם המערכת מגיעה לתוצאות שהגוף רצה עבורה, זה עובר דרך ההסברתיות.

מצד אחד היה לנו את המשקולת הזו שההסברתיות היא כלי להגשמת הרבה אינטרסים משפטיים. מצד שני, אי אפשר להתעלם מהמציאות ש מבחינה טכנולוגית ההסברתיות לא שם. היא כנראה מתקדמת, אנחנו יודעים יותר דברים ממה שידענו בעבר.

אבל כמערכת סטטיסטית שעובדת על נתונים כל כך רבים, יהיה מאוד קשה לזקק מדוע המערכת קיבלה את ההחלטה שהיא קיבלה. אני חושבת שניסינו להלך בין הטיפות,

כלומר לאזן בשני הדברים, מצד אחד הצורך בהסברתיות ומצד שני להתחשב בקושי הטכנולוגי. יש על זה הרבה אמירות, למשל ב-gdpr, שיש שם איזושהי זכות להסברתיות, שהיא יכולה להוות חסם מאוד משמעותי לשילוב של בינה מלאכותית בכל מיני דברים שאולי היינו רוצים שיהיה בהם בינה מלאכותית. אז זה רק כדי לספר על מערכת השיקולים והאיזונים שמיפנו ושהנחו אותנו. זה פרק, שהיה אחד הכי סוערים בדיוני הצוות, בניסיון להגיע להסכמות מאוזנות. אני מקווה שגם הצלחנו.

אני קצת אפרט על ההמלצות ועל השיקולים. קודם כל בספרות וברגולציה יש הרבה מונחים שמתחלפים, יש שקיפות, יש הסברתיות, יש אינטרפיליות, ואנחנו ניסינו רגע לשים בצד את כל המונחים שהנחיתו עלינו, ולחשוב איך אנחנו מתייחסים לזה מהותית. והתוצאה הייתה, אני מקווה

שהיא מספיק ברורה, הפרדנו בין הסברתיות כללית, שהיא יותר נוגעת לאופן פעילות המערכת, היא דומה יותר לשקיפות אם תרצו, לבין הסברתיות פרטנית, שזו הסברתיות ביחס להחלטה ספציפית - למה ההחלטה התקבלה, איך היא התקבלה, ומה השיקולים שהנחו כאלה ומה היו השיקולים המרכזיים. הסברתיות פרטנית היא כמובן הגולדן סטנדרט שקשה להגיע אליו, ולכן לא רצינו ללכת לתוצאה שאומרת חייבים הסברתיות פרטנית בכל מקרה ומקרה. רצינו ליחד את ההסברתיות הפרטנית הזאת למקרים שבאמת היא נדרשת בגלל היותה חסם. כלל אצבע הראשון שהצבנו לעצמנו, הוא מה הדין הקיים אומר. אם בדין הקיים יש לפרט זכות לקבל הסבר פרטני, אז זה משהו אחד. אם אין היום חובה לקבל הסבר פרטני, קשה היה לנו להצדיק למה בינה מלאכותית משנה מהאיזון הזה. כלומר אם היום אין בדין הסבר, חובה לתת הסבר להחלטת אשראי. עלינו לשאול

האם כניסתה של בינה מלאכותית משנה מהאיזון הזה? אז כלל אצבע ראשון שקבענו, הוא איפה שלא נדרש היום בדין, בינה מלאכותית לא אמורה לשנות. כל רגולטור לשקול את הנסיבות כן, אבל לא כאיזשהו כלל חוצה.

איפה שאנחנו גם שצריכה הסברתיות פרטנית, זה בהתאם למהותיות ההחלטה. למה אנחנו רוצים בעצם הסברתיות פרטנית? כדי שהפרט יוכל להשיג על ההחלטה, להגיד רגע, טעיתם בענייני, לא נתתם משקל לאיזשהו שיקול מסוים. בלי הסבר הוא לא יוכל לעשות את זה. וגם שאולי הפרט יוכל לשנות את ההתנהגות שלו לעתיד. תחשבו בהחלטות אשראי, אם הפרט יודע מה הייתה הסיבה שלא קיבל אשראי, הוא יוכל לקחת את זה לתשומת ליבו ולשנות את ההתנהגות שלו באופן שישפיע על החלטות עתידיות. בהחלטות מהסוג הזה, בהחלטות שמהותיות לפרט, נראה שיש משמעות להסברתיות פרטנית.

דבר אחרון, וכמובן, מה החלק של הבינה מלאכותית בקבלת ההחלטה. אם בסוף יש אדם שמקבל את ההחלטה, אם היא נמצאת ב-back office, אז לא משנה מה האיזון שקיים היום,

ולכן גם לא נדרשת איזושהי מערכת עם יכולת הסברתית. עוד דבר שאמרנו, שאם יש היום שני שירותים, יש לי שירות אנושי ושירות בינה מלאכותית. אמר שלומי שזה עניין דורי. בשבוע שעבר הרציתי באוניברסיטה ואמרתי את זה, הסטודנטים הסתכלו עליי, שבועדאי שהם יעדיפו שמערכת בינה מלאכותית תקבל החלטה בעניינם.

בכל אופן, עד שייסגר הפער הדורי, כשיש מערכת בינה מלאכותית ולצידה חלופת שירות אנושית שניתנת באותם תנאים, יכול להיות שאני יכולה לקחת מערכת שאין בה הסברתיות, כי הפרט יכול לבחור בין שני השירותים. ובנוסף, אמרנו שגם אם מכניסים מערכת בינה מלאכותית בלי הסברתיות, אפשר לפצות על זה בכל מיני חלופות אחרות,

כמו מעורבות אנושית יותר אדוקה, בקרה על התוצאות. זו הייתה הדרך שלנו, להתמודד עם הסוגיה הרגישה הזאת. ניסינו גם להניח בתפקיד איך חושבים על הסוגיה הזאת, ולא רק לקפוץ להמלצות,

בגלל זה הוא כזה ארוך. אבל גם אם לא מסכימים, אני חושבת שהוא מניח את התשתית שמאפשרת את החשיבה על הסוגייה הזאת. אני גם אגע בקצרה גם כן על סוגיית האחריות.

למשפטים אנחנו לא עסקנו בדו"ח באחריות נזיקית, עסקנו יותר בשאלה מי הגורם שאחראי לעמוד בכללי הרגולציה.

דובר כאן על החוק האירופאי, החוק האירופאי ממש פיצל את שרשרת השחקנים הרלוונטיים בפיתוח ושימוש במערכות בינה מלאכותית וחילק את האחריות הזו. אנחנו חשבנו שעוד לא בשלה העת להגיע למצב הזה, ושהבינה המלאכותית באופן שהיא משתלבת היום, לא אמורה לשנות כלל מאוד בסיסי.

כשהגוף הפיננסי, גם אם הוא עושה מיקור חוץ וגם אם הוא משתמש במערכות של גורם חיצוני, נותר אחראי כלפי הלקוח, וזה כלל שצריך לשמר. קיבלנו על זה קצת הערות ציבור.

בין היתר היכולת של משתמש הקצה אולי להתעלם מאזהרות ביחס לשימוש במערכת או להתעלם מדיסקליימרים מסוימים, מחייב אולי כן לחלק את האחריות באופן שונה בין המוסד הפיננסי לבין הפרט. אנחנו מסכימים שיכולים להיות מצבי קיצון כאלה, ונתנו לכך דגשים בדו"ח. אבל כן חשבנו שצריך לשמר את הכלל הזה. אנחנו בדו"ח לא חשבנו, לא עסקנו בעצם בשאלת אחריות של המפתח לצורך העניין. שוב, אנחנו הסתכלנו על הדברים בצד הרגולטורי.

המוסד הפיננסי, אין הבדל אם הוא מכניס מערכת טלפוניה מתקדמת, אם הוא מכניס מערכת שיחות מתקדמת. בעינינו בשלב הזה, זה דומה אם הוא מכניס מערכת בינה מלאכותית. לצורך העניין כמו שאומרים במשפטים, מונע הנזק היעיל ביותר, נותר המוסד הפיננסי. עליו לדאוג שיכניס את המערכת הכי טובה שהוא יכול, אבל בוודאי בוודאי שאם הוא שגה, הוא לא יכול לשלוח את הפרט להתדיין מול מפתח המערכת, וזה ברמה, ברמת העסק. אז זה ככה טעימה משתי סוגיות.

אביבה וייס

שלום לכולם, אני אביבה מרשות שוק ההון, ביטוח וחסכון.

יחד עם יעל רגב היינו הנציגות של רשות שוק ההון בצוות הבין המשרדי שהפעילות שלו עזרה לקדם ולפתח את החשיבה שלנו כרגולטור בתחום הזה של בינה מלאכותית, ובאמת הדו"ח פורס את הסיכונים שאנחנו רואים, כמובן יחד עם היתרונות.

וגם שינויים רגולטוריים שאפשר להתחיל לאמץ אותם כדי להתמודד. הקופסא השחורה היא סיפור חדש, אנחנו לא מכירים אותו ממערכות קודמות. וגם יש פה משהו שמאוד מלחיץ אותנו כרגולטורים. למה זה ככה?

יש מנגנון מאוד מאוד בסיסי, מנגנון בקרה שיש אצל כל ארגון, בכל תחום, בכל גודל, שהוא כמעט מובן מאליה, זה המנגנון של מדרג קבלת החלטות. מה זה אומר? זה אומר שיש דברים שעובדים עושים עצמאית, יש דברים שעולים לרמת מנהל מחלקה, יש דברים שמתיעצים איתם על עוד גורמים, יש דברים שרק הנהלה יכולה לקבל החלטה ויש החלטות שחייבות להגיע לדירקטוריון.

מי שרוצה לקדם את ההחלטה בא, מסביר, מציג את האירוע, אומר מה השיקולים, מציג חלופות, מציג איזושהי המלצה, והכל על בסיס הנתונים והמידע.

ופה, כשנכנסת לפה, הקופסה השחורה, היא מכניסה לתוך כל הדבר הזה עמימות. איך עכשיו תתקבל החלטה בעצם? כי הקופסה השחורה בעצם, הבינה המלאכותית נותנת לנו את התשובה. או את ההמלצה או את הכיוון, אבל אנחנו לא יודעים בדיוק איזה מהנתונים שהוזנו אליה השתמשה, מה היה הנתון שיותר השפיע, מה היה פה יותר חשוב. זה מכניס איזושהי עמימות בתוך תהליך קבלת ההחלטות. ועמימות בדרגות מסוימות יכולה גם להגיע ממש לאיבוד שליטה. וזה משהו שאנחנו כרגולטור כמובן לא יכולים לאפשר אותו. אנחנו חושבים שהדבר הזה עדיין לא פתיר עד הסוף, זה משהו שצריך להיות פתיר, צריך לטפל בו כדי שבאמת יהיה אפשר לאמץ את הבינה המלאכותית. אז איזה כלים בכל זאת יש כדי להתמודד עם הבינה המלאכותית, כדי להכניס אותה לשימוש בטוח, נכון, ואפילו יעיל יותר. אז יש הרבה מאוד כלים, אני ככה בקוצר הזמן אני אדבר על שתיים שלוש נקודות.

קודם כל חיזוק היכולת טכנולוגית. גופים פיננסיים הם הגופים שמתבססים על משאבי אנוש, על כוח אדם. עכשיו יש לנו פה גורם חדש, טכנולוגיה, שהיא הולכת ונכנסת לכל סוגי התהליכים. צריך להקדיש לו את המשאבים הדרושים כדי שגם היא תעבוד כמו שצריך בשבילנו, כדי שתהיה יעילה. זה משהו שדורש השקעת משאבים. אז זה העניין של חיזוק היכולת טכנולוגית, זו הנקודה הראשונה. דבר נוסף זה עדכון תהליכי העבודה. יש לנו פה בעצם מנהלת מחלקה שקיבלה עובדת חדשה, בינה מלאכותית. מה היא תעשה עם התוצרים של אותה עובדת חדשה? היא צריכה לדעת להשתמש בהם בתוך תהליכי העבודה הקיימים.

אם היא צריכה להוסיף להם משהו, להוסיף להם מידע, לעבד אותם בצורה מסוימת לפני שהם משתמשת בהם. איך היא יודעת לבקר אותם, איך היא יודעת מה רמת הדיוק שלו, שהעובדת הזאת נתנה לה.

איך היא נותנת לה משוב? כל הדברים האלה, אם אנחנו נדע לפתור אותם, בעזרת כל הכלים שיש, נדע לפתור את הדברים האלה וזו יכולה להיות תחילת הדרך לפתור את הסוגייה הזאת של הקופסא השחורה.

אמיר ברנע

מה את מתכוונת איזה מודל זה?

אביבה וייס

יש הרבה סוגים של מודלים.

אמיר ברנע

מודל זה קשר בין משתנים. מה הקשר? אם את לא יודעת, זו קופסה שחורה, עדיין לא הבנתי מה עומד מאחורי, איך את מפרקת את הקופסה השחורה? גם לצורך בקרה, וגם לצורך בניית המודל.

אביבה וייס

זאת השאלה שאני מפנה גם אל מי שרוצה לאמץ את הטכנולוגיה.

שלומי שוב

רגע, אמיר, אביבה. שרית, את רוצה לעזור?

שרית פלבר

אני מקווה שזו תהיה עזרה טובה. אני חושבת שמה שניתן לעשות בתוך, זה לתת את הקווים המנחים לאיך מאמצים מודל, מה הגבולות הגזרה.

אמיר ברנע

היה מודל, אני שואל, מודל זה ספציפיקציה של משתנים. יש משתנה תלוי, משתנים מסבירים, והקשר ביניהם זה המודל. יש מודל?

אביבה וייס

בוודאי שיש מודל, אתה פשוט לא יודע איך הוא עושה את הקשרים. אתה יודע מה אתה מכניס לו, אבל אתה לא יודע מה ההקשרים.

אמיר ברנע

אבל אם אני אדע, זה משנה מהותית את כל הטכניקה הזאת.

לימדנו בבנקאות, מודלים לאשראי. יש מודל ספציפי, יש משתנה תלוי, תן, אל תיתן, יש משתנים מסבירים, שכנעתי את הלקוח, כל הדברים האלה. פה מדברים איתנו על מודל כשאין מודל. מדברים על מודל שהוא מין קופסה שחורה, שאנחנו לא יודעים מה הדברים שמשפיעים. אם נדע אותם, אז זה יוצא מכללי הבינה המלאכותית. אנחנו מדברים על משהו אחר.

זאת אומרת, איך שם מצד אחד תביא להם קופסה שחורה שבנויה על דאטה, דאטה ענקית, והיא מגיעה לקשרים בתוך הדאטה, אני לא יודע עד כמה הם יציבים, אבל נניח, אבל הם לא ספציפיים. הקשר הוא לא ספציפי. אין נוסחה. תגידו לי אתם. אולי יפה לכתוב את זה במלל ויש לו, אני לא אומר, יש 200 עובדים. אבל אני אומר, אני יודע. אני אומר, מה הצעד הבא? האם יש גם מערכת, בממשק בין מוסף פיננסי ובין לקוח, מערכת שאומרת, המחשב אמר לי, זה המצב. זהו. מצטערים.

מאיה רפפורט

זה לא המצב, כן? יש פתרונות היום. היי, אני מאיה רפפורט, אני מנהלת את תחום הדאטה AI במגדל. רק תשובה לדבר הזה, יש היום פתרונות, אנחנו משתמשים למשל במגדל, לתוך מודל כזה, אנחנו מייצרים מודל של אקספריביליטי, שבעצם יודע לקחת את הפרמטרים העיקריים שהשפיעו על ההחלטה של המודל, נכון, אתה לא מפרט עכשיו את כל ההקשרים והדאטה וזה, אבל אתה לוקח את הפרמטרים המרכזיים ואתה מייצר הסברות עבור כל מודל שאתה מייצר. זאת אומרת שאם אנחנו מייצרים מודל מסוים שהוא machine learning וזה מגיע לסתם human in the loop שיש לו איזשהו קבלת החלטה, יש לו לצד המודל ולצד התוצאה ממש הסבר של למה ה core הזה למשל, הגיע. יש היום פתרונות טכנולוגיים שמשלבים.

אמיר ברנע

רק תסבירי לקהל פה, כשאת מגדירה את המשתנים, כאילו העיקריים, האם יש להם תוכן כלכלי? זה שם המשפחה שלו, או העיר שלו, מה המשתנים שאת מגדירה כמשתנים שמשפיעים על ההחלטה?

מאיה רפפורט

אז קודם כל זה נורא תלוי בסוג המודל, אם זה supervised model או unsupervised model. אנחנו מדברים בעיקר על האנסופרווייזד מודל, זה מודל שהדאטה בסופו של דבר הוא זה שנתן לך את התוצאה, ואנחנו אומרים שזה קופסה שחורה. אבל גם כאן, בניתוח של זה, יש משקולות לפרמטרים שונים שאתה יכול להוציא החוצה ולהגיד אוקיי, אלה הפרמטרים המרכזיים שהשפיעו, כי כשאתה בונה מודל אתה זורק המון המון דאטה, לא כולם משפיעים או לא כולם משפיעים באותן משקולות, ואם אתה משלב לזה עוד מודל שיודע לקרוא את הדאטה ובצורה llm -ית להסביר בדיוק מה הוא רואה, אתה מייצר הסברות למודל.

שלומי שוב

מה את עושה אם יש עכשיו פרמטר לא politically correct?

מאיה רפפורט

אז אצלנו במגדל יש לנו תהליך של פעם בחודש, אנחנו לא מחכים לרגולטור אנחנו מאמצים את המצפן שניתן ולא הכלוב. יש לנו פעם בחודש פורום של ניהול סיכונים, משפטי, פרטיות. כשאנחנו עוברים על כל פעילויות ה AI שיש לנו, אנחנו מסתכלים על ה DATA - ות. שאנחנו משלבים. אם אנחנו רואים שיש משהו שהוא בסיכון להטעיה, ואין לו איזשהו משקל הסבר משמעותי, אנחנו באים ומסתכלים מה עולה ואיזה דאטה אנחנו משתמשים וכולי. יש לנו מודל אחד למשל שהוצאנו פרמטר.

שלומי שוב

העדפת להפסיד כסף, זאת אומרת להיות פחות מדויקת, מאשר ליצור הטעיה.

מאיר לוי

איסור להפלות קיים היום בדין. הקושי הגדול במערכות ה-AI, יהיה לזהות את האפליה וההטיה, הן בגלל האימון של המערכות על דאטה גדול ומשמעותי הכולל באופן מובנה מידע מוטעה ומפלה כלפי אוכלוסיות מסוימות, והן בשל שימוש של המערכת בנתוני פרוקסי שאינם מפלים כשלעצמם, אך שיש להם קורלציה משמעותית עם משתנים מפלים.

שלומי שוב

זאת אומרת, מה שאתה אומר, מאיר, שכנראה לא תהיה פה יכולת אכיפה.

מאיר לוי

לא, מה שאני אומר זה שהאתגר יהיה לייצר את יכולת האכיפה ולהתמודד למשל עם סוגית הקופסה השחורה. יהיה צורך לנקוט בשורה של בקורות ואמצעים לזיהוי ומניעת הטיות לכל אורך חיי המערכת, ובפרט ביחס לתוצאות שאליהן מגיעה המערכת. אנחנו כבר מכירים הרי סיפורים שקרו בארצות הברית, שהייתה אפליה בין בני זוג שביקשו אשראי. והם היו באותו מצב כלכלי, והאישה קיבלה אשראי בתנאים הרבה פחות טובים מהגבר, והמערכת טענה, אנחנו לא עושים אפליה, כי אנחנו לא שואלים בכלל מהו המגדר. אבל זה ברור שהמערכת התבססה על נתוני פרוקסי שהיתה להם קורלציה למשתנה המגדר.

שלומי שוב

תהיה להערכתך עמידה באתגר?

מאיר לוי

אני חושב שכן. ראשית, אני חושב שתהיה התמקדות משמעותית בבקרה תוצאתית ויהיה שימוש בכלי AI כדי לזהות מקרים של אפליה בפועל. שנית, אני חושב שבסופו של דבר הרגולציה, ככל שהיא מתקדם וככל שהמערכות יתקדמו, תדע לדרוש דרישות ולהפעיל בקורות מוגברות על פעילות המערכות באופן שוטף, גם בשלב ההטמעה והיישום, במטרה להתמודד עם הסיכון לאפליה, כולל דרישה להסברתיות במקומות המתאימים כדי לוודא שהמערכת לא התבססה על קריטריון זהותי מוגן באופן מפלה.

משה ברקת

אני חושב שהבעיה היא לא אפליה, אלא הבעיה המרכזית היא זמינות המוצר או השירות לאוכלוסיות מסוימות. המערכות הללו ידעו בדיוק את הפרופיל של מי שמקבל שירות, ולכן המצב יכול להיות, למשל בביטוח, שזה כבר לא יהיה ביטוח, זה יהיה פשוט scheme של מימון, מכיוון שאני יודע לזהות בדיוק לגבי כל אדם מה הסיכון שלו. עכשיו אם אני לא רוצה לבטח אדם עם סיכון גבוה זו אפליה? זולא אפליה.

הבעיה תהיה של זמינות: של אשראי, של ביטוח, לגורמים כאלה שהמערכת מזהה אותם ככאלה שאני לא רוצה לעשות איתם עסקים, משום שהם מסוכנים או מסיבות אחרות.

שלומי שוב

אנחנו רוצים להתקדם בדיון. שלי נשמח אם תיתן לנו את הזווית האסטרטגית בקצרה, כדי שנפתח את הדיון.

שלי תשובה

בוקר טוב לכולם, שמי שלי תשובה. אני שותף בדלוייט. כמה דקות באמת על זווית שאנחנו מסתכלים עליה. היה לנו מזל גדול להצטרף לקבוצת מחקר גלובלית בדלוייט שעוסקת ב-AI, תחת הקונספט של NEO, next evolutionary organization, כשהסתכלות היא גם ברמה הרוחבית וגם ברמה הענפית. ואני חייב להגיד שאני בדיון הקצר שהתפתח בין שלומי ועמית בפתחה, הרבה יותר קרוב לגישה של שלומי, ואני מאמין גדול מאוד במהפכה שאנו הולכים לעבור. אבל אני גם מאמין מאוד גדול ב-Hype Cycle של גרטנר, שלמי שלא מכיר אותו בא ואומר שכל חדשנות טכנולוגית עוברת את התהליך הזה שבו מגיעים ל-Hype ואז נופלים לאותו תהום של התפכחות, ורק אז עולים בהדרגה אל הצמיחה הנורמטיבית.

ואני חושב שאנו נמצאים היום בעיצומה של הצניחה לתוך התהום בהיבטים של ההייפ. ולכן אני חושב שאנחנו הולכים להיכנס לתקופה מאוד מאוד מעניינת,

בשנת 24 ובתחילת 25 היה היה קונצנזוס רוחבי בארגונים, שכולם חייבים לנסות פרויקט AI ולא היה מנכ"ל שלא שחרר לצורך כך תקציבים. במהלך 2025 אנו רואים את ה"אופוזיציה" מרימה ראש. ניצנים לכך אפשר לראות במחקר הגדול של MIT, ש-95 אחוז מהפרויקטים לא מראים את ה-ROI הנדרש, ושהמעבר מ-POC לייצור נתפס כקשה הרבה יותר. אני חושב שזאת הולכת להיות העונה הכי מעניינת בהיבטים העסקיים, כי קבלת ההחלטות העסקיות צריכה להיות הרבה יותר ממוקדת ומבוססת על תועלות, ולא רק על ההיפותזה הראשונית שהעולם משתנה. זאת נקודה אחת שרציתי לחדד. בטרמינולוגיה, דיברנו עליה, אני חושב שכל העולם של הצי'אט בוטים הוא כבר, אנחנו בתוכו, לא רוצה להגיד מאחורינו, אבל לכולנו ברור האימפקט שלו.

תחום ה-אייג'נטים הוא הנושא העכשווי. אייג'נט זה אותו מודל AI שמחליף תפקיד מסוים, ואנחנו רואים את הניצנים של זה. בדלוייט יש כ-100 איש שכותבים אייג'נטים לחברות אמריקאיות, לכן מי שתוהה האם זה אמיתי, אז אני יכול להגיד שזה אמיתי לחלוטין. יש אייג'נטים שפועלים בחברות אמריקאיות, ועושים תפקידים שאנשים עשו בעבר.

שלומי שוב

100 אנשים בישראל?

שלי תשובה

כן, מפתחים אייג'נטים. הדיון הפילוסופי הגדול הוא מסביב ל-AGI מודל הבינה המלאכותית שמחליפה אותנו כבני אדם, ועל זה אני לא אתעכב עכשיו, כי זה לא באמת הדיון שכאן. יש פה תפיסה מעניינת שאפשר להסתכל עליה, בעיקר על הרובד למטה, וזה אומר איפה ה-AI נמצא מולנו כאנשים. ואם התחלנו בהתחלה כבן אדם שנעזר ב-AI למהלכים מסוימים, לבדיקות מסוימות, וזה מה שאנחנו רואים היום כתופעה רוחבית, אנחנו רואים את מה שקרוי פה Cognitive Enhancement, ה-AI הוא דרייבר, הוא עושה פה, לאחר מכן הוא יעשה תפקיד, זה העולם של ה-Agents, ואנחנו כבר מדברים על Agent of agents. זאת אומרת, אנחנו מדברים כבר על התפיסה שיש יכולות AI שמנהלות למעשה מחלקה, עד כדי זה שבאמת הארגון יהיה ארגון אוטונומי, כמו שהרכב הוא אוטונומי. בפילוסופיה הזאת אפשר לנהל דיונים שלמים איפה זה נעצר.

אולי הנקודה האחרונה שאני רוצה להגיד עליה לפני שנגיד משהו על הצד הפיננסי, זה שלתפיסתנו, והרבה פעמים ארגונים טועים בזה, החשיבה לא צריכה לבוא רק מהכיוון האסטרטגי, איזה יוזקייסים אני עושה וכו', אלא לתת משקל שווה לכל שלושת המרכיבים האחרים.

כאן יושבים הרגולטורים ולכן מבחינתנו כל עולם המשילות הוא קריטי ומה שאתם עושים פה הוא עבודה מבורכת מכיוון שאנחנו צריכים את לדעת מהם הכללים שאנחנו עובדים עם לקוחות כדי לבנות את זה נכון,

הטאלנט הוא הרובד השלישי, וכמובן כל הצד הטכנולוגי. זאת אומרת, אם כל ארבעת הדברים האלה לא שלובים, אנחנו נמצאים בבעיה.

ובהיבט הבנקאי, או כלל המגזר הפיננסי קיימים אין סוף דוגמאות ל-use cases. אני רוצה להראות דוגמה אחת, בסופו של דבר, היתרון שלנו שאנחנו יכולים לעבור על הרבה מאוד יוזקייסים בכל מיני מקומות בעולם, אז יש לנו מאגר גלובלי שיש בו היום 200 יוזקייסים הכוונה לא 200 רעיונות שאנשים העלו, אלא 200 פתרונות שנמצאים או בתהליך בנייה או שכבר עלו לאוויר, לכל אורך שרשרת העולם הבנקאי, כמובן יש שרשרת דומה בביטוח. הדברים המרכזיים שעושים, וזה נאמר מקודם, אני חושב שבשנתיים שלוש הקרובות אלה יהיו היוזקייסים המרכזיים, זה היוזקייסים שמחברים אנשים ומכונות. אנחנו קוראים לזה "the age of with" - "העבודה המשותפת", אבל אפשר לראות שהרבה מאוד מהיוזקייסים האלה, הם יוזקייסים שבהם המודלים כן מביאים הרבה מאוד ערך, אבל בקצה יושב בן אדם, והבן אדם הזה לוקח את האחריות. ואז אני חושב שהסיכונים יורדים גם בהיבט הרגולטורי וגם בהיבט הפנימי של אותו גוף. החומר נשלח אליכם, לכן אפשר יהיה לעבור, אני לא אעבור עכשיו.

שלומי שוב

בואו נשמע כעת את הגופים הפיננסיים עצמם.

יאייר קפלן

שלום, שמי יאייר קפלן, אני מנכ"ל בנק ירושלים. בקצרה, בנק ירושלים הוא בנק לא מחמשת הקבוצות הבנקאיות, אלא הבנק העצמאי השישי בגודלו, ולמעשה תסתכלו על זה, גם היחיד שנשאר אחרי שבעבר היו עשרות בנקים. והסיבה שבנק ירושלים מצליח ונשאר וגם מתפתח זה כי הוא נישתי, סקטוריאלי והוא מוצרי. כלומר הוא יותר מתמקד בניהויות מסוימות באוכלוסייה, הוא מתמקד יותר בסקטורים ספציפיים ובמוצרים, ולא דווקא בעו"ש, הוא עושה דברים שהם ברוב המקרים מורכבים יותר.

עכשיו, כשאני מסתכל על AI, אז אני מסתכל על זה בשלושה מישורים.

אולי אני קצת אשלב בין מה ששלומי אמר למה שאמיר, אני מסכים עם שניהם, אני חושב שהאמת היא באמצע. שלושת המישורים אחד זה סיכון, אז AI זה דבר מסוכן. ובעצם מה שהדו"ח הזה אומר, זה שהאחריות היא עלינו, על המנכ"לים. כי אוקיי, אפשר לעשות הרבה דברים, ולבנות מעטפות של בקורות וביקורות וגופים והכל, אבל בסוף מי שלוקח את האחריות זה ההנהלה, ואין הגנה על זה.

כשאני מסתכל על הדו"ח, למעשה, אם אני מקבל החלטה לא נכונה, אם זו אפליית גיל או מין או משהו שאני לא יודע להסביר, או שאני עושה איזושהי טעות רגולטורית כי המכונה או הקופסה השחורה נתנה איזה החלטה, אז הסיכון הוא על הבנק, וצריך להביא את זה בחשבון.

שלומי שוב

צריך ביטוח AI, כמו ביטוח סייבר.

יאייר קפלן

זה לא כזה מצחיק, אני בחודש הבא נמצא בלונדון, הולך לדבר עם המבטחים שלנו על ביטוחי AI, בדיוק ככה. בביטוח הבנקאי אני חושב שצריך להסתכל גם על סיכוני AI, בהחלט, ביחד עם סיכון הסייבר וסיכון המעילות וההונאות, נכנסנו גם לסיכון הזה. הדבר השני זה העלויות. AI זה דבר יקר. מי שחושב ש AI מוזיל, אז לא, קודם כל הוא מעלה את העלויות, הוא לא מוריד אותן. צריך לבנות תשתיות, תשתיות של ענן, תשתיות של מודלים, להתחבר copilot מי שצריך, או משהו אחר. להטמיע את העננים, להטמיע את מודלי השפה, אנחנו עובדים עם שלושה מודלים ChatGPT, Gemini, Claude ו Claude אנחנו צריכים להטמיע אותם. קודם כל זה עולה כסף, אחרי זה נראה מה עושים. והדבר השני, זה באמת הפוטנציאל לצמיחה. דווקא אנחנו כבנק קטן, אם אנחנו יכולים להשתמש ביכולות של AI לאו דווקא להתייעל, אלא להגדיל את הצמיחה שלנו, ולחפות בכך על הסיכונים שלנו ועל החסרונות לגודל, ואז אני מסתכל על זה בשני מישורים אחד זה מכפיל כוח והשני זה מכפיל ידע כלומר, בשימוש ב AI, אני אחד, יכול להכפיל ידע, ברור, כל הידע שבעולם נמצא עכשיו גם אצלי, שזה דבר מדהים, קצת משתלב עם הבנקאות הפתוחה, והדבר השני, זה באמת מכפיל כוח, אני יכול אולי, דרך שימוש באייג'נטים או בכל מיני פתרונות אחרים, לעבוד מעבר למגבלות הגודל שלי, שזה דבר

שאנחנו בוחנים, ואז אנחנו בעצם בונים כרגע, מפת דרכים ל-AI, כמובן שעובדים על אסטרטגיית AI, כל המערכת הבנקאית עובדת על אסטרטגיית AI לא רק אנחנו.

אמיר ברנע

הייתה לך בעיה מול הרגולציה?

יאיר קפלן

גם הרגולציה כותבת את עצמה עכשיו. יש היבטים שמכוסים אבל אין עדיין נב"ת ייחודי ל-AI במערכת הבנקאית.

שלומי שוב

הפיקוח על הבנקים היה חלק בדו"ח, הוא גם פרסם קווים כלליים.

יאיר קפלן

אבל זה קווים מנחים, אין עדיין נב"ת AI, יש קווים מנחים כמו שאמרתם, אבל אין נוהל בנקאי תקין מחייב. אז כמובן שאנחנו מסתכלים על יישום AI ודאטה, והיום אפשר לדבר עם הדאטה בייס לא צריך לעשות SQL או שאילתות אנחנו פשוט יכולים לדבר עם הדאטה בייס, שואלים אותו בשפה פשוטה שכל אחד מבין והוא רץ על כל הדאטה בייס של הבנק ומביא תוצאות זה דבר מדהים.

הדבר השני כמובן זה אוטומציה. אם פעם השתמשנו ברובוטים, בוטים או אר.פי.איי, אז היום זה עובד בצורה של אייג'נטים, לדוג' בשירות לקוחות כמו שכולכם מכירים, כמובן.

השלישי זה פיתוח, רוב הפיתוח בעולם יעבור ל-AI אין סיבה היום באמת לפתח לבד, אתם מכירים את זה, ה-AI היום מפתח כבר ברמה של דוקטור, אולי זה יהיה פרופסור אני לא מהאוניברסיטה אז אני לא יודע מה יש אחרי.... וכמובן שהיישומים הכי פשוטים זה דוקימנט AI. כלומר, לקחת את היישומים התפעוליים שיש בכל בנק המון, יש המון המון תפעול בנקאי, ולהעביר אותו ל-AI זה הדבר הכי מתבקש.

אמיר ברנע

תסביר לי משהו. על כל הבנקים, מקורות המידע הזה שאתה מדבר עליהם, הגלובליים, הם לרשות כל אחד. איך אתה יוצר את הייחודיות של בנק ירושלים? בהקשר להחלטות שלך, כאשר אתה משתמש במערכות אוטומטיות כאלה, מערכות ידע שקיימות לכולם. איפה הייחוד של בנק ירושלים? איפה השיקולים שלך? זאת אומרת, איפה הייחודיות שלך? אחרת זה מגיע למצב סטנדרטי של מוצר אחיד.

יאיר קפלן

אז אמרתי שאני לא בצד של שלומי ולא בצד שלך, אני באמצע. אני לא חושב שה-AI יחליף את האנשים, ואני לא בטוח שאני אתייעל מאוד. בסוף הידע הייחודי הוא כן עדיין אנושי, אבל תמיד הידע נשען על

איזשהו דאטה ועל איזשהם מודלים. אז גם אם מודל החיתום שלי יהיה AI, או מודל קבלת ההחלטות, עדיין תמיד תישאר הזווית האנושית, עם הניסיון האנושי וההיכרות בעיניים, ואני מהאולד פאשן שעדיין מעדיפים אולי את החיתום ואת מדרגות האשראי בעצמי, יודע שיש גם לניסיון הנצבר וגם למדרגי הניהול הרבה מאוד ערך, גם אם מודל הדירוג הוא AI אז יש תמיד שילוב של בקרה אנושית. אגב גם לפני זה זו הייתה רגרסיה לינארית שבנויה על דאטה, יהיו מודלים יותר מתוחכמים ויותר חכמים, ויהיה יותר דאטה, אבל המשקל האנושי תמיד יכריע.

שלומי שוב

ההחלטה בסוף היא אנושית, אני לא אמרתי משהו אחר. בואו נשמע כעת גם צד של בנק גדול.

מאיה רביע

אהלן, אני מאיה, מנהלת אסטרטגיה בבנק לאומי. אני לא אחזור על הרבה דברים שנאמרו, אנחנו בגישה, שמסתכלת בצורה מאוד יישומית על נושא הAI. מבחינתנו, הAI הוא לא עומד בפני עצמו, הוא לא המטרה, הוא אנייבלר. הוא מאפשר לנו כבנק להשיג את היעדים האסטרטגיים. וכמובן לצד זה אנחנו גם בודקים איפה הוא מעמיד בסיכון את היכולת שלנו לעמוד בתוכנית העסקית שלנו, מה הם האזורים, מה אנחנו צופים שיגיע, איך אנחנו רואים באמת את הצרכים של הלקוחות, ואת הזירה התחרותית משתנה בעקבות כניסה של פתרונות AI. אבל אנחנו בעיקר עוסקים, לפחות בתוך הבנק, באסטרטגיה, במה שאנחנו יודעים, פחות במה שאנחנו לא יודעים. ברור לנו שזו מהפכה מאוד גדולה שתשפיע מאוד על הצרכנים כמעט בכל היבט אפשרי. נעשתה השוואה לבנקאות פתוחה, אז אנחנו מבינים שבנקאות פתוחה זה בעולם הבנקאות, שהוא מאוד מעניין את הלקוחות, אבל עדיין בפרופורציות אחרות מאשר כל הסקטורים האחרים ודרך חיים שלהם. ולכן כשאנחנו מדברים על משהו כזה, זה קצת יותר כמו להשוות למהפכת האינטרנט או מהפכת התחבורה, סמראטפונים ודברים הרבה יותר מהפכניים וגדולים שהם הרבה מעבר לתחום הבנקאות. ובהקשר הזה, ברור לנו שיהיו לזה אדוות שחלקן, כמו שלא היה אפשר לחזות במהפכות שציינתי, גם לא כל כך ניתן לחזות היום. עכשיו אנחנו לא מאוד מרחיקי לכת ונסתכל עוד 10 שנים מהיום איפה נהיה, אלא מסתכלים בצורה יישומית כרגע איך אנחנו יכולים למנף את הAI, כדי בסופו של דבר, בשורה התחתונה, לשמר יתרון תחרותי ולפגוש את הלקוחות שלנו והציפיות שלהם לשירות מהיר, מדויק ועדכני מהבנק.

ואז כל היתר, לנו כן יש מטרות על יעול תהליכים, זה כן משהו שהוא KPI, אבל הם סייד אפקטס, בסוף מהיכולת לפגוש את הלקוחות בציפיות שלהם.

שדרך אגב, בעינינו, לתפיסתנו זה באותם אזורים שכבר יש ציפיות היום. זה פשוט קטליזטור לציפיות הרבה יותר מואצות. כבר היום יש ציפייה ללקוחות שאנחנו נהיה פרסונליים, שאנחנו ניתן להם כמעט את כל, לא כמעט, את כל השירותים הבנקאיים בזום, בדיגיטל, בצורה פשוטה, בשני קליקים ואפילו שניים כבר טו מאצ'. לפגוש אותם בשירות 24/7, לפגוש אותם איפה שנוח להם לפגוש אותנו ולא איפה שלנו נוח לפגוש אותם.

אני לא יכולה להגיד שלגמרי הבנקים נמצאים שם כבר היום, אבל בסוף מבחינתנו AI זה קטליזטור לציפיות. הציפייה היא כבר ליותר מהר, יותר פשוט ויותר דיגיטלי. זה בכלל לא שאלה של אם נכון להיות שם או לא, זה הכרח.

שלומי שוב

מאיה, אני רוצה רגע לפרש את מה שאת אומרת, בגדול, את אומרת שיש פה משהו דרמטי, אבל את כרגע עם העיניים על הכדור?

מאיה רביע

חד משמעית, אנחנו בכל זאת כן מצמידים ROI לפרויקטים. אבל כן אני אגיד שבהקשר העיניים על הכדור, ברור לנו ואני חושבת שאני מתחברת למה שכבר נאמר, שבסוף נדרש פה להשקיע בגדול. רמת ההשקעה, בשביל שכל יוזקייס קטן כגדול יצליח, היא מאוד גדולה. צריך לסדר את הדאטה, צריך לחבר תשתיות, צריך לתפור את הפתרונות גם במערכות הליבה המאוד מורכבות של גוף גדול ומסורתי. אז כן נדרשת פה איזושהי השקעה מאוד גדולה, אפרופו גם רגע היכולת קטנים מול גדולים.

שלומי שוב

בדיוק, איפה את רואה את החסרונות? בטווח הזה של יתרונות חסרונות

מאיה רביע

עוד פעם אני אגיד, ברור שההשקעה היא מאוד מאוד גבוהה ולכן פה זה יתרון לשחקנים הגדולים. שיש להם את היכולת להכיל השקעות מאוד גדולות. מצד שני, וכפי שציננת בעצמך בדברים שלך, את כמות הדאטה והמידע. אני אגיד מצד שני, בסוף הכל יקום וייפול על יכולת האקסקיושן. בתיאוריה, אנחנו יודעים כבר היום יש לבנקים הרבה מאוד מידע, השאלה כמה הם באמת משתמשים בזה, כמה הם פרסונליים. AI הוא קטליזטור, אבל היכולת הייתה קיימת גם לפני, המידע עצמו הוא לא חדש. אז נכון שזה מנגיש, זה מסדר, זה מאפשר, זה הרבה יותר מהיר. זה מקצר תהליכים, כמו שאמרתי, לתפור הרבה מערכות ברקע, את כל החוטים, אבל עדיין, אם בסוף הוא ניוון מספאם, הוא יוציא החוצה ספאם. אז זה לא חוסך את עבודת התשתית שנדרשת לעשות.

אמיר ברנע

אני רוצה לשאול אותך משהו כל המערך הזה של הבנק עם AI, עבר את הדירקטוריון? הדירקטוריון טיפל בנושא הזה וקבע את ההנחיות? איך זה מתבצע פה, אומרים, הדרג הנמוך, הדרג הגבוה, איך זה היה בבנק?

מאיה רביע

יש לנו אסטרטגיית AI מאושרת דירקטוריון. ספציפית בלאומי הנושא נלקח מאוד בחשיבות ותשומת לב, הוקם מרכז AI ייחודי אני חושבת במערכת הבנקאית. זו יחידה ייעודית, כך שהחשיבות שהבנק

מקדיש לנושא ברורה. הנושא מנוהל ברמות הכי גבוהות שיש ומפוקח באופן הולם על ידי הדירקטוריון.

שלומית ווגמן

בימים אלו ארגונים מתחילים להקים פונקציה שנקראת Chief AI Officer בחברות, שבעצם יוצרים את המבנה הזה כדי להציג אותו לדירקטוריון. אנחנו מלמדים כאן ברייכמן קורס כזה להכשרה בתחום.

שלומי שוב

מאיר, לך יש זווית מעניינת כי אתה מגיע מהטכנולוגיה, כסמנכ"ל טכנולוגיה במזרחי טפחות, אני אשמח לשמוע את הזווית שלך, על מה שאנחנו מדברים.

מאיר אהרוני

קודם כל נעים מאוד, מאיר אהרוני מבנק מזרחי טפחות. תראו, אני חושב שיש כאן איזשהו בלבול, שהולך ומתעצם ככל שעולים בדרגים בארגון בגלל השימוש בילה מודל, מבלבלים בין מודל שפה לבין מודל נטישה או מודל חיתום, וכולם רואים בעיני רוחם אייג'נט שיחליט האם אתה תקבל ממני אשראי, כן או לא? רגע, כמי שעושה את זה בפועל, זה לא עובד ככה, זה רחוק מזה. אני אתן דוגמה מעשית כדי להבין את המורכבות. יש לי משכנתא, ופתאום קפץ לי גובה החיוב בצורה משמעותית, ואני רוצה לשאול למה. אם הציפייה כאן, ש אצור קשר טלפוני, או אפנה ל לערוץ דיגיטלי, אולי אפילו דרך אייג'נט, ואני אקבל תשובה פשוט על סמך הדאטה שלי, בנושא, אז בחלומות הלילה. אני מזמין אתכם לקחת את כל המידע שלכם מהבנק שלכם, לשים אותו על אקסל, אפילו לטרוח ולהסביר לגבי כל שדה מה הוא אומר, ולהעלות את זה לציט ג'יפיטי, ג'ימיני, דיפסיק, כל אחד מה שהוא מעדיף, ולשאול למה עלה לי גובה החיוב. אם התשובה תהיה תשובה - א' נכונה, ב' מלאה וג' עקבית, אנחנו בבנק נממן לכם את מיזמי ה-AI לשנה שלמה.

זה לא עובד ככה, בסדר? כדי שזה יעבוד ככה, צריך לטרוח מאוד. אמרתם אייג'נט אופ אפייג'נטס, נשמע מאוד מאוד מרשים, כבר נמצא תפקיד במשאבי אנוש של מי שמגייס ומעסיק ומשמר סוכנים נכון? אבל בסוף על מה מדובר? בסוף יש פה איזושהי אוטומציה מתקדמת, שנשענת על יכולות שלא היו שם קודם. לדוגמה, אני עושה בזה נפלאות אגב בפיתוח, זה עולם מאוד מאוד תחום, אתה פשוט לקחת איזושהי יכולת, אימנת אותה על המון המון טקסט רלוונטי, ונאמרו פה SQL, קובול, דוטנט, אנגולר, תבחרו מה. כבר לא קשה לעזור למפתח לעשות הרבה מאוד עבודה שחורה שהוא היה עושה, ואז גם לבדוק אותה. דוגמה נוספת - אוטומציה ארגונית, לנתב פנייה כתובה מלקוח למקום הנכון בארגון כדי לקצר זמן של בק אופיס Slai. אז מה עושים? איך זה נראה מאחורי האייג'נט הזה? יש את האייג'נט הראשון שהתפקיד שלו לקבל פנייה כתובה, להבין מה כנראה היא אומרת, זה כבר לא קשה, כולנו ראינו שזה עובד. נמצא שם איזושהו טקסט, תעביר בבקשה כסף מחשבוני לחשבון אחר של הבן שלי, לבנק אחר, אני גם ארשום את זה בצורה כזאת. מה שיעשה אותו האייג'נט שכבר הבינ שכנראה

מדובר בבקשת העברה, הוא יפעיל כלי נוסף, אייג'נט נוסף, שהתפקיד שלו זה להשלים מידע מובנה מהארגון שלי, בסדר? שישלים מה זה חשבונני, מי זה הבן שלי, וכיוצא בזה. אייג'נט להשלמת המידע.

ובסוף, אם אין פה בעיה של סמכויות, ואם אין פה חשבון שהוא באיזושהי בעיה משפטית, או שאין יתרה מספקת בחשבון, הדבר הזה יומר לאיזושהי אוטומציה שתקרה. אבל היינו צריכים לעבוד המון כדי שהדבר הזה יקרה. אין פה איזושהו קסם במסגרתו יתרנו את העובדים בארגון ופשוט עשינו את זה, צריך לקחת איזושהו עולם תוכן ולוודא שהוא עובד. לעומת זאת כשמדובר בפרונט אופיס, באינטראקציה עם לקוח בקצה, זה עולם הרבה יותר מסוכן. אני אסביר למה, ולכל מי שפה, פשוט יש המון אנשים מעולם הבנקאות, אז הם יתחברו לדוגמה. ב 2019 החל לפעול מאגר שנקרא מאגר נתוני אשראי, כולם מכירים, נכון? עשה דרמה אגב בתחום הזה. אבל, אם הייתם רואים, בטח בשלבים הראשונים וגם היום, איזה שטויות הגורמים שהיו צריכים לדווח למאגר הזה דיווחו אליו, מה שנקרא זר לא יבין זאת. למה? כי המידע שזורם בעורקים של ארגון, במערכות התפעוליות שלו, בטח בארגונים הגדולים לא נועד לתמוך תשובות חשופות unfiltered לצד ג'. הם נועדו לצורך מסוים, המידע נשמר בצורה מסוימת, ואין ערובה לאיך הוא יראה כשהוא יצא ככה בצורה גולמית.

ולכן לפני שאתה שם מישו שעונה תשובה, שצריך לעמוד מאחוריה, אני לא מדבר על החלטה על מתן אשראי בכלל, שאלה שרירותיות פשוטה, כדאי שתוודא שהמידע שלך שמור בצורה נבונה, שתהליך קבלת ההחלטות של אותו אייג'נט היה תהליך נכון והתשובה שניתנה היא תשובה סדורה, עוד לפני עניין ההסברות, בסדר? כדי שלא תביך את עצמך. עכשיו, נאמר פה עוד איזה משהו על איך יראה העולם, ריבוי סוכנים, סוכן אחד. אז גם כאן הייתי שם בפרופורציה. אני לא מבוגר מדי בגיל, אבל אני זוכר לדוגמה, כשנהיו אדונים חדשים לכל עולם המסחר, פייסבוק, גוגל, הייתי בטוח שזהו, כל הנדל"ן המסחרי, הקניונים ושטחי המסחר, זהו, אבד להם, נכון? יש ערוצי הפצה חדשים, יש אדונים חדשים לכל העולם הזה, למה לשלם שכירות בקניון, נכון? קודם כל תבואו לקניון ליד עיר מגורי ביום שישי, אי אפשר לחנות ואי אפשר להכניס סיכה.

אבל בגדול, זה מה שאתה אמרת, זה עניין ההייפ. אחרי שיורד ההייפ ופוגשים מציאות, אז מבינים שגם כאן נכנסים אינטרסים, וההצעה הכי טובה שלי לכולם זה תלכו אחר הכסף. בסוף אף אחד לא עושה שום דבר לשם שמייס, גוגל לא נותנת מידע לשם שמייס, ולא מקדמת תכנים לשם שמייס, אותו דבר פייסבוק, וכך גם כל ערוץ אחר. תשאלו כל עסק קטן שנאנק למצוא לקוחות ולמכור, העלות הכי גדולה שלו היא עלות ההמרה, הרכשת הלקוח והמכר. יש לנו שחקנים חדשים שלוקחים המון המון כסף. סביר, סביר, בסדר? למרות גילי הצעיר ולמרות שאינני נביא, שגם כאן יקרה משהו דומה. בסדר? בסוף, הדבר הזה עולה המון המון כסף, עולה המון חשמל, המון נדל"ן, ואנשים שמו שם איזושהם מתווכים סופר חכמים שיודעים לקצר ולעשות המון המון פעולות במקומך, סביר שהם יחפשו את המודל המסחרי שלהם.

שלומי שוב

איך הרגולציה צריכה לפקח על הנושא הזה? יש לה יכולת להיכנס לצורך העניין לתוך הקופסא השחורה?

מאיר אהרונ

קודם כל, לנו אין יכולת, בתור התחלה. שאלו פה הרבה על אשראי ועל חיתום. אני בא מחברת כרטיסי האשראי. אני בא ממקס כמו גם מאיה הנהדרת שדיברה לפני, וכמה שאנחנו עבדנו על לייצר מודל חיתום חדש, ולתקף אותו, ולראות שהוא עושה מה שהתכוונו, ושהדיפולטים שאנחנו חוטפים אחר כך, הם משהו שהתכוונו אליו, או שלא התכוונו אליו, זה תהליך של חודשים רבים. זאת אומרת, הציפייה שישב שם איזשהו יצור ויקבל און דה פליי החלטה על סמך הסושיאל מדיה ודברים שהוא חש באווירה של השיחה, אז בחלומות הלילה, זה רחוק משם. עכשיו לגבי הרגולטור, הוא כן עשה פה משהו שהוא מאוד מאוד בוגר. באיזה אופן? הוא לא בא להקשות או להיות דווקאי, הוא אומר בואו נתקדם בגישה, מבוססת סיכון. כמו עם ענן אגב, אנחנו מאפשרים לכם, אבל תנהגו באחריות. מדירקטוריון ומטה, תראו מה אתם עושים, בסוף אתם אחראים לתוצאה, אנחנו לא מגבילים אתכם אבל אתם אחראים לתוצאה. אני חושב שזה לוקח גופים כמונו, שהם שמרנים בסופו של דבר, וגורם לנו לנהוג באחריות.

שלומי שוב

אוקיי, תודה. אני רוצה, אורן, לשמוע אותך עכשיו. הזווית שלך מאוד מעניינת, כי אנחנו מדברים קצת באוויר, ואתה מגיע מתוך השטח, מתוך הפרקטיקה של סוכני הביטוח.

אורן כהן

בוקר טוב, שמי אורן כהן, אני מנהל את הסוכנויות של הפניקס. בסופו של דבר, בואו אני ארגיע את כל האפוקליפסה שיש פה לעניין הפיטורים והצמצומים. עדיף לעשות הרבה יותר עסקים באמצעות הטכנולוגיה מאשר לפטר אנשים, זה מה שקורה אצלנו בארגון, אנחנו רק צומחים בכוח האדם, תוך כדי שימוש בטכנולוגיות הקיימות, ויש טכנולוגיות קיימות.

בעסק הזה יש הרבה סיכונים לצד ההזדמנויות, ויש פה אחריות מאוד גדולה. האם אני יכול להגיע לרשות שוק ההון ולהגיד לה, כמו שעכשיו אתה התחלת בדברים ואמרת שאנחנו קנינו חברה, אנחנו הופכים אותה להיות סוכנות ביטוח דיגיטלית, AI, אוטאר, שממש מדבר עם מר ברנע בערב או בלילה, מתי שיבוא לו, אבל האם אני יכול לפנות לרשות שוק ההון ולבקש רישיון? עבור גברת "ללי", אנחנו קראנו לה "ללי", אני לא יכול לפנות לרשות שוק ההון, להגיד לה, יש לי פה איזו אייג'נטית חמודה, קוראים לה "ללי", בואי תני לה רישיון. הרשות תגיד, אדון אורן הנחמד, זה בסדר, הרישיון הוא לך באופן אישי על התעודת זהות שלך, אתה אחראי ל"ללי".

וזה מתחבר לשאר הדברים שנאמרים פה, שבסוף האחריות פה היא קריטית. כלומר, והם עושים נכון, הם אומרים, חברים, אני לא יודע מה אירופה, מה אמריקה, אתה רוצה AI? אתה רוצה להשתמש בטכנולוגיה, אתה רוצה לשפר את השירות, אתה רוצה למכור יותר, אתה רוצה לגדול, אתה רוצה להתפתח, אתה רוצה להיות בקדמה, קח אחריות. אז הצעדים הם הרבה יותר איטיים, שמרניים, כי אתה חושש, לא רק מהרגולציה, אלא לעשות נזק. כי בסוף בעולם הביטוח, והוא עולם מאוד מסובך, אני לא יודע לגבי הבנקאות, אבל בעולם הביטוח, הצמתים הם אין סופיים. בסוף הנושא הזה של לקחת את הרובוט, לאתר את הלקוח למוצר הספציפי שאתה רוצה לשווק לו, ולאתר את הצרכים הספציפיים של אותו לקוח, כי זאת החובה הרגולטורית שלנו - אני לא יכול למכור פוליסה סתם בלי לאתר צרכים. אחרי זה למכור לו את הפוליסה, ואחרי זה ללוות אותו בכל עולם השירות והתביעות, ואחרי זה גם לחדש לו את הפוליסה ולהציע לו הצעת ערך נוספת, כל האירוע הזה, זה נשמע מאוד מאוד פשוט, ו-AI וטכנולוגיה, אבל אם אין מישהו מאחורה שבדק כל תהליך כזה, ועושה עליו סוג של וי, אני אומר לכם, יקבל מכתב מרשות שוק ההון, יגיד לי אדוני קיבלנו פה, 12 אלף תלונות על הללי הזאת. כי זה, זה הסיפור האמיתי, כי אין כרגע טכנולוגיה תומכת בכל החלום היפה הזה וכו'. יכול מאוד להיות שעוד 5 שנים, 10 שנים, המחשב החדש, הקוונטי, קראתי עליו, זה נשמע לי משהו שהוא בכלל יוצא דופן ובכלל לא קשור למה שאנחנו היום מדברים עליו, כי הכל יהיה הרבה יותר מהיר. אבל ההזדמנויות הן מאוד גדולות.

לקחת את המוצרים הקיימים, להטמיע אותם, לארוז אותם. אבל, כל זה לצד ה-HUMAN, לצד האנשים שיובילו את התהליך הזה ולא רק ילכו הביתה בעקבות הכניסה של אוטאר ללי, אלא ילוו אותה בתוך התהליך, ואם אתה רוצה גם להרוויח מזה כסף, אתה צריך לעשות את זה במסות הרבה יותר גדולות. וכמי שחורט על דגלו להוביל את שוק ההפצה, והוא השוק באמת,

שמתאים לעולם AI בצורה דרמטית, כי הפצה זה לא רק הפצה, בוא יש לי מוצר של מגדל או הפניקס או כלל, ולבוא ולמכור אותו לשלומי. הוא פשוט לוקח אותך לכיוונים אחרים לגמרי. זה מהצד של הצרכים והשירות, התביעות וכו'. זה לא עולם כזה מסובך אם היום עובדת מתוך 1,600 עובדים שלנו, היא עושה לדוגמה 100 חידושים, אז באמצעות ה-AI, באמצעות ללי, היא תעשה 300 או 250.

שלי תשובה

אז זה לא ייתר את חלק מהאנשים?

אורן כהן

זה לא. כי אתה תגדיל את הפעילות ואתה תיתן ערך יותר גדול ללקוח.

שלי תשובה

כן, אבל השוק מוגבל

אורן כהן

השוק הוא לא מוגבל. בעולם הביטוח, כן? אני לא מכיר את עולמות הבנקאות. מה שמעניין אותי- עולם הביטוח, אנחנו בישראל במאקרו, כן? הצריכה הביטוחית של אזרח פה היא בערך 50 אחוז מה-OECD אוקיי? כלומר, יש לנו בעולם הביטוח לאן לגדול ולעשות הרבה יותר. כמובן, כשנתסכל על סוף הדרך, בסוף המספרים מתכנסים. יחד עם זאת, מישהו צריך לנהל את העולם הזה. והעולם הזה צריך להיות מנוהל על ידי אנשים.

שלומי שוב

אני מסכים עם שלי, אני רק שואל אותך מזווית אחרת- התפקיד של סוכן הביטוח, הרי מן הסתם התהליך הזה שאתה מדבר עליו, חלק משמעותי מהערך יעבור ללקוח, יחסוך עלויות ללקוח, על חשבון מי זה יבוא? כלומר, מה יהיה פה במערכת היחסים שבין חברת הביטוח לסוכני הביטוח? הרי מנהלי חברות ביטוח אומרים שהסוכנים מרוויחים יותר מהם. איך אתה רואה את התפקיד בטווח הרחוק של סוכן הביטוח, ובשיווי משקל הזה בין סוכן הביטוח לבין חברת הביטוח, האם יהיה שינוי?

אורן כהן

אני לא מתעסק בצד של היצרנים, אני מתעסק בצד של ההפצה, ושם את הלקוח במרכז. כשאני שם את הלקוח במרכז, ואני מצליח ליצור איתו מח"מ אחר לגמרי ממה שהוא מכיר היום, משך חיים ממוצע, של הפעלה שלי מול הלקוח, אחד, המטרה שלי שהוא ישלם הרבה פחות, והמטרה השנייה שלי שהוא יישאר איתי לזמן הרבה יותר ארוך. היצרן מהצד שלו, יש לו את האירועים שלו, הוא צריך להרוויח ממקומות אחרים. זה שאני, באמצעות ה-AI, ככל הנראה אקטין גם את הפרמיה שהלקוח ישלם, זה כבר עניין שלי מול הלקוח, וזה הערך המוסף שאני נותן לו.

שלומי שוב

יכול להיות שהיצרן, אולי במידה מסוימת יכול גם לדלג עליך, בעקבות התהליכים האלה?

אורן כהן

לא, לא, ממש לא קשות. פשוט העובדות הן מאוד פשוטות. בעבר, כשנכנסו הביטוחים הישירים, הצעד הנוסף שלהם היה, ללכת לכיוון מערכות מידע, ובוא, וחמש חמש חמש, ותעשה ביקור בדיגיטל ונעשה לך הנחה... בסוף יהיה AI, אבל יהיה מישהו שילווה את ה-AI, כי בסופו של דבר, מה זה AI? המכונה, הדבר הזה?

לקחת את השכל של הבן אדם, ולהפוך אותו לשכל של מחשב. וזה בסופו של דבר, אם זה לא יהיה משולב ביחד, זה לא יקרה. אז ככה אנחנו רואים את הדברים, ואני מאוד מאמין ומקווה שהרגולטור, והוא קריטי פה, הוא קריטי, קריטי, ברגע שהוא ייתן לנו את הכלים ואת האפשרויות לפתח את התהליכים האלה ביתר שאת, בסוף המטרה שלו היא זהה לשלי-שהלקוח ישלם פחות, ויישאר איתנו הרבה יותר זמן.

יעל רגב

אני אשמח לשמוע ממך, מה חסר כרגע ברגולציה?

אורן כהן

אחד, חסר שיח. אז בואי נתחיל בזה. חסר שיח. כי אני לא ראיתי אתכם פונים לעולם ההפצה ואומרים לו בואו, בואו תסבירו לנו איך אנחנו באמצעות טכנולוגיה יכולים

לגרום ללקוח לשלם פחות או לחלופין לקבל שירות טוב יותר. זה, ואני מדבר, אני המפיץ, בוא נגיד אחד הכי גדולים שיש, אוקיי? והמשמעותי ביותר שיש. ובנוסף לעניין הזה, השילוב שלנו יכול לשרת את כולם.

גם את היצרנים. משום שבסוף היצרנים המסורתיים בישראל, ש80 אחוז מהפעילות בישראל היא על ידי הפצה של בני אדם, להלן סוכני ביטוח, אתם קוראים לו סוכן ביטוח אבל אנחנו מנהלי סיכונים. למעשה לוקחים את הלקוח, מנהלים עבורו את הסיכון, ובסופו של דבר הקשר הזה בין החברות המסורתיות, שמיוצגות כאן, יש פה משאר חברות הביטוח חוץ מהפניקס, אוקיי? הקשר הזה הוא הקשר כביכול, שכולם שמים עליו את העין, לאו דווקא עם הרגולטור. ההשפעה שלכם הרבה יותר גדולה, מההשפעה שאתם חושבים שיש לכם על התהליך.

יעל רגב

אבל איפה הרגולציה כרגע חוסמת אותך מלפתח את הכלים האלה? או, מה חסר לך, מה היית רוצה לראות מהרגולטור אחרת?

אורן כהן

אני לא אמרתי חוסמת. אני פעם ראשונה שאני, אני היחידי אגב, שיכול להרשות לעצמו להשקיע עשרות מיליונים, בתחום הזה כמפיץ. עוד לא קם סוכן ביטוח שאמר אני רוצה להיכנס לעולם הזה, כי אגב, זה הכל מתחיל ממה שאמרה הגברת מבנק לאומי, זה הכל ROI, זה הכל איך אני מחזיר את ההשקעה, ומה הדרך שלי להיקשר ללקוח בצורה יותר משמעותית וארוכת טווח, וזה ייתן לי, איזשהו אקסטרע על השוק. תמיד אנחנו הובלנו את השוק, אם זה על ידי מכוונות שעשו ציטוטים למחירים וכו' וכו', וגם בתחום הזה אנחנו נוביל, ואין לי ספק שברגע שזה יקרה על ידי הפניקס סוכנויות, או לא משנה, הזרועות שלה, גם חברות אחרות יכנסו לזה, וזה יועיל מאוד ללקוח.

שלומי שוב

אתה סוג של מכניס דיסטרפן לשוק.

אורן כהן

רגע, ודבר אחרון, כן. השיבושים האלה הם תמיד נכונים, כל התחום הזה של המחשב, הוא מקדים את העולם, את הקומודיטיס. הוא מקדים את עולם הבנקאות הפיננסית והביטוח. אבל דבר אחד אני כתבתי לעצמי פה.

שלומי שוב

אתה תפחית את העמלה על הקומודיטיס?

אורן כהן

יכול להיות כן, אבל אני אגדיל את הפעילות, ואני אמצא את הדרך לתגמל את עצמי בצורה כזאת שהרווח שלי לא יקטן. אבל אני רוצה לסגור את מה שאני אומר בדבר מאוד פשוט. בסוף ההפצה, עולם ההפצה, הוא המוצר העיקרי. אוקיי? בואו נשים דברים על השולחן, איך אמר הבחור, בסוף להביא את הלקוח, זה הסיפור האמיתי.

מאיר אהרוני

אני רציתי להגיד לך שהמח"מ יתקצר אבל נדבר על זה תכף..

אורן כהן

תתפלא, תנסה אותי, תראה איך אתה מאריך אותו. אבל בואו, הפצה ממה היא נובעת? היא מגיעה מאנרגיה. הפצה היא מגיעה מאנרגיה. ולדעתי דווקא עולם ההפצה הוא האחרון להיפגע בשונה ממה שאתם חושבים, אוקיי? כי יהיה מאוד מאוד קשה למחשב להכניס את האנרגיה של המפיץ כדי להביא את הלקוח אליו. זה הקלוז'ר לסיפור.

אמיר ברנע

מה דעתך על מסקנות הוועדה?

אורן כהן

אני האמת לא קראתי את כל המסקנות, אבל אני אתעמק בזה.

שלומי שוב

שלומית, עד כמה את רואה את הכניסה של הפינטקים לתוך הדיסטרפן כאן?

שלומית ווגמן

אז אני באמת חושבת שאני מדברת אולי כמה וכמה שנים קדימה ממרבית האנשים שדיברו כאן, מנקודת הבת של הפינטקים העולמיים, שבימים אלו מתחילים להשתלב גם בישראל. לפינטקים יש יתרון תחרותי עצום בהטמעת בינה מלאכותית מאשר גופים פיננסיים ותיקים: קודם כל אין להם את

ה LEGACY הכבד של מערכות ישנות שמושכות אותם אחורה כמו שתואר כאן, ומטבען הן הרבה יותר חדשניות, יש הרבה יותר טכנולוגיה, "TECH" ולא רק FIN, ולכן אנחנו רואים שם אימוץ הרבה יותר מהיר. אני עד לאחרונה עבדתי ב-RAPYD, היום אני עובדת בפינטקים ברחבי העולם, ואני גם מוציאה עכשיו ספר על Agentic Commerce, ומה שאנחנו עושים שם זה באופן מואץ, קליטה של AI כדי בעצם לפשט וליעל תהליכים. בעולמות הפיננסיים, מכיוון שהמודל העסקי ברובו מתבסס על גביית פרומיל אחוזים מהעסקאות, כל התייעלות תפעולית מובילה לרווח מהיר. זה שוק מאוד תחרותי, מי שעושה את זה הכי טוב מנצח בתחרות הזאת.

מנסיוני, ישנם מספר CASES USE מרכזיים של הטמעת בינה מלאכותית בגופים פיננסיים, והם: אוטומציה (שלא חייבת להתבצע באמצעות AI, אבל זה הדבר שהכי מאיץ את זה), פרסונליזציה של השירותים ללקוח, אופטימיזציה והמעבר הכללי למסחר אוטונומי שמבוצע באמצעות סוכני בינה מלאכותית – AGENTIC COMMERCE.

כדוגמה, במקרה של ה-AI שהטמענו ב-RAPYD, החלפנו את תהליך ה-ONBOARDING של הלקוחות במכונות ואייג'נטים. המהלך הוביל לצמצום של 90 אחוז מכוח האדם, ולביצוע תהליך שלקח שעות ולפעמים שבועות – לשניות בודדות. לקחנו תהליך שיש לו POLICY רגולטורי מאוד ברור של מה צריך לעשות. חשוב להדגיש שזה תהליך עם המון נהלים ופרטים, כשבפועל העובדים הזוטרים שמבצעים אותו, ה-low cost employees, מבצעים טעויות כל הזמן, ב-X אחוז מהמקרים, וקשה לנטר ולתקן את זה. ומה שקיבלתי כאן מבחינת יכולת ניהול תהליך עי AI זה מצב שבו גם אם הבינה המלאכותית עושה טעויות, עדיין זה בד"כ הרבה פחות טעויות מאשר אנשים שהיו עושים את העבודה, ויש לי יכולת בקרה הרבה יותר גבוהה, וניתן לתקן את התהליך בשינוי אלגוריתמי רוחבי בפעם אחת. אז כן זה תהליך מאוד מקצועי ומורכב ללמד את המכונה לעשות את זה, וצריך מישהו עם כתפיים רחבות שיוודע איך להתאים את זה לדרישות הרגולציה וליצור את התהליך באופן נכון. כי כפי שתואר, זה מאוד לא פשוט לכתוב prompts לתהליך ויש הרבה בעיות של הטיות ומניפולציות, ושלל נושאים של AI.

אבל ברגע שיוודעים להכיל אותם וכן מנהלים תהליכים באופן נכון, אפשר להגיע להישגים יוצאים מן הכלל. יכולת קליטת הלקוחות שלי עלתה עשרות מונים והתכווצה לשניות במקום תהליך של שעות וימים. אגב, פה ברייכמן, אנחנו מלמדים עכשיו קורס- Chief AI Officers. להכשיר את אנשי המקצוע בארגונים לפתח ולאמץ AI באופן אחראי שעולה בקנה אחד עם הרגולציה.

עוד דילמה שאני רואה בכל הגופים הפיננסיים ברחבי העולם, גם בנקים וגם פינטקים, זה שאלת ה-Build Vs. Buy האם אני אעשה את זה בעצמי בתוך הארגון או שאני אקח מומחה?

החדשות הטובות הן שבחלק גדול מהמוסדות והגופים אפשר לעשות Build, לעשות פנימית בתוך הארגון כברת דרך ארוכה ומועילה. אבל החל משלב מסוים, צריך הרבה פעמים כבר את ה-EXPERT שיתמכו בכך.

בנוסף, מהניסיון שלי אני רואה שאם זה לא יוזמות שמובלות ומתועדפות על ידי המנכ"ל עם פורום ייעודי ל AI, זה לא עובד טוב. יש המון יוזמות AI בתוך הארגון, הבעיה שהיוזמות האלו מבזבזות לכם דאטה, משאבים, אנרגיה ולא משפיעות מספיק על שורת הרווח. נכון, צריך שהארגון ילמד איך לאמץ AI וישתמש בתחום, אבל כל היוזמות המקומיות grassroots, אם הן לא בעלות יכולת להזיז את המחט העסקי, הן לא נותנות את הערך המתאים, את ה ROI שצריך. והארגון מפספס. חשוב גם לדעת איך לבחור מלכתחילה פרויקטים שגם יתאימו למה שהגולטור רוצה, וגם יתאימו לעסק, זה קריטי.

זה טוב שעובדים יתנסו, אבל לא להתבכש עם זה יותר מדי, כי פשוט תראו שאתם נופלים בסטטיסטיקה שאנחנו רואים ב-MIT שכ- 85-95% מהפרויקטים נכשלים, כי לא בוחרים את הפרויקט הנכון או שלא מוביל את זה האדם הנכון. זה לא יכול להיות יוזמה של הפרודקט, של ה-CTO של ה-Compliance, צריך גורם אסטרטגי עסקי שבסוף יהיה מגובה על ידי המנכ"ל ולשנות את התרבות.

לגבי הרגולציה: מה שאנחנו רואים שמתהווה בעולם, והרגולטורים פה הישראלים שנמצאים בוועדה, הולכים בדיוק בכיוון הנכון. הדגש צריך להיות שלא בודקים את ה-AI אתה עושה, באיזה טכנולוגיה, אלא את ה"מה" – מה התוצאות מהתהליך. בהנחה שארגון מבצע קודם כל את ה-governance הפנימי, שרשרת התכנון, המדיניות והאחריות, ויש לך human in the loop, ואתה מתעד את התהליכים ומאפשר Explainability. הבחינה היא מהותית לגבי התהליך והבקורות בו, מה אתה עושה ולא איך. כמו שאני משתמשת בתוכנה כזו או אחרת, אם יש לי אנשי מקצוע שיודעים לעשות oversight ובקרה ופיקוח ואני מטמיעה ומפתחת נכון, כפי שלמדנו קודם את המודלים, עם האקספלנביליות וכל המונחים, אפשר ליצור מודל שעומד בצורה טובה ברגולציה, או לכל הפחות כרגולטורית לשעבר, סליחה לא הצגתי את עצמי, אולי ראש הרשות לאיסור הלבנת הון לשעבר..

אמיר ברנע

בלי להסביר איך אתה עושה?

שלומית ווגמן

לא, לא, אתה מטמיע פנימה הרבה מאוד מנגנונים, שהופכים את הזירה הזו לניתנת להסברה, פיקוח ובקרה, ויש לך את האחריות הזו להראות ולעשות איך עוברים את זה. על זה העבודה שלי בהרווארד בפינטקים. אפשר להגיע, יש כברת דרך לא קצרה, אבל עם אנשי תוכן ומומחים אפשר להגיע לשם, ובעולם הפיננסי, ואולי עם זה אני אסיים, זה מקום מעולה להשתמש ב-AI, מכיוון שזו זירה שנמצאת תחת רגולציה כבדה, כל דבר נבדק בקפידה, ויש רגולטור שידפוק לך בשלב מסוים בדלת, ויבקש לראות איך ומה עשית. מי שמכין את זה צריך לעשות את זה באחריות. ויודע להחיל את ה-policies באופן מדויק ותקין. פעמים רבות בדקתי את זה מול אנשים, בפחות טעויות מאנשים, או טעויות מסוג אחר, שיש לי יכולת לתקן. אבל זה אולי כבר לדיון הבא.

רן שחם

אני אשאל שאלה את הרגולטורים, רן שחם, אלשטולר שחם. גם אנחנו משתעשעים עם AI כמו כולם, מי שפחות משתעשע עם AI, זה הפרודיסטים. הם פחות כפופים לרגולציה, ואנחנו חווים בחודשים האחרונים..

הפרודיסטים, כל הרמאים למיניהם. רמאים שמשתמשים בAI ברמאויות כאלה ואחרות. לא יודע למי יצא לראות את גילעד אלטשולר, שותף שלי ממליץ על מניות סנט בכל מיני פורומים, אנחנו חווים גל מאוד גדול של ניסיונות FRAUD כלפי הלקוחות שלנו, אני חושב שזה לא רק אנחנו אני חושב שזו כל המערכת הפיננסית חווה את אותו דבר.

הייתי רוצה לשמוע, אנחנו פונים לרגולטורים, למשטרה, הגשנו כבר כמה תלונות במשטרה, פייסבוק, טיקטוק, וואטאבר, אנחנו לא מרגישים שאנחנו מקבלים מענה כמו שצריך, וזו סכנה אמיתית.

אמיר וסרמן

גם אותנו סיכון הדיסאינפורמציה מאוד העסיק, ובשנה החולפת אפשר היה לראות איך הוא מתחיל להתממש ביתר שאת בעולם ובישראל. אפשר לראות למשל כיצד בונים באמצעות כלי AI זיופים של דמויות ידועות שממליצות לכאורה על השקעות. אלה הרבה פעמים הונאות בסיסיות מסוג pump and dump הנעשות כעת בעזרת בינה מלאכותית. סיכון נוסף שאנו מודאגים מהתממשותו הוא הונאות של גניבת זהות, שאף הן מתגברות בעולם. הבינה המלאכותית מגבירה סיכונים שהיו קיימים גם בעבר כיוון שהיא מאפשרת בקלות, ובעלות נמוכה, לזייף זהויות. היא גם מאפשרת להפיץ מהר את המידע השקרי, וזה באמת אחד הסיכונים שהבלטנו בדוח.

ההתמודדות עם סיכון הדיסאינפורמציה אינה קלה ותצטרך להיעשות בכמה מישורים. מישור אחד הוא חינוך פיננסי לציבור על מנת להגביר ערנות ודרכי התנהלות נכונות. מישור שני הוא היערכות על-ידי הגופים הפיננסיים עצמם, בדומה להיערכות שלהם בנושאי סייבר, כך שיוכלו להתמודד בין היתר עם הונאות מסוגים חדשים. מישור שלישי הוא טיפול בדרכי ההפצה של המידע, וכפי שהתפרסם הרשות מנסה להביא לכך שמידע שקרי יטופל כבר ברמת הרשתות ופרסומים מזויפים יצומצמו ויוסרו באופן מהיר.

מאיר אהרונ

בסוף יש פה אוזלת יד מדינתית.. בוא נלך הכי פשוט, פשינג. פעם אם הייתה מודעת פשינג מזויפת, בטח אם היא מחו"ל הייתה רואה את הכתב העילג, הייתה רואה משהו לא טוב הועתק ב look and feel של אתר הארגון שמעתיקים אותו, אבל הבינה המלאכותית השתפרה, וזה נהיה הרבה יותר קל. זאת אומרת, כל שני וחמישי, אני מניח, לכל בית השקעות או מוסד פיננסי אחר, עולה איזה אתר מתחזה, שולח הודעה ללקוח, עם הודעה תמימה בסגנון אנא אמת פרטיך. ותמיד יהיו את הלקוחות שיפלו. אפשר לתת מענה בצעדים הכי פשוטים. קשר בין הרגולטור שלנו לבין משרד התקשורת. בא fraudster,

בסדר? מישו שעושה עבירה, עושה sim swap. גונב לך את הסיים (שזה אירוע הרבה פחות מחמיר מול חברת תקשורת לעומת מול הבנק). המינימום, תנו אינדיקציה, כדי שבפעם הבאה שהוא יזדהה מולנו, שנדע שזה SIM חדש.

דבר נוסף, אני אענה על הנושא הקודם. דיברת על מח"ם (אורך כהן) ואמרנו את זה רגע בצחוק. תראו, בסוף כשאנחנו התחלנו להטמיע דיגיטל בצורה אינטנסיבית, חשבנו שהדבר הזה יאפשר לנו להתייעל במוקדים ולהיות הרבה יותר קטנים, אבל מה שזה עשה זה פשוט העלה ביקושים. כי זה עשה אותנו הרבה יותר נגישים, ולקוחות התחילו לפנות בהרבה מאוד נושאים כשלפני זה היה החסם של שיחה ושל המתנה. כפועל יוצא אנחנו בעלויות לא יורדים. עכשיו, דיברנו על מפיץ לעומת יצרן. בסוף, עבור יצרנים, להעמיד שירות או מוצר, עולה כסף.

אם כל האירוע הזה בסוף יסתכם בזה שלקוחות כל היום יושבים, בוא נלך עוד פעם על המח"ם. לקוחות לא כל כך מבינים בהשקעות, אז יהיה מישו, ידען חכם וחרוץ שכן מבין, יודע להשוות בין ביצועים של בתי השקעות, מבין מה זה קרן כספית, יודע כל מיני דברים שהציבור לא יודע ולרוב גם משתעמם מהם, ופתאום כל שני וחמישי ישווה ויבקש להחליף. מה זה יעשה לעלויות? יוריד אותן או יעלה אותן? בסדר? ואז יכול להיות שתראה תופעה שהיא לא בהכרח כזו שאתה חושב.

שלומי שוב

ואז מה זה יעשה למפת התחרות?

מאיר אהרוני

בסוף היצרנים וגם מי שנמצא בתווך, אולי יוכלו להתייעל במענה לשירותים בנאליים, שם ירד הצורך בגלל האוטומציה, אבל פתאום הלקוחות פונים הרבה יותר, עוברים כל שני וחמישי, משווים כל שני וחמישי ואז העלויות דווקא עולות והמח"ם אגב יורד.

אני לא יודע מה זה יעשה, אבל אני יכול לדבר על הציפיה להתייעלות במהפכות קודמות. במהפכת הדיגיטל, היינו בטוחים - אני מסרתי הבטחות כאלה ליותר ממנכ"ל אחד, שאם נעשה א' ב' ג' ונשקיע בערוצים דיגיטליים, נשקיע בקונטקט סנטר חכם, אנחנו נוריד ביקושים ולכן נוכל להתייעל. בפועל זה לא מה שקרה.

שלי תשובה

בעשור ירדו 12 אלף איש מהמערכת הבנקאית, הוצאות השכר של הבנקים היו FLAT נומינלי בעשר שנים, איך אפשר להגיד שהבנקאות לא התייעלה בגלל הדיגיטל?

מאיר אהרוני

בסדר, יש לי מה להגיד על זה.

שלומי שוב

הוא רוצה להגיד שהיא התייעלה, הייתה צריכה להתייעל בלי קשר. אני אפרש אותך.

שלומית ווגמן

אבל יש שינוי בכישורים של העובדים. צריך מעט מאוד עובדים זוטרים, הם מוחלפים במכונות, ויש צורך בהרבה יותר אנשי מקצוע מומחים שידעו לבדוק באופן ביקורתי ומקצועי את התוצרים של המכונה. מחד צריך להעסיק פחות עובדים, אבל מאידך השכר הממוצע של העובדים שנדרשים - גבוה יותר.

משה ברקת

אני חושב שדני ימין פתח בזה שבעצם השאלות המרכזיות הן לא סוגיות יישומיות או מקצועיות בהכרח. השאלות המרכזיות הן מוסריות, פילוסופיות ואסטרטגיות. יש פה שאלות מאוד גדולות.

מבחינת לקוחות, האם אנחנו נוכל לשרת את כל הלקוחות? ה- Law of Large Numbers למשל, האם הוא ימשיך להתקיים? האם חלק מהאנשים יסבסדו אנשים אחרים, או שזה כבר לא יהיה אפשרי? האם חברות גדולות עם יותר מידע, יאכלו חברות קטנות? בסופו של דבר, בנק לאומי מול בנק ירושלים, זה לא כוחות מבחינת יכולת השימוש ב- AI וגודל ה- database שעליו אפשר לאמן את המודל. עם היתרון הזה בנק לאומי לוקחים עכשיו את כל הלקוחות. אז איך הדברים האלה ישתנו מבחינת תפיסת התחרות, שימוש במידע, ועוד?

אני חושב שאלה שאלות גדולות שיעסיקו אותנו יותר מאשר השאלות היישומיות, כי השאלות היישומיות מבחינה טכנולוגית מוכרות לכל מי שהתעסק ב-SOFTWARE. אז היום הסקטור הוא קצת יותר מתוחכם, אבל הדבר שהוא שונה הוא בעיקר האלמנט הזה של הג'ינרטיב. עכשיו, בהקשר הזה, הסברתיות זה דבר יפה מאוד, אבל אתה לא יודע להסביר מה שאתה לא יודע להסביר. כי הג'ינרטיב יכול לקחת אותך למקומות שלך תסביר לאן הוא לקח, כשאתה לא יודע לאן הוא לוקח ולאן הוא ייקח.

טל הפלינג

בוקר טוב לכולם, טל הפלינג, שותפה במשרד עו"ד שיבולת וגם עמיתת מחקר במכון לפנסיה וכלכלה באוניברסיטת בן גוריון וגם לשעבר חברת הנהלת רשות שוק ההון.

רציתי להוסיף שתי מילים על מה שנעדר לי מהדו"ח - אני הסתכלתי עליו גם בעיניים של משפטנית, ברמת האותיות קטנות ורמזים לאסדרת המשך, וגם בעיניים של חוקרת אקדמית, ובשתי הפריזמות - חסרה לי התיחסות להשפעות AI על הפנסיה. מוצר הפנסיה בישראל הוא הר הכסף הכי צומח, שגובהו כשלושה טריליון שקל בערך נכון להיום, בקצב צמיחה של כ-10 אחוז בשנה, וזה מוצר שמשפיע על החיים של כולנו.

כדי לנסות להבין עמדה משתמעת לגבי הפנסיה, הסתכלתי על פרקי הביטוח ועל הבנקאות שבדוח. אלה יוזקייסים שיש להם דמיון אחד לשני- כי כש the house wins the costumer lose – גם בחיתום אשראי וגם בחיתום ביטוח, אבל זה לא דומה לפנסיה, כי בפנסיה מדובר בדמי ניהול מסך הנכסים המנוהלים והמוצר מתנהל אחרת מבחינת הטעות שהמערכת עושה- אם יש בעיית ניהול השקעות מבוססות AI, התשובה של הלקוחות תפגע, ותהיה השפעה עקיפה על ההכנסות של החברה שמנהלת את הפנסיה. אז אני לא יכולה לקחת את הניתוח לגבי בנקאות וביטוח וליישם על הפנסיה.

הסתכלתי על הפרק של ניהול תיקים, אבל זה גם לא ממש אותו דבר. קודם כל, ניהול תיקים היא תעשייה קטנה יותר, ומוצר ניהול תיקים תחום יותר- מבחינת אורך חיי תיק ההשקעות וסוגי הנכסים המנוהלים. לא פחות חשוב- בניהול תיקים ללקוח יש הרבה יותר מידע על הרכב התיק והתשובות לעומת החיסכון הפנסיוני שלו, ואז יש יותר נתיבי בקרה לאתר טעויות שנובעות, למשל, מניהול השקעות מבוסס AI לא מבוקר דיו. אז מה קורה עם הטמעת AI בתהליכי ניהול ההשקעות בפנסיה, שהם פחות ברי עקיבה על ידי הלקוח?

חלקם יהיו פשוטים יחסית- יש רגולציה שכוללת כללים טכניים מגבלות כמותיות על התיק- ואז קל לאתר כשלים. אבל רוב התיק הוא תיק חופשי מבחינת סוגי נכסים וההחלטות על ההשקעה בהם, כשעל כל התהליך חלה חובת נאמנות כללית. מאיר לוי דובר קודם על חובות נאמנות כלליות כאלה, שאתה לא בדיוק יודע איך תהיה אכיפה שלהן, ולכן לא ברור מה הציפיה לגבי הטמעת AI בתהליכי ניהול השקעות בפנסיה.

במצב בו אין איתות רגולטורי בכלל, ואני עכשיו צריכה לשבת פה כיועמ"ש של גוף מוסדי, ושואלים אותי אם מטמיעים או לא מטמיעים AI בתוך תהליך ניהול השקעות כזה או אחר, התשובה מורכבת. אני עלולה לעשות אחת משתי הטעויות. טעות מסוג אלפא זה להגיד- אני שמרנית. לא כתבו מילה על הנושא בדוח, סימן שאסור. כי אני מחפשת את הרמזים את האותיות הקטנות. או שהייתי עושה טעות מסוג בטא- לא אמרו לי שאסור, סימן שמותר, וננסה לבדוק מה באמת קרה, כשמישהו ישים לב שמישהו השתבש.

שתי הטעויות לא טובות לגופים המוסדיים וללקוחות שלהם. אני חושבת שראוי, ולו במילה, להגיד מה הציפיה לגבי הטמעת AI בניהול השקעות בפנסיה. זה לא חייב להיות קונקרטי, זה יכול להשתמש באותם כלים שהדוח עובד איתם לאורך כל הדרך- עקרונות של הסברתיות, תמיכה בהחלטה לעומת ביצוע מלא וכו', כי הנזק בשני סוגי הטעויות עלול להיות גבוה.

יעל רגב

בשתי מילים תן לי רק להגיד, שזה צוות שעבד באופן רחבי, והסתכלנו על סוגיות שהן רחביות בכל הסקטור הפיננסי, ונכונות לכל הרגולטורים, וככה גם בכתב המינוי, סוגיות כמו אסדרת שינויים של מערכת הפנסיה שהיא נטו ברשות שוק ההון, כמובן לא נכנסות לצוות בין משרדי, אבל דיברנו פה על

עקרונות שנכונים לכולנו, ואחרי זה כל אחד יישם אותו ברגולציה שלו, בהתאם למה שנכון לאותו תחום.

אבי בן נון¹

לראיית AI היא מהפכה של ממש אבל בניגוד להייפ שהיה בתחילה צריך להבין שישנה אבולוציה, אין פה קסם נדרש זמן אימון ניכר, יש לא מעט יוזמות שנכשלות, ללא ROI מספק כאשר איכות המידע ועומק המידע הינם משמעותיים להצלחת הפרויקטים. לטעמי יש לחלק את ההסתכלות האסטרטגית לשלושה קבועי זמן שלכל אחד ערכים מוספים שונים.

בטווח הקצר עיקר האירוע יסכם להערכתי ביעילות תפעולית בעולמות שירות הלקוחות, תפעול, פיתוח ובדיקת קוד בעיקר.

בטווח זמן הבינוני נהנה משיפור בפריז, ניהול סיכונים, תמחור, ניהול תביעות ותמיכה אקטוארית טובים ומוותאמים יותר ובזמני תגובה מהירים יותר כאשר תמהיל כח העבודה יהפוך היברידי תוך שילוב סוכני AI עם כח העבודה האנושי ושינוי מסוים בתמהיל המשרות.

בטווח זמן הארוך ובשילוב גם עם מחשוב קוונטי נראה שינויים במודלים עסקיים, במודלי הפצה ובטעמי הציבור.

לאור זאת ומכיוון שהתועלות העתידיות עשויות להיות משמעותיות נדרשת רגולציה מאפשרת כמובן תוך התייחסות לסיכונים הייחודיים שמודלי AI מביאים עמם אבל יש לאפשר לגופים ניהול סיכונים פנימי בהתאם לתאבון הסיכון של כל גוף, תוך קביעת ובניית סט בקורות בהתאם לסיווג הסיכון של יוזמות AI מתוך הבנה שכל אחד ממקטעי הפעילות שרלוונטי בטווח הנראה לעין למודלי AI נתון מלכתחילה לרגולציה קיימת ללא תלות בשיטת ההפעלה.

אמיר ברנע

מילה אחת. למה קוראים לזה בינה? היכן כאן הבינה? האם היא תחליף טלפון לבנק, דבר עם היועצת שלך, האם מגע אישי כבר לא יתקיים יותר, טלפון לחברת הביטוח, להתייעץ עם הסוכן, הסוכן לחברת הביטוח, כל הדברים האלה שהם לב הפעילות, הממשק בין לקוח לגוף פיננסי. האם כל אלה ישתנו? האם השאיפה לתגובה אנושית בהחלטות פיננסיות היא אנכרוניזם?

שלומי שוב

דברי הסיכום שלי בדיוק מתחברים לדברים של אמיר, כי כאן לדעתי נעוצה הבעיה במחשבות ובדיונים של הדירקטוריונים היום. קשה לנו להבין את זה אבל הצעירים היום, שהם הילדים או במקרים רבים תלוי בגיל הנכדים שלנו, מעדיפים בכלל לא לדבר טלפונית עם נציגי שירות, שלא לדבר על לא להגיע פיזית לסניף.. מעדיפים שהכל ייסגר דרך האפליקציה - מה שנקרא "ללא מגע יד אדם"...

¹ אבי בן נון, משנה למנכ"ל ומנהל סיכונים ראשי כלל ביטוח, התווסף לאחר הדיון כי היה צריך לצאת

ודירקטוריונים היום חייבים להיכנס לראש הזה כי אחרת יפספסו את הרכבת... כלומר את החשיבה האסטרטגית על העתיד.

נקודה קריטית שעלתה כאן לדעתי היא הנושא הרגולטורי של היכולת של הגופים הפיננסיים להזרים לרובוטי ה-AI מידע רגיש שנמצא ברשותם ויכול להוביל למשל לתמחור וחיתום הרבה יותר נכון ומדויק בחברות ביטוח אבל גם יש לו פוטנציאל להוביל לאפליה. כלומר, השאלה היא עד כמה מונעים מראש שימוש במידע ולא נותנים שיקול דעת/אחריות לגופים עצמם שלא לבצע אפליה? אני מבין שבעקרון כבר כיום קיימות מגבלות מסוימות על אפליה ועל נתונים בהם מותר ואסור להשתמש ואני מבין את מה שהרגולטורים אומרים שהגופים אחראים ליישם את הכללים האלה גם כשמכניסים AI לתהליכים – אבל יש הבדל גדול מהניסיון שלי עם כלי ה-AI באיזה שלב מונעים את השימוש בהם... ככל שמונעים אותם כבר בשלב של נתוני קלט אז די מעקרים את הכוח של ה-AI... זה הבדל שם שמיים וארץ. מניח שאלו דברים שיתחדדו בקרוב במסגרת המרוץ של הגופים הפיננסיים לעבר ה-AI בשורה התחתונה, אני חושב שבאמת היה לנו דיון משמעותי וחשוב על מהפכה בהתהוות – ממש בתחילתה. אני חושב ששמענו לא מעט מהגופים ואני לא מופתע שכל גוף לוקח את זה לכיוון אחר, ברמת דאגה ודחיפות אחרות, כאשר במקביל צריך לשמור על "עיניים על הכדור" כדי לא לפספס את המצב הנוכחי וזה לא פשוט. זה מובן, כי יש פה הרבה אי וודאות ודני ימין השווה את זה להשלכות על האנושות של גילוי אמריקה או על גילוי חיים במאדים... וזה ברור לי שהגופים גם לא יכולים להיות חשופים לגמרי, כל אחד שומר על האינטימיות שלו, על היתרונות התחרותיים שלו והכל בסדר. דבר אחד בטוח – כל דירקטוריון והנהלה חייבים בימים אלה להיות בפקוס על הנושא האסטרטגי הזה ולהכריח את עצמם לחשוב אחרת בראיה של הדור הצעיר. תודה רבה.

על הפורום:

פורום שווי הוגן (Fair Value Forum - FVF) הוקם בכדי לתרום לאיכות המידע בשוק ההון ולייצר שיח רב דיסציפלינארי פורה בסוגיות מקצועיות הקשורות לשוק ההון ולעולם העסקי שנמצאות על סדר היום הציבורי.

הפורום, בייסודם של פרופ' אמיר ברנע, הדיקן המייסד של בית ספר אריסון למנהל עסקים ושלומי שוב, ראש תכנית חשבונאות, סגן דיקן בית ספר אריסון למנהל עסקים וראש מכון הפניקס לחקר שוק ההון, פועל במסגרת מכון הפניקס לחקר שוק ההון באוניברסיטת רייכמן.

הפורום כולל מומחים מובילים מהאקדמיה ומהפרקטיקה בתחומי החשבונאות הפיננסית והמימון וכן משתתפים בו נציגים של מדווחים, אנליסטים וגופי הרגולציה הפיננסית בישראל.

סיכומי הדיונים שמתפרסמים לציבור הרחב מבוצעים על ידי הצוות המקצועי של הפורום המורכב מבוגרים מצטיינים של התכנית בחשבונאות.

לאתר הפורום: [לחץ כאן](#)