



הפינה המקצועית של ירדן פלד

מומחה לטרנספורמציה דיגיטלית מבוססת
בינה מלאכותית, מנהל עדיף Tech

עדיף
TECH

האם שימוש בבינה מלאכותית עלול לפגוע בתדמית המקצועית של עובד

בין חדשנות לפרדוקס חברתי: תובנות ממחקר רחב היקף
על סטיגמה בשימוש ב-AI והשפעתו על עולם הביטוח

בעולם שבו הבינה המלאכותית הולכת ותופסת מקום מרכזי בתהליכי עבודה, קבלת החלטות, שירות לקוחות ושיווק – מתברר כי לצד ההזדמנויות העצומות, קיים מכשול לא צפוי: התפיסה החברתית כלפי המשתמשים בטכנולוגיה הזו. תפיסת העובד את עצמו ואת ההשפעה של שימוש בכלי בינה מלאכותית ותפיסת העובד על המנהל שלו וסביבת העבודה הקרובה.

• תוצאה זו חזרה בכל קבוצות הגיל, המגדר והמקצוע, דבר המראה על סטיגמה החוצה מגזרים ומקצועות.

3. ניסוי שלישי – השפעה על גיוס עובדים

• מנהלים שהיו צריכים לבחור בין מועמדים דיווחו על נטייה לגייס את מי שאינו משתמש ב-AI במיוחד אם גם הם עצמם לא משתמשים בו לעומת מנהלים שכן משתמשים בכלי AI, דבר המראה על הטיה בגיוס המושפעת מאופי המנהלים והשימוש שלהם בכלי AI.

4. ניסוי רביעי – הקשר למשימה ולמגייס

• כאשר השימוש ב-AI היה מוצדק וקשור למשימה דיגיטלית – הסטיגמה נעלמה.

• כאשר השימוש היה למשימה ידנית – המשתמש נתפס כלא מתאים, בין היתר כי נתפס כעצלן. כלומר, כשיש הצדקה למשימה הסטיגמה יורדת וכשאין הצדקה היא עולה.

המסקנה מהמחקר ברורה: השימוש ב-AI טומן בחובו תג מחיר חברתי במקום העבודה ולכן העובדים נמנעים מלהשתמש או לדווח על כך, אלא אם כן הסביבה והמנהלים מכירים בערך ובנחיצות הכלי.

הרלוונטיות של המחקר לענף הביטוח

בעולם הביטוח – ובמיוחד בחברות ביטוח ובבתי סוכן גדולים – החלה בשנים האחרונות הטמעה של כלים מבוססי בינה מלאכותית, אך דווקא בחברות שמעסיקות עובדים שכירים, כגון מוקדנים, אנשי שירות, מנהלי שיווק ומיישבים – עולות דילמות זהות לאלו שתוארו במחקר:

• האם מותר לעובד להשתמש בצ'אטבוט פנימי או ב-GPT כדי לנסח תגובה ללקוח?

מאמר זה מבוסס על מחקר רחב היקף שפורסם במאי 2025 בכתב העת היוקרתי PNAS, תחת הכותרת "Evidence of a Social Evaluation Penalty for Using AI" מאת Jessica A. Reif ואחרים.

המחקר בדק דרך סדרה של ארבעה ניסויים בהשתתפות למעלה מ-4,400 משתתפים, כיצד נתפסים עובדים שמשתמשים בכלים מבוססי בינה מלאכותית בארגונים – והאם השימוש בטכנולוגיה ובינה מלאכותית פוגע בתדמיתם המקצועית בעיני אחרים.

ממצאים מרכזיים של המחקר

החוקרים ביקשו לבחון האם קיימת "ענישה חברתית" כלפי עובדים שעושים שימוש בבינה מלאכותית. ארבעת הניסויים שבוצעו מגלים תמונה מדאיגה:

1. ניסוי ראשון – תפיסה עצמית של עובדים

• משתתפים שהתבקשו לדמיון שימוש בכלי AI בעבודתם ציפו שיראו בהם פחות חרוצים, פחות מקצועיים ויותר ברי החלפה ולכן גם הרגישו בנחיתות להשתמש בכלים אלו.

• הם גם דיווחו כי לא ימהרו לשתף את השימוש בכלי בינה מלאכותית מול מנהלים או עמיתים במקום העבודה.

2. ניסוי שני – תפיסות של אחרים

• משתתפים שקראו תיאורים של עובדים ש"נעזרו ב-AI" לעומת עוזר

המשך בעמוד הבא <<<

למידה ושיפור, ולא כ"ידיד אישי", ואף תגמול על ביצועים מיוחדים.

5. מדידה לפי תוצאה, לא רק לפי דרך

• מעבר מחשיבה של "כמה זמן עבדת על זה" ל"מהי איכות התוצאה" - ללא קשר לכלי שנעשה בו שימוש.

6. התאמה למשימות ספציפיות

• יש להבחין בין משימות שדורשות שיקול דעת אנושי לבין משימות שחוזרות על עצמן וניתנות לאוטומציה וייעול ושיפור באמצעות כלי AI.

סיכום ודעה אישית: לנקות את המראה, לא את הבינה

כמי שמלווה תהליכי שינוי והטמעת AI בענף הביטוח, אני פוגש מנהלים שרוצים להתייעל, אך לא תמיד מבינים מדוע העובדים "לא זורמים" עם הכלים החדשים.

המחקר שהוצג כאן נותן תשובה: העובדים פוחדים להתפס כלא מקצועיים ומפחדים שיוחלפו בבינה מלאכותית. הם פועלים לעיתים מתוך דאגה לתדמיתם ולא בהכרח מתוך התנגדות לטכנולוגיה.

אם נרצה להוביל טרנספורמציה אמיתית, לא מספיק לספק כלים, צריך גם לאפשר מרחב חברתי וארגוני שמקבל את השימוש בהם ושיתוף מידע שגם יעודד את השימוש, גם יוריד לחץ וחרדה וגם יסייע לקבלת תוצאות טובות יותר - ומקרים רבים יותר בשימוש של כלי AI בקרב העובדים.

השימוש ב-AI הוא לא סימן לחולשה אלא להסתגלות ולגמישות ועובדים שמפעילים כלים חכמים. הם לא מבקשים לקצר דרך, אלא לייעל תהליך או לקבל רעיונות וביצועים טובים יותר.

בתחילת דרכי עם הבינה המלאכותית גם אני התביישתי בכך שסיכמתי והתחלתי לכתוב מאמרים וניתוחים באמצעות בינה מלאכותית. היום אני גאה בכך ואף מודה שאין מאמר או תכנית עסקית או פוסט שלא התחיל בשיח עם כלי מבוסס AI.

זה הזמן להנהלה לקחת אחריות: להוריד את הסטיגמה, להעלות את השקיפות - ולהפוך את השימוש בבינה המלאכותית לחלק טבעי, מובן מאליו, ומקובל בעבודת היום-יום, ולקצור את הפירות של הכלים החדשים הזמינים לעובדים.

המשך מהעמוד הקודם <<<

• האם שימוש במערכת ניתוח שיחות חכמה נתפש כהעצמה מקצועית או כחוסר שליטה?

• האם מיישב תביעות שמבצע ולידציה רגולטורית אוטומטית נתפס כמקצועי או כ"עוקף תהליך"?

בפועל, עובדים רבים מדווחים - בשיחות עומק ובהדרכות - כי הם חוששים להודות שהם משתמשים ב-AI. הם מפחדים להיתפס כעצלנים, כלא בקיאים בתוכן, או כמי שאפשר להחליפם בבוט.

הפער בין התוצאה (שמושפעת לחיוב מהשימוש) לבין התפיסה (שמושפעת לרעה מהשימוש) הוא בדיוק אותו פרדוקס שמתאר המחקר, והפחד מכך שהבינה המלאכותית תחליף בקרוב חלק מהעובדים משפיע על הביצועים ועל הדיווח בשימוש בכלים אלו - ובמקום להעצים, הם לפעמים מקטינים את העובד.

המשמעות הניהולית: איך מנהלים צריכים להגיב?

המסקנות מהמחקר צריכות להדליק נורה אדומה דווקא בקרב מנהלים - ולא בהכרח רק אצל המשתמשים עצמם אשר מאמצים כלי AI בהיקפים הולכים וגדלים.

הנה המלצות פרקטיות להנהלה בחברות ביטוח וסוכנויות גדולות:

1. לקבוע מדיניות ברורה ומוצהרת לשימוש ב-AI

• לדוגמה: הגדרת תהליכים ונהלים שבהם מותר ואף מומלץ להשתמש בבינה מלאכותית בעבודה, כגון כתיבת טיוטות, מענה ללקוחות, הפקת דוחות או ולידציה רגולטורית.

2. שילוב AI בהנחיות עבודה רשמיות

• למשל, בהנחיית חידושי ביטוח רכב - לציין כי ניתן ורצוי להשתמש בצ'אטבוט, תוך הגדרה של השלבים שבהם נדרש פיקוח אנושי.

3. הכשרות ייעודיות למנהלים ולעובדים

• הכשרה לא רק על השימוש בכלים, אלא גם על האופן שבו נכון לנהל דיון פתוח סביבם, ולמנוע תפיסות שגויות כולל ישיבות שבועיות והצגת מקרי בוחן.

4. שקיפות כתרבות ארגונית

• לעודד עובדים לשתף כאשר השתמשו בכלים חכמים כחלק מתהליך

