

20

המבקר האחרון? בינה מלאכותית ועתיד מקצוע הביקורת

פבל קריינין

מר פבל קריינין הוא מנכ"ל A-Impact, חברה המתמחה בתיווך בין מומחיות אנשית לטכנולוגיית בינה מלאכותית, ומלווה ארגונים במגזר הציבורי והעסקי בתהליכי הטמעת AI.

המאמר בקצרה

- 57% משעות העבודה בכלכלה האמריקאית ניתנות כבר היום לאוטומציה - 44% על ידי סוכנים דיגיטליים ו-13% על ידי רובוטים פיזיים. יכולות סוכני הבינה המלאכותית מוכפלות בכל כמה חודשים.
- השאלה אם בינה מלאכותית עתידה להחליף את המבקר האנושי חדלה להיות ספקולציה עתידית והפכה לשאלה מעשית שמוסדות הביקורת בעולם כבר מתמודדים עימה. מאמר זה מפרק את תהליך הביקורת לרכיביו, בוחן את יכולת הבינה המלאכותית לבצע כל אחד מהם ומציג תרחישי שימוש ממוסדות ביקורת מובילים בעולם - GAO האמריקאי, NAO הבריטי ו-ANAO האוסטרלי.
- ייחודו של המאמר אינו במיפוי הטכנולוגיה, כי אם בהיפוך שאלת המוצא. הדיון הציבורי עוסק בשאלה כיצד תוכל הבינה המלאכותית להשתלב לצד המבקר האנושי ולהעצים את עבודתו. המאמר טוען שזו שאלה זהירה מדי, כמעט מתחמקת.
- השאלה הנוקבת היא אחרת: האם יש בכלל הצדקה להותיר את הביקורת בידיים אנושיות? כיום, חלקים מן המגזר הציבורי אינם מבוקרים כלל. הדגימה מצומצמת, הבדיקה נעשית רק בדיעבד, והדוחות מתפרסמים לעיתים שנים אחרי האירוע. יוצא אפוא שגם המבקר המיומן ביותר סורק רק חלק זעיר ממה שראוי להיבדק. מכונה יכולה לשנות כל אחד מהפרמטרים הללו - היקף, תדירות וזמן תגובה: מדגימה אל כיסוי מלא, מרטרוספקטיבה אל ביקורת רציפה ומדוחות שמגיעים לאחר שנים אל תוצר שמגיע בתוך ימים.
- אם תכלית הביקורת היא לשרת את הציבור, הרי שההיצמדות לגורם האנושי עלולה להתגלות לא כנאמנות לערך הביקורת אלא כפגיעה בו.
- מוסדות הביקורת שיקדימו להיערך לא רק ישמרו על הרלוונטיות שלהם - הם יגדירו מחדש את הסטנדרט שלפיו תימדד ביקורת המדינה בעשורים הבאים.

מבוא | מיפוי פעולות הביקורת ובחינתן ביחס ליכולת הבינה המלאכותית | יישומי הבינה המלאכותית כיום | תרחישי שימושי AI בעבודת הביקורת בפועל | ביקורת על גופים מבוקרים המשתמשים בבינה מלאכותית | מצ'ט לסוכן: רמת האוטונומיה של הכלי | התהליך מול התוצר: האינטרס הציבורי בקיום ביקורת אנושית | האם הבינה המלאכותית תייתר את עובדי הביקורת? | סיכונים ואתגרים | ממתודולוגיה ליישום: מפת דרכים | מתווה אופרטיבי לשילוב בינה מלאכותית בעבודת הביקורת | סיכום

מבוא - העין הפקוחה

נוספות ניתנות לאוטומציה על ידי רובוטים פיזיים. אלו אינן שתי סטטיסטיקות נפרדות, אלא שני חלקים של אחד, המשקפים יחד את הפוטנציאל הכולל של האוטומציה.

התמונה אינה שונה כשבוחנים את עולם הביקורת. זוהי גם תמונת המצב שעולה מדוח Anthropic⁵ ממרץ 2026 על השפעות AI על שוק העבודה⁶. על פי דוח זה, תפקידי משרד ומינהלה, מקצועות עסקיים-פיננסיים ומקצועות משפטיים - שאליהם משתייכים גם תחומי הביקורת והחשבונאות - משתייכים לקטגוריות שלהן חשיפה תיאורטית של 80% - 94%. כך עולה גם מדוח IIA Global Internal Audit Standards 2024⁷, הטוען כי תפקידים הכרוכים בניתוח מידע, הערכת ציות, בדיקת נתונים כספיים ועריכת דוחות נמנים עם המקצועות בעלי החשיפה הגבוהה ביותר לאוטומציה. אלה בדיוק היכולות שדורשת ליבת עבודת המבקר.

בד בבד, המגזר הציבורי, בעיקר בארצות הברית, בבריטניה ובאוסטרליה, כבר מאמץ בינה מלאכותית בקצב מואץ. סקר של ה-National Audit Office הבריטי (NAO) משנת 2024 מצא כי 37% מגופי הממשלה כבר פרסו מערכות AI, ו-70% נוספים מתכננים לעשות כך או מפתחים יישומים⁸. בארצות הברית קפץ מספר תרחישי השימוש ב-AI בממשל הפדרלי מ-2,133 ב-2024 ל-3,611 ב-2025 - גידול של כ-70% בשנה אחת בלבד, שבמהלכה הצטרפו 15 סוכנויות נוספות ומספרן הכולל עלה ל-56. מתוך 3,611 התרחישים, 1,818 כבר פרוסים או בפיילוט ו-445 מסווגים כ'השפעה

מהות הביקורת - הן ביקורת המדינה והן הביקורת הפנימית - היא "העין הפקוחה": מנגנון שנועד להבטיח לציבור שמי שמנהל את ענייני עושה זאת כראוי, ביעילות, בחיסכון ובאפקטיביות. בהיסטוריה האנושית המוכרת, העין הזו הייתה תמיד אנושית. מבקרים אנושיים נכנסו לגופים מבוקרים, בחנו מסמכים, ראיינו עובדים, ניתחו נתונים, ניסחו ממצאים והגישו דוחות. מערכת שלמה של ידע, ניסיון ותפיסת עולם מקצועית נבנתה סביב התהליך הזה. השאלה שמאמר זה מעלה אינה אם העין הזו חשובה - היא חשובה מתמיד - אלא אם היא חייבת להיות רק אנושית. המסגרת הנורמטיבית של הביקורת הציבורית מעוגנת בתקן ISSAI 100 של INTOSAI¹ בחוק-יסוד: מבקר המדינה (ס' 2) ובחוק הביקורת הפנימית, התשנ"ב-1992.

מכון המחקר העולמי של מקינזי - MGI (McKinsey Global Institute)² פרסם בנובמבר 2025 דוח שטלטל את עולם התעסוקה וקבע כי 57% משעות העבודה במקצועות הקוגניטיביים והמשרדיים בכלכלה האמריקאית ניתנות לאוטומציה באמצעות הטכנולוגיה הקיימת היום³. אין מדובר בתחזית עתידית או בהערכה תיאורטית. מדובר ביכולת טכנית שכבר קיימת, כאן ועכשיו. אך כדי להבין נכון את הנתון הזה, יש להביט בפילוח שמציג הדוח⁴: 44% משעות העבודה ניתנות לאוטומציה על ידי סוכני בינה מלאכותית (להלן גם - AI) - תוכנות אוטונומיות שמבצעות משימות קוגניטיביות ודיגיטליות. 13% שעות

1 International Organization of Supreme Audit Institution (INTOSAI), ISSAI 100 - Fundamental Principles of Public-sector Auditing (2019), https://www.intosai.org/fileadmin/downloads/documents/open_access/ISSAI_100_to_400/issai_100/ISSAI_100_EN.pdf תקן בין-לאומי המגדיר את עקרונות היסוד של הביקורת במגזר הציבורי כביקורת עצמאית, אובייקטיבית, ממלכתית ומלווה באחריותיות.

2 דו"ע המחקר של חברת הייעוץ McKinsey & Company שפרסומיה משרתים מקבלי החלטות במגזר הציבורי ובמגזר הפרטי.

3 McKinsey Global Institute, *Agents, Robots, and Us: Skill Partnerships in the Age of AI* (2025), <https://www.mckinsey.com/mgi/our-research/agents-robots-and-us-skill-partnerships-in-the-age-of-ai>

4 שם, בעמ' 14. הנתון מתייחס ליכולת הטכנית בלבד ולא לתחזית ההחלפה בפועל.

5 Anthropic היא חברת בינה מלאכותית מובילה, מפתחת משפחת מודלי השפה Claude.

6 Anthropic, *Labor Market Impacts of AI: A New Measure and Early Evidence* (2026), <https://cdn.sanity.io/files/4zrzovbb/chrome-extension://efaidnbmnnnibpcajpcglclefindmkaj/https://cdn.sanity.io/files/4zrzovbb/website/2b5bbaf2c1eb81dbf6e6fb813c1a24e35a64d376.pdf>

7 The Institute of Internal Auditors (IIA), *Global Internal Audit Standards* (2024), theiia.org/en/standards/2024-standards/global-internal-audit-standards אפקטיבי מ-9 בינואר 2025.

8 National Audit Office (UK), *Use of Artificial Intelligence in Government*, HC 612 (2024), nao.org.uk/reports/use-of-ai-in-government. הסקר נערך בקרב 87 גופי ממשלה בבריטניה. ראוגם: House of Commons Committee of Public Accounts, *Use of AI in Government*, HC 356 (2025), <https://publications.parliament.uk/pa/cm5901/cmselect/cmpubacc/356/report.html>

לשינוי אלא אף להוביל אותו. מוסד ביקורת המדינה הוא בעל סמכות רחבה ויכולת השפעה על הממשל כולו ויכול להיות חלוץ בתחום זה. משכך, הרלוונטיות של הדיון הזה לישראל היא מיידית ומעשית.

מאמר זה מבקש לבחון את הטענה שעבודת הביקורת ניתנת לאוטומציה - לא כספקולציה עתידנית, אלא כמציאות טכנולוגית שמתהווה מול עינינו ושראוי כי אנשי מקצוע בתחום הביקורת יכירו אותה לעומקה. המאמר יפרק את תהליך הביקורת לרכיביו, יבחן את מצב הטכנולוגיה, יציג תרחישי שימוש אמיתיים ויטען כי דווקא בתחום הביקורת שימור האדם במוקד התהליך עלול להיות מנוגד לערך שהביקורת אמורה לשרת ולא לשמור עליו.

מה באמת עושה מבקר? מיכוי פעולות הביקורת ובחינתן ביחס ליכולות הבינה המלאכותית

תהליך הביקורת - הן ביקורת המדינה והן ביקורת פנימית - ניתן לפירוק לתשעה רכיבי יסוד המשקפים את מבנה תהליך הביקורת כפי שהוא מעוגן בתקני INTOSAI ISSAI (Guidelines on Central Concepts for Performance Auditing) ובתקני IIA Global Internal Audit Standards (2024), הכוללים שלבי תכנון, ביצוע, דיווח ומעקב. כל רכיב מייצג סוג שונה של פעילות אינטלקטואלית, וניתן לבחון את יכולתה של הבינה המלאכותית לבצע כל אחד מהם בנפרד.

תכנון

שלב התכנון מתחיל בהערכה עמוקה של הנושא המבוקר: מהו תחום הפעילות, מהם הסיכונים הטבועים בו, איפה נדרשת העמקה מיוחדת והיכן יש להשקיע את משאבי הביקורת. זהו שלב שדורש הבנה של סביבה רגולטורית, זיהוי רבדים טעונים (חוקיים, פוליטיים וחברתיים) ותיעודף רציונלי של המשאבים המוגבלים של צוות הביקורת.

גבוהה⁹. המשמעות היא שהגופים שמסודות הביקורת אמורים לבקר הופכים בעצמם לגופים מונעי AI. מוסד ביקורת שאינו שולט בטכנולוגיה הזו לא יוכל לבקר את מי שמשתמש בה, שכן לא יוכל להעריך את אמינות התוצרים של המערכות המבוקרות, לאתר הטיות אלגוריתמיות או לזהות סיכונים ממשי AI ייחודיים.

ישראל בעידן הבינה המלאכותית

גם בישראל קיימת מודעות גוברת לחשיבות הנושא. במשרד מבקר המדינה נוציב תלונות הציבור הוקמו אנשים ייעודיים לביקורת הבודקת טכנולוגיות דיגיטליות וסייבר. מבקר המדינה המכהן ר"ח מתניהו אנגלמן, מגלה עדיפות לביקורת טכנולוגית בתיאום בין-לאומי¹⁰. בדוח מבקר המדינה משנת 2024 בנושא המוכנות של ישראל לעידן הבינה המלאכותית נמצאו ליקויים ניכרים בנושא זה, בין היתר היעדר גורם מרכזי המוביל את אסטרטגיית ה-AI בממשל, חוסר תיאום בין יחידות הממשלה וחוסר בהירות במדיניות.

עדות למיקומה המרכזי של ישראל בשילוב AI בעולם הביקורת ניתן לראות בין היתר בכנס בין-לאומי שהתקיים באבו דאבי בפברואר 2026 בהובלת מבקר המדינה בתפקידו כנשיא EUROSAI. בכנס הוצג דוח ביקורת בין-לאומי בנושא AI שהכין משרד מבקר המדינה בשיתוף 11 משרדי ביקורת עליונים אירופיים נוספים. על המאמר היו חתומים, בין היתר, מזכ"ל ה-OECD מתיאס קורמן ומבקר המדינה של איחוד האמירויות חומיד עובד אבו-שיבס¹¹.

ישראל היא מדינה המתהדרת בחדשנות טכנולוגית. כך, לפי מדד Tortoise AI Global Index 2024, ישראל מדורגת במקום ה-9 מבין 83 מדינות, היא בעלת אקו-סיסטם עשיר בחברות הזנק (סטארט-אפ) ונהנית מנוכחות טכנולוגית גבוהה גם בעולמות ה-AI (לפי דוח Startup Nation Central לשנת 2024, בישראל מעל 2,500 חברות AI פעילות; חלקן של חברות אלה בגיוסי הון עמד על כ-30% מסך גיוסי הטכנולוגיה). כלומר, ביכולתה של ישראל לא רק להגיב

9 U.S. Office of Management and Budget (OMB), 2025 Federal Agency Artificial Intelligence Use Case Inventory (2025), github.com/ombegov/2025-Federal-Agency-AI-Use-Case-Inventory.

10 משרד מבקר המדינה דוח ביקורת מיוחד - ההערכות הלאומיות בתחום הבינה המלאכותית (2024) EUROSAI IT Working, library.mevaker.gov.il/sites/DigitalLibrary/Documents/2024/2024.11-Cyber/2024.11-Cyber-107-AI.pdf; Group, ITWG Newsletter (2024), <https://euosai-it.org/news/newsletter>. (משפט כלכלי) עקרונות, מדיניות, רגולציה ואתיקה בתחום הבינה המלאכותית 2023 (2023), https://www.gov.il/he/pages/ai_23.

11 משרד מבקר המדינה נוציב תלונות הציבור, "מזכ"ל ה-OECD למבקר המדינה: 'הממשלות חייבות לאמץ את יתרונות הבינה המלאכותית ולמזער את הסיכונים'" (3 בפברואר 2026), www.mevaker.gov.il/newsroom/2026-02-03-eurosai.

12 International Organization of Supreme Audit Institution (INTOSAI), ISSAI 3100: Standard for Performance Auditing (n.d.), <https://www.issai.org/pronouncements/issai-3100-performance-audit-guidelines-key-principles/>. התקן מגדיר עקרונות לביקורת ביצועים שניתנים ליישום גם בהקשר של מערכות AI.

מסוגל ללמוד במלואו ולעומקו בתוך זמן קצוב. זהו שלב בעבודת הביקורת שכבר התחוללה בו תמורה בשנתיים האחרונות: ככל שגופים מבוקרים עוברים לדיגיטציה, המידע הופך נגיש יותר לכלים אוטומטיים. סוכן AI המבוסס על מסגרת כמו Microsoft AutoGen¹⁸ או LangChain Agents ומחובר למערכות המידע של הגוף המבוקר יכול לשאוב חומרים בהיקף ובמהירות שאין להם תקדים - אלפי מסמכים, מיליוני רשומות ושנים של נתונים בתוך שעות ספורות בלבד.

קביעת סרגל

הסרגל - הנורמה שעליה יתבססו הממצאים - הוא התשתית לקביעת ממצאי הביקורת, והוא עשוי לכלול מצע קשיח של חיקוקים, פסקי דין והחלטות שיפוטיות, נהלים, הוראות רגולטוריות, הנחיות מדיניות ותקנים מקצועיים בין-לאומיים. התשתית האמורה עשויה לכלול גם נורמות רכות, בין היתר מספרות המחקר ומדוחות מקצועיים. קביעת הסרגלים ויישומם דורשים בקיאות במקורות ובתימוכין ולימוד תכנים רלוונטיים לנושא הנבדק באופן מעמיק.

במסגרת ביצוע שלב זה בעבודת הביקורת עשוי להיות לבניה המלאכותית יתרון מוכח על הבניה האנושית. מודלי שפה גדולים כגון Claude¹⁹ Opus 4.7, Gemini 3.1²⁰ Pro ו-GPT-5.5²¹, מסוגלים לסרוק, לנתח ולהצליב מאות מסמכים נורמטיביים - חקיקה ראשית ומשנית, פסיקה, נהלים פנימיים ותקנים בין-לאומיים כדוגמת תקני ISSAI²² - ולהפיק מהם סרגל ביקורת מדויק ומעודכן בתוך דקות. יתרה מכך, הכלי יכול להתעדכן בזמן אמת בשינויי חקיקה ורגולציה, יכולת שמבקר אנושי לא תמיד נהנה ממנה.

כלי AI מתקדמים כבר מסוגלים לנתח מאגרי מידע נרחבים, לזהות דפוסי סיכון ולהציע תיעודי מבוסס נתונים. כך למשל, Thomson Reuters CoCounsel Audit מקצר תהליכי סקירת workpapers בעד 70%,¹³ MindBridge Ai Auditor מאפשר בדיקה של 100% מהעסקאות בעת ניתוח סיכונים,¹⁴ AuditBoard/Optro הוא סוכן בדיקה לבקורות פנימיות אשר משמש יותר מ-50% מחברות Fortune 500 ו-Microsoft 365 Copilot¹⁵ מציע תסריטים מוכנים לשימוש בתהליכי הביקורת.^{17,16} לעיתים קרובות כלים אלה מציעים כיסוי רחב יותר מזה שצוות אנושי יכול להשיג - כפי שהוכיח ה-GAO עצמו במערכת פנימית שפיתח לצורך סינתזה של דוחות ביקורת קודמים ולסריקת מסמכים של הקונגרס. מערכת AI כגון MindBridge Ai Auditor, שעברה תיקוף על ידי UCLC ועומדת בדרישות תקני SAS 99, ISA ו-CAS 240, יכולה למשל לסרוק את כל דוחות הביקורת מהשנים האחרונות, לזהות נושאים חוזרים שלא טופלו ולהציע תיעודי אוטומטי על בסיס תדירות, חומרה וציר זמן. כך למשל פועל סוכן הבדיקה של AuditBoard/Optro.

איסוף חומרים

הגישה לחומרים המצויים ברשות הגופים המבוקרים היא אתגר מקצועי ראשון במעלה עבור עובדי הביקורת. מדובר בחומר כתוב (מסמכים אקטואליים או ארכיונים) המפוזר במקומות אחסון מגוונים ובמערכות מידע שלא תמיד מכילות רכיבי תיעוד ויצוא נתונים נוחים ונגישים לשימוש בעבודת הביקורת. כדי להשיג את החומרים נשואי הביקורת, על המבקר להשיג הרשאות גישה, לנווט במערכות שונות ולהתמודד עם היקף של מידע שלעיתים צוות אנושי אינו

Thomson Reuters, *CoCounsel Audit - Legal Research and Audit Workpapers*, <https://tax.thomsonreuters.com/en/products/cocounsel-audit> 13

MindBridge AI, *Platform Overview - Finance-Native AI Trained on 260+ Billion Transactions*, mindbridge.ai 14

Optro (formerly AuditBoard), *AI Platform for Internal Controls* (2026), optro.ai 15

Microsoft, *Microsoft 365 Copilot - Audit and Risk Use Cases*, adoption.microsoft.com/en-us/scenario-library/ 16

U.S. Government Accountability Office (GAO), *Artificial Intelligence: GAO's Work to Leverage Technology and Ensure Responsible Use* 17
GAO-24-107237 (2024), gao.gov/products/gao-24-107237. GAO מתאר במסמך את פריסת מודל השפה הגדול ככלי לסינתזה של דוחות עבר וסריקת מסמכי קונגרס.

Microsoft Research, *AutoGen: Multi-Agent Conversation Framework*, microsoft.github.io/autogen 18

Anthropic, *Claude Opus 4.7 - 1M Token Context Window* (April 2026), platform.claude.com/docs/build-with-claude/context-windows. 19

Google, *Gemini 3.1 Pro - Long Context Window of 1M Tokens* (April 2026), ai.google.dev/gemini-api/docs/long-context. 20

OpenAI, *Introducing GPT-5.5* (April 23, 2026), openai.com/index/introducing-gpt-5-5. 21
טוקנים.

ISSAI (International Standards of Supreme Audit Institutions) הם תקנים בין-לאומיים של INTOSAI (ארגון הגג של מוסדות הביקורת העליונים בעולם) המגדירים את עקרונות הביקורת במגזר הציבורי. מוסדות הביקורת העליונים החברים ב-INTOSAI כפופים לתקנים אלה.

ניתוח - הפער בין המצוי לרצוי

לצורך זיהוי הפער נדרש תחילה שלב מקדים של קביעת המצב המצוי והעובדות - איסוף ראיות, אימות מסמכים וראיונות - כמעגל ב-ISSAI 3100 Performance Audit Guidelines ובכרך התכנון של IIA Global Internal Audit Standards 2024.

השוואה שיטתית בין המצוי (כיצד פועל הגוף בפועל) לבין הרצוי (הנורמה) מצויה בליבת עבודת הביקורת. הסוכן - המבקר - צריך לאסוף נתונים על מצב העניינים בפועל (דוחות פיננסיים, מדיניות בפועל, החלטות מינהליות והיבטי ביצוע) ולהשוות אותם בשיטתיות לרף ההתנהלות שהציבה הנורמה. זוהי משימה שדורשת זיהוי פערים, הערכה של חשיבותם והשפעתם וניתוח של ההשלכות. תהליך זה מעוגן בתקני IIA (5 תחומים, 15 עקרונות מנחים, 52 תקנים ומעל 200 דרישות) ובתקני ISSAI 3100.²³

כלי בינה מלאכותית כגון Thomson Reuters CoCounsel, MindBridge Ai Auditor, AuditBoard/Optro, Microsoft 365 Copilot ואחרים מסוגלים לצלוח את השלב הזה באופן מיטבי, לרבות השוואת מערכי נתונים גדולים (Big Data), זיהוי חריגות וסטיות, סיווג פערים לפי חומרה ודחיפות ותייעוד שיטתי של הממצאים. סוכן AI מסוגל עקרונית לבדוק כל עסקה, כל חוזה וכל החלטה - לא מדגם, אלא 100% מהאוכלוסייה. יכולת כזו יש לדוגמה למערכת MindBridge, אשר מאומנת על יותר מ-260 מיליארד עסקאות ביותר מ-3,000 מערכות ERP, עם יותר מ-8,000 חוקי GAAP מוטמעים.²⁴ בעת השימוש במערכת הבקרה האנושית נדרשת בעיקר לשם אימות חריגות שהכלי מסמן לבדיקה.

גיבוש ממצאים, מסקנות והמלצות

המדריך לביקורת המדינה²⁵ מגדיר ממצא ביקורת כ"מידע עובדתי רלוונטי הנובע מהשוואה בין הנורמה למצב בפועל, כפי שעולה מראיות הביקורת". ממצא המבטא פער שלילי ביחס לנורמה מחייבת הוא ליקוי. המדריך מבחין בין שני שלבים: השלב הראשון הוא גיבוש הממצא - שלב שבו אוספים את הראיות, בוחנים את ההשוואה לנורמה, את השפעות הליקוי לטווח הארוך ואת השלכות הרוחב שלו

ובדקים אם יש אינדיקציה לבעיה שורשית ומהי חומרת התמשכותו של הליקוי. השלב השני הוא ניסוח הממצא בדוח - תרגום הקביעה העובדתית לטקסט ברור, מאוזן, מנומק ומבוסס ראיות, המשקף את עמדת מוסד הביקורת בלשון מדויקת ובהירה לקורא הסביר. כל אחד מהשלבים דורש כישורים מקצועיים נפרדים: הראשון - שיקול דעת ביקורתי וניתוח עובדתי, השני - מיומנויות ניסוח, עריכה ובהירות. בשלב הניסוח ניתן להיעזר במודלי שפה גדולים כעורכי משנה: יצירת טיוטה ראשונית מתוך הממצאים שגובשו, התאמת רמת הכירוט לקהל היעד, בדיקת עקביות מול דוחות קודמים ואיתור עמימות לשונית. על המבקר האנושי יוטל לשמור על שיקול הדעת המקצועי ועל האחריות לתוכן.²⁶

מודלי שפה גדולים כגון Claude Opus 4.7, GPT-5.5 ו-Gemini 3.1 Pro הם כותבים מיומנים. הם מסוגלים לנסח ממצאים בהתאם לפורמט קבוע, לנמק את האמור על בסיס הנורמה שהופרה, להצליב נתונים עם ממצאים דומים מביקורות קודמות ולהציע המלצות מותאמות. כך, המבקר האנושי יעבור מתפקיד כותב לתפקיד עורך - שינוי שיחסוך זמן רב וישפר את עקביות הדוחות. כלומר, המבקר האנושי יבחן, יאמת וישפר טיוטה שחיברה מערכת AI, וכך יתפנה להתמקד בשיקול הדעת המקצועי ובקבלת ההחלטות, במקום להשקיע במלאכת הניסוח עצמה.

שלב הפנייה - קבלת תשובות הגופים המבוקרים והתייחסות אליהן

לאחר השלמת טיוטת הדוח, המבקר שולח אותה לעיון הגוף המבוקר ולמתן הערות על תוכן הטייטה וסגנונה. הגוף המבוקר מגיב לאמור בדוח, מתייחס לתכנים ומעיר עליהם, ובמקרים רלוונטיים לשיטתו מנסה להוכיח שהתמונה העובדתית שהציגה הביקורת אינה הולמת את המציאות בפועל וממצא הביקורת שגוי, אינו מדויק או חסר ביסוס. במקרים מסוימים עשוי הגוף המבוקר לטעון כי הגם שממצאי הביקורת נכונים לכאורה, הרי שהערות הביקורת והמלצותיה אינן רלוונטיות או שיישומן בפועל אינו מעשי. בהתייחסו לתגובת הגוף המבוקר, על המבקר להעריך את היקף טענותיו ואת מידת נכונותן, מהימנותן, תקפותן והשפעתן על נוסח ממצאי הדוח.

The Institute of Internal Auditors, Global Internal Audit Standards 2024, <https://www.theiia.org/en/standards/2024-standards/global-internal-audit-standards>, אפקטיבי מ-9 בינואר 2025.

MindBridge AI, *Platform Overview - Finance-Native* 24, לעיל ה"ש 15.

25 משרד מבקר המדינה ונציב תלונות הציבור *מדריך לביקורת המדינה* (ללא ציון תאריך), <https://www.mevaker.gov.il/state-audit/roles-work/state-audit-guide>

26 שם. המדריך מגדיר את מתודולוגיית הביקורת הממלכתית ואת שלבי עבודתו של המבקר.

מותאמים לקהלי יעד שונים (תקציר מנהלים בעמוד אחד, תקציר תקשורתי ותקציר מקצועי מעמיק), ייצור פורמטים נגשים (יצירת הסכת [פודקאסט] מהדוח באמצעות מערכות כגון ElevenLabs, הפקת סרטוני וידאו קצרים, הכנת אינפוגרפיקות אינטראקטיביות ועוד), תרגום אוטומטי איכותי לשפות נוספות (ערבית, אנגלית, רוסית ואמהרית) לטובת אוכלוסיות בישראל ובחו"ל וכן ניתוח של הכיסוי התקשורתי ותגובת הציבור ברשתות החברתיות לאחר הפרסום בעזרת כלים כגון Brandwatch ו-Meltwater. הפרסום הוא חלק אינטגרלי מתפקיד הביקורת - דוח שלא נקרא ולא הופנם חוטא לתכליתו הציבורית.²⁸

מעקב ותיקון ליקויים

המעקב אחר יישום המלצות הביקורת הוא תהליך מתמשך ולעיתים סיזיפי. בדרך כלל, מוסדות הביקורת מבצעים בדיקות מעקב ספורדיות על יישום המלצות דוחות הביקורת ותיקון הליקויים שתועדו בהם לאחר פרקי זמן משתנים (החל בשנתיים לאחר פרסום הדוח). אולם עד ביצוע ביקורת המעקב עשויה המציאות בשטח להתעדכן באופן שאינו מתיישב עם ממצאי הדוח המתיישן, ולכן ההמלצות בדוח עלולות להפוך ללא רלוונטיות או שיישומן עלול לפגוע במועילות תפקודו של הגוף המבוקר.

לעומת המעקב האנכרוניסטי, סוכן AI שמחובר למערכות הארגון יכול לבצע מעקב רציף - לא בדיקה תקופתית אחת לשנה, אלא ניטור מתמשך המתריע בזמן אמת על סטיות מתוכנית התיקון. יכולת זו ממומשת כיום בכלים כמו AuditBoard/Optro ו-MindBridge Ai Auditor, המאפשרים ניטור רציף של תהליכים פיננסיים, זיהוי חריגות בזמן אמת והתרעה אוטומטית על סטיות מבקורות מאושרות.²⁹

כפי שעולה מהאמור עד כה וכפי שיורחב להלן - בשנת 2026 סוכני בינה מלאכותית שכבר פותחו בשדה הטכנולוגי הם בעלי יכולת תפעול של כל שלבי עבודת הביקורת, לכאורה ללא מגע יד אדם. לעבודה האנושית תפקיד בניטור כשלים פוטנציאליים ובבקרה על התוצרים, אך גם תפקידים אלה עשויים אולי להתייטר בשנים הקרובות עם התפתחות כלי הבינה המלאכותית.

זהו שלב שמערב אינטראקציה בין אנשים ובין ארגונים, ויש בו כדי לאתגר את הבינה המלאכותית. אך גם בשלב זה ניתן לדמיין מערכת AI שמנהלת את התהליך - גם אם נכון להיום מדובר בתרחיש עתידי המובא כדי להדגים את גבולות הדיון הרגולטורי ולא כיכולת הקיימת באופן מיידי. מערכת כזו תשלח טיטות, תקבל תגובות, תנתח את מידת הרלוונטיות שלהן, תזהה אילו הערות דורשות שינוי מהותי ותנסח את הגרסה הסופית.

ייעוץ מקצועי משלים

לפני סגירת הדוח, המבקר נעזר בייעוץ חיצוני בתחומים הרלוונטיים לממצאים: ייעוץ משפטי לבחינת המסגרת הנורמטיבית והחשיפות, ייעוץ כלכלי וחשבונאי לאומדן השלכות תקציביות, ייעוץ סטטיסטי לאימות מתודולוגיית הדגימה וייעוץ מומחים לנושאים ספציפיים שהדוח עוסק בהם (לדוגמה סייבר, רפואה, הנדסה או סביבה). בשלב זה ל-AI תרומה פוטנציאלית ניכרת: מודלי שפה גדולים מסוגלים לאתר פסיקה וחקיקה רלוונטית בתוך דקות (בכלים כגון Thomson Reuters Westlaw Precision AI או Lexis+), לסכם חוות דעת מומחים מקבילות מתוך מאגרים בין-לאומיים, להציג זוויות מקצועיות נוספות שמבקר אנושי עלול להחמיץ ולהכין מסמכי רקע לדיוני ועדת הביקורת. השלב מסתיים בהטמעת תובנות הייעוץ בדוח עצמו ובחיזוק הביסוס המקצועי של הממצאים.²⁷

השלב האחרון הוא עריכת לשון, עיצוב והכנה לפרסום, תחומים שבהם הבינה המלאכותית כבר פועלת ביעילות רבה ברחבי העולם, מסיכום ועריכה ועד ליצירת מסמכים מובנים בפורמטים מגוונים, כולל גרפיקה, טבלאות ותרשימים.

הפרסום והשיווק של הדוח

לאחר השלמת הדוח הסופי, מוסד הביקורת אחראי לפרסומו ולתרגומו למסרים שיגיעו לקהלים השונים: תקשורת, מקבלי ההחלטות בכנסת ובממשלה, הציבור הרחב, הקהילה המקצועית והגוף המבוקר. תרומת ה-AI בשלב זה יכולה לכלול יצירה אוטומטית של תקצירים

Thomson Reuters, *Westlaw Precision with CoCounsel - AI-Assisted Legal Research*, thomsonreuters.com; LexisNexis, *Lexis+ AI 27* Platform, lexisnexis.com. כלי הייעוץ של AI נסקרים גם ב-*International Bar Association, AI for Legal Professionals: A Practitioner's Guide* (2025), www.ibanet.org/document?id=IBA-Legal-Agenda-2025.

ElevenLabs, *AI Voice Generation Platform*, elevenlabs.io; Brandwatch, *Consumer Intelligence Platform*, brandwatch.com; Meltwater, 28 INTOSAI-P 12, *The Value and Benefits of the Audit*, meltwater.com. על חשיבות הפרסום ככלי ביקורת ראו *Media Intelligence and Social Listening*, meltwater.com. *of Supreme Audit Institutions* (2019), <https://www.issai.org/pronouncements/intosai-p-12-the-value-and-benefits-of-supreme-audit-institutions-making-a-difference-to-the-lives-of-citizens/>.

MindBridge AI, *Continuous Monitoring and Anomaly Detection* (2025), mindbridge.ai/blog/modernizing-internal-controls-over- financial-reporting-with-ai; AuditBoard, *Optro AI Agent Platform*, auditboard.com

יישומי הבינה המלאכותית כיום

כדי להבין את עומק השינוי שעומד בפתח, יש לבחון שלושה מחקרים מרכזיים שפורסמו לאחרונה. כל אחד מהם מאיר היבט אחר של המציאות הטכנולוגית, ויחד הם מציירים תמונה שלא ניתן להתעלם ממנה.

הפוטנציאל הטכני

כאמור, דוח של מכון המחקר העולמי של מקינזי מציג במאמר בשם "Agents, Robots, and Us: Skill Partnerships in the Age of AI" ניתוח שיטתי של פוטנציאל האוטומציה בשוק העבודה האמריקאי. הממצא המרכזי המובא בו הוא כי 57% משעות העבודה ניתנות לאוטומציה כבר באמצעות הטכנולוגיה הקיימת. סוכני AI (תוכנות שמבצעות משימות קוגניטיביות באופן אוטומטי) יוכלו להחליף 44% מהשעות, ורובוטים (מערכות פיזיות) יוכלו להחליף 13% מהן³⁰. כ-43% משעות העבודה דורשות ביצוע משימות פיזיות, אינטראקציה חברתית-רגשית ושיקול דעת מקצועי מורכב, ולפי הדוח הן ימשיכו לדרוש מעורבות אנושית מרכזית בעתיד הנראה לעין. הדוח אומנם מדגיש שמדובר ביכולת טכנית של הבינה המלאכותית ולא בהחלפת עובדים בפועל - אבל הפער בין היכולת לביצוע הולך ומצטמצם עם כל רבעון שחולף.

דוח נוסף של מכון המחקר העולמי של מקינזי, "A New Future of Work: The Race to Deploy AI and Raise Skills in Europe and Beyond", מוסיף נתונים אירופיים ונותן אומדן זמנים לשינוי באירופה: עד 2030, כ-30% משעות העבודה באירופה יעברו אוטומציה, עם צורך ב-12 מיליון שינויי תעסוקה - הכפלה של קצב השינויים לעומת העשור שלפני מגפת הקורונה³¹.

הפער בין הפוטנציאל למציאות

על אף האמור לעיל, תמונת המצב בשטח שונה. בסקר

ה-NAO הבריטי על השימוש ב-AI בממשלה דיווחו 37% מתוך 87 הגופים הממשלתיים שנסקרו כי הם משתמשים ב-AI³², ו-70% מהנתרים היו בשלבי הפיילוט או התכנון³³. בארצות הברית קפץ מספר יישומי ה-AI המשמשים בסוכנויות פדרליות מ-571 ב-2023 ל-1,110 ב-2024³⁴. אם 57% משעות העבודה ניתנות לאוטומציה, מדוע זה לא קורה כבר עכשיו? מחקר פורץ דרך של Anthropic ממרץ 2026 "Labor Market Impacts of AI: A New Measure, and Early Evidence", מספק לשאלה זו תשובה מדויקת ומפתיעה. המחקר מציג לראשונה מדד חדש: "חשיפה נצפית" (observed exposure). מדד זה אינו בודק מה הבינה המלאכותית מסוגלת לבצע באופן תיאורטי, אלא מהן יכולותיה האקטואליות, הנקבעות על בסיס ניתוח מיליוני אינטראקציות אמיתיות.

הנתונים חושפים פער דרמטי בין הפוטנציאל למציאות: בתחום המחשוב והמתמטיקה, 94% מהמשימות ישימות תיאורטית בעזרת AI, אך רק 33% מבוצעות כך בפועל. בתחומי המשרד והמינהלה מדובר ב-90% מול כ-25%. בתחומים עסקיים ופיננסיים - 85% מול כ-20% ובתחום המשפטי - 80% מול כ-15%³⁵. הפער הזה אינו חולשה של הטכנולוגיה - הוא חלון זמן. חלון זה הולך ונסגר ככל שהכלים משתפרים וככל שהמחירים יורדים כתוצאה מעלייה ביעילות העיבוד של מודלי השפה לצד ירידה בעלות הקריאה ל-API. כך למשל, מחיר עלות שאילתה אחת מהמודל, הנמדד לפי מספר הטוקנים שעוברים בקלט ובפלט של מודלי שפה, ירד ביותר מ-90% במהלך 2023 - 2024.

מחקרים אחרונים מצביעים על כמה גורמים עיקריים לאי-הטמעת כלי AI בפרקטיקה הארגונית המצויה אף שהטכנולוגיה מאפשרת זאת:

א. חשש מחוסר דיוק - ארגונים רבים מדווחים על חשש מחוסר דיוק אשר מעכב הטמעת AI. אם AI עלול לטעות, גם

30 שם; McKinsey Global Institute, *Agents, Robots, and Us*, לעיל ה"ש 3.

31 McKinsey Global Institute, *A New Future of Work: The Race to Deploy AI and Raise Skills in Europe and Beyond* (2024), mckinsey.com/ mgj. הדוח מצביע על 12 מיליון מעברי תעסוקה נדרשים באירופה עד 2030, מהלך שמהווה הכפלה של קצב המעברים בעשור שלפני מגפת הקורונה (COVID-19, 2020 - 2022).

32 House of Commons, National Audit Office (UK), *Use of Artificial Intelligence*, לעיל ה"ש 8. הסקר נערך בקרב 87 גופי ממשלה בבריטניה. ראו גם Committee of Public Accounts, *Use of AI in Government*, לעיל ה"ש 8.

33 National Audit Office (UK), *Use of Artificial Intelligence*, לעיל ה"ש 8.

34 U.S. Federal, *AI Use Case Inventory*, לעיל ה"ש 9. נתוני השוואה בין דיווחי 2023 ל-2024 מבוססים על הרישום הפדרלי הציבורי של יישומי AI בסוכנויות. ראו גם U.S. Government Accountability Office (GAO), *AI Use Case Inventory*, לעיל ה"ש 10.

35 Anthropic, *Labor Market Impacts of AI*, לעיל ה"ש 6. הנתון מבוסס על ניתוח שימוש בפועל (observed exposure) ולא על הערכה תיאורטית בלבד. המדד "חשיפה נצפית" מודד במפורש מה AI עושה בפועל ולא מה הוא יכול לעשות.

במסגרת "שותפויות כישרוים" (Skill Partnerships), כפי שמגדיר McKinsey: התכלית אינה ייתור העובדים אלא שינוי תמהיל המיומנויות, כך שלעובדים יש תמריץ אמיתי להכשיר עצמם לתפקידים מועצמי. AI עם גופים אלה נמנים יחידות ביקורת פנים וביקורת המדינה בשירות הציבורי בחלק ממדינות המערב (בין היתר בבריטניה, ארצות הברית ואוסטרליה, כמתועד בדוחות ANAO ו-NAO, GAO המצוטטים בהמשך). יחידות אלה הן שמרניות מטבען ולכן זהירות מאוד במעבר לשיטות עבודה חדשות הדורשות השקעה משמעותית בהדרכה ושינויי פרדיגמות ברמה הארגונית.

הפער בין הפוטנציאל התיאורטי של ה-AI לבין שיעור האימוץ בפועל מתועד היטב בנתוני דוחות שפרסמה חברת לינקדאין⁴⁰. בדוח "Work Change Report: AI Is Coming to Work"⁴¹, שפרסמה החברה בינואר 2025, נמצא כי 88% מהבכירים הגלובליים רואים בהאצת אימוץ ה-AI יעד מרכזי. אולם, דוח שוק העבודה של החברה מינואר 2026, "Building a Future of Work That Works"⁴², מראה שרק כ-3% מחברי הרשת בארצות הברית מציגים כישורי AI בפרופיל שלהם, והם מרוכזים בעיקר במקצועות ההנדסה, ניהול המוצר והמחקר. כלומר, הציפייה הניהולית מתקדמת מהר יותר מהיכולת הארגונית לספק אותה - וזהו הקושי האמיתי שניצב בפני יחידות ביקורת המבקשות להטמיע את הטכנולוגיה. החזון הניהולי לאימוץ AI מתקדם יותר מהמבנה הארגוני הנדרש כדי לתמוך בו.

כל אלה הם אתגרי עומק שאינם טכנולוגיים ונוגעים לתרבות ארגונית, לניהול סיכונים, לנשיאה באחריות עמוקה לגורל הארגון ולהטיות קוגניטיביות של פרטים ושל קבוצות

אם בתדירות קטנה יותר מבני אדם, הרי שגורמי המקצוע שונאי הסיכון יהססו לשאת באחריות למימוש הסיכון. זהו פרדוקס מוכר: החברה מגלה סובלנות רבה יותר לטעויות אנוש מאשר לטעויות מכונה - תופעה הידועה בספרות המחקר כ"סלידה אלגוריתמית" (algorithm aversion)³⁶. הדוגמה של הרכב האוטונומי ממחישה זאת היטב: הציבור דורש מרכבים אוטונומיים סטנדרט בטיחות חמור בהרבה מזה הנדרש מנהגים אנושיים, גם אם הטכנולוגיה מפחיתה סיכונים ביחס לנהיגה אנושית. תאונה שגרמה לה מכונה זוכה לתשומת לב ציבורית ולתגובה פסיכולוגית בעוצמה גבוהה בהרבה מתאונה זהה במאפייניה שגרם לה נהג אנושי³⁷.

ב. אי-ודאות רגולטורית - בסקר Deloitte לרבעון 4 של 2024, שכלל 2,773 מנהלים בכירים מ-14 מדינות, בלט הציות הרגולטורי כחסם המוביל העומד בפני הטמעת טכנולוגיה - שיעור המנהלים שציינו אותו כחסם עלה מ-28% ברבעון 1 ל-38% ברבעון 4³⁸. החסם הזה מורגש במיוחד במוסדות הביקורת הציבוריים: GAO האמריקאי, NAO הבריטי ו-ANAO האוסטרלי - כולם פועלים בסביבה רגולטורית של AI שעודנה מתעצבת, ולכן נוקטים זהירות יתרה בהטמעת AI. חסם רגולטורי נוסף ונפרד מתעצב סביב שאלות של פרטיות, אתיקה, הטעיות ודיפ-פייק - חסם המתגבש במקביל לכניסת EU AI Act לתוקף ולפרסום NIST AI Risk Management Framework³⁹. כלומר, לרשויות ממשלתיות יש סיבה טובה לנהוג בזהירות יתרה.

ג. פערי כישרים וידע - המנהלים שנסקרו בסקר Deloitte מדווחים על פערים בידע וביכולות בתחום ה-AI המונעים הטמעת כלים מתקדמים בארגונים. אתגר זה נפתר רק

Berkeley J. Dietvorst, Joseph P. Simmons & Cade Massey, *Algorithm Aversion: People Erroneously Avoid Algorithms After Seeing Them* 36 *Err.*, 144(1) *J. of Experimental Psych. General* 114 (2015).

Azim Shariff, Jean-François Bonnefon & Iyad Rahwan, *Psychological Roadblocks to the Adoption of Self-Driving Vehicles*, 1 (10) *Nature* 37 (2017).
Azim Shariff, Jean-François Bonnefon & Iyad Rahwan, *How Safe Is Safe Enough?* *Psychological Human Behaviour* 694 (2017).
Mechanisms Underlying Extreme Safety Demands for Self-Driving Cars, 126 *Transportation Rsch. Part C* (2021), art. 103069.

Deloitte AI Institute, *The State of Generative AI in the Enterprise: Generating a New Future* (Quarter 4 Report 2025) 38
www.deloitte.com/content/dam/assets-zone3/us/en/docs/services/consulting/2024/us-state-of-gen-ai-q4.pdf.
הסקר נערך ביולי ספטמבר 2024 בקרב 2,773 מנהלים בכירים ב-14 מדינות.

Regulation (EU) 2024/1689 Laying Down Harmonised Rules on Artificial Intelligence (AI Act), OJ L of 12 July 2024 39
(entered into force 1 August 2024).
התקנה האירופית לרגולציה של בינה מלאכותית, האוסרת את השימוש ב-AI למטרות דיורג חברתי, בעלת דרישות שקיפות ובקרה אנושית על מערכות בסיכון גבוה.

40 הרשת המקצועית הגדולה בעולם שבה יותר מ-1.3 מיליארד חברים. הרשת מפעילה מכון מחקר (Economic Graph Research Institute) המדווח על מגמות בשוק העבודה.

LinkedIn Economic Graph Research Institute, *Work Change Report: AI Is Coming to Work* (2025), economicgraph.linkedin.com 41

LinkedIn Economic Graph Research Institute, *Labor Market Report: Building a Future of Work That Works* (2026) 42
economicgraph.linkedin.com



בארגונים ובמערכות שטרם בשלו להטמעת שינויים משמעותיים בקרבם^{43, 44}.

ממצא נוסף מהמחקר של Anthropic ממרץ 2026 ראו לתשומת לב מיוחדת: עדיין אין עלייה מובהקת באבטלה בקרב עובדים בתחומים חשופים (תחומים שבהם יותר מ-80% מהמשימות ניתנות לביצוע תיאורטי על ידי AI לפי המדד של Anthropic - בראשם מקצועות מחשב, מזכירות, עסקים ומשפט), אך ניכרת האטה מובהקת בגיוס של צעירים בני 22 - 25 לתחומים אלה⁴⁵. על פי דוח LinkedIn Labor Market Report מינואר 2026, הגיוס למשרות מפתחים ירד ב-20% - 35% לעומת תקופת טרום הקורונה, בעיקר בשל חוסר ודאות כלכלית לצד שינויים בתמהיל הכישורים הנדרש⁴⁶. כלומר, השוק כבר מגיב וחווה שינוי, אם כי בצורה שקטה: לא פיטורים המוניים, אלא צמצום הדרגתי של גיוס עובדים חדשים לתחומים בעלי חשיפה גבוהה ל-AI. גם מחקר מעבדת Stanford Digital Economy מצביע על ירידה של 13% - 20% בתפקידי entry-level בתחומי הנדסת התוכנה ושירות הלקוחות.

הכפלה כל 89 יום

ארגון METR (Model Evaluation and Threat Research)⁴⁷, שעוקב באופן שיטתי אחרי יכולתם של סוכני AI לבצע משימות מורכבות באופן אוטונומי, מודד זאת באופן שיטתי ומעדכן את "אופק הזמן" (Time Horizon) - כלומר קובע מהו משך המשימה האנושית שסוכן AI מצליח לבצע באופן עצמאי בהצלחה של 50% לפחות. לפי המדידה העדכנית של הארגון (Time Horizon 1.1, ינואר 2026), זמן ההכפלה משנת 2023 ואילך הוא כ-131 יום (לעומת 165 יום תחת המתודולוגיה הקודמת), והוא ממשיך להתקצר - בנתוני 2024 ואילך הוא עומד על כ-89 יום בלבד⁴⁸. כלומר, לא רק שהיכולות גדלות - הן גדלות בקצב שהולך ומואץ. לעיתים קרובות תחזיות בעניין קצב התפתחות הטכנולוגיה מתיישנות בתוך חודשים בודדים.

בשנת 2020 היו מערכות AI מסוגלות לבצע משימות אנושיות שארכו כמה שניות, כגון סיווג פריט אחד, תיקון שגיאת הקלדה או תשובה למשימת שפה קצרה; בתחילת שנת 2023 הן הצליחו במשימות הדורשות מאדם כמה דקות (כגון סיכום מסמך קצר, ניסוח פסקה או פתרון בעיית תכנות בסיסית); בסוף שנת 2024 הן הצליחו במשימות הדורשות מאדם כמה עשרות דקות (כגון סקירת מסמך ארוך, ניתוח טבלת נתונים מורכבת או כתיבת טיוטת סעיף דוח). בהתאם לקצב המואץ של ביצוע המשימות כמתואר, סוכני בינה מלאכותית עתידים לבצע משימות הדורשות שעה מלאה בעוד כחצי שנה, יום עבודה מלא - בעוד כשנה ושבוע עבודה - בעוד כשנתיים.

על פי שלושת המחקרים האמורים, הטכנולוגיה בת זמננו כבר מסוגלת לבצע למעלה מ-28% מהמשימות הכרוכות בעבודת הביקורת באופן עצמאי לחלוטין, ולכן הפער בין היכולת הטכנולוגית ליכולת הביצועית הולך ומצטמצם בקצב מואץ. דהיינו, בעוד שלוש עד חמש שנים תוכל הבינה המלאכותית לבצע את עבודת הביקורת בשלמותה, בתנאי שכל הגופים המבוקרים יטמיעו במערכותיהם כלי בינה מלאכותית מתקדמים בממשק מאובטח עם המערכות שבשימוש משרד מבקר המדינה ונציב תלונות הציבור.

תרחישי שימוש: AI בעבודת הביקורת בפועל

הדיון על יכולות ה-AI אינו תיאורטי בלבד. הן מוסדות ביקורת עליונים ברחבי העולם - Supreme Audit Institutions (SAls), המאוגדים תחת ארגון INTOSAI הבין-לאומי המונה 193 חברים כדוגמת ה-GAO בארצות הברית וה-NAO בבריטניה, והן יחידות ביקורת פנימית כבר מפתחים ומיישמים כלים מבוססי בינה מלאכותית. להלן חמישה תרחישי שימוש מרכזיים, כל אחד מלווה בדוגמה ממשית מגוף ביקורת שכבר התחיל בשלב היישום.

Jin Li, Feng Zhu & Pascal Hua, *Overcoming the Organizational Barriers to AI Adoption*, Harvard Bus. Rev. (2025), 43 <https://hbr.org/2025/11/overcoming-the-organizational-barriers-to-ai-adoption>; Cole Stryker, Amanda Downie & Matthew Finio (42%), הצדקה עסקית חלשה (42%) וחששות לפרטיות או לסודיות (40%).

5 *Biggest AI Adoption Challenges for 2026*, IBM (n.d.), <http://www.ibm.com/think/insights/ai-adoption-challenges>. 44 לפי סטרייקר, דאוני ופיניו, חמשת החסמים הגדולים לאימוץ AI ב-2026 הם: חשש מדיוק או הטיה (45%), מחסור בנתונים קנייניים (42%), חוסר במומחיות Gen AI.

Anthropic, *Labor Market Impacts of AI*, לעיל ה"ש 6. ההאטה בגיוס נצפתה בקרב עובדים בני 22 - 25 בתחומים עם חשיפה גבוהה ל-AI.

LinkedIn Economic Graph, *Labor Market Report* 46 לעיל ה"ש 44. הדוח קובע במפורש: "hiring remains 20-35% below pre-pandemic levels in advanced economies... AI isn't driving the hiring slowdown".

47 מכון מחקר עצמאי המודד שיטתית את יכולות מודלי AI.

METR, *Time Horizon 1.1: Task-Completion Time Horizons of Frontier AI Models* (2026), <https://metr.org/blog/2026-1-29-time-horizon-1-1/>. 48

ניתוח מסמכים רחב-היקף

אחד האתגרים הגדולים בביקורת הוא היקף החומרים: אלפי חוזים, פרוטוקולים, מכתמים ונהלים שיש לבחון. המבקרים האנושיים קוראים חלק מהחומרים האלה - אולם בעידן ה-AI, החומרים רבים כל כך שגם דגימה סטטיסטית נכונה עלולה ליצור הטיה בממצאים. משכך, לעיתים קשה לדעת אם אכן נבחן מדגם מייצג.

ה-GAO האמריקאי (Government Accountability Office) מפעיל מערכת המבוססת על מודל שפה לסינתזה של דוחות ביקורת קודמים, סיוע בסקירות עריכתיות וסריקת מסמכי קונגרס⁴⁹. פעילות שדרשה בעבר הקצאת משאבי זמן וארכה שבועות ארוכים של קריאה, ניתוח ועיבוד של מידע כתוב נעשית כעת בתוך שעות, עם כיסוי מלא של כל החומר הרלוונטי. סקירת הבינה המלאכותית מבססת את היכולת של הביקורת לבנות עבור הציבור תמונת מצב עדכנית ומדויקת של המציאות הנבדקת בגופים המבוקרים.

זיהוי חריגות ודגלים אדומים

ביקורת פיננסית שמרנית כוללת בדיקה של היקף עסקאות מסויים. הערך המוסף של פעילות הבינה המלאכותית הוא בסריקת כל הנתונים ללא מגבלה וזיהוי דפוסים וחריגות שבינה אנושית נוטה להתעלם מהם או שהייתה מתקשה לזהות בשל ההיקף העצום של הנתונים. יכולות אלה תועדו במחקר אקדמי על data-driven auditing ובמסגרת מסמך הגדרת האחידות של GAO ל-AI במגזר הציבורי⁵⁰.

Australian National Audit Office - מוסד הביקורת העליון של אוסטרליה (ANAO) בחן כיצד רשות המיסים (ATO) משתמשת ב-AI לזיהוי הונאות ומצא שהכלים מזהים חריגות שנשמטו ממבקרים אנושיים בשל מגבלות הזמן והיקף הנתונים. סוכן AI שמנתח "כל עסקה" לא רק עושה זאת במהירות רבה יותר - הוא מבטל לחלוטין את מגבלת המשאבים, כלומר שעות העבודה והתקציב של מוסד

הביקורת, שבהכרח מוגבלים יחסית להיקף הפעילות המבוקרת⁵¹.

ניסוח ממצאים והמלצות

פרסום ממצאי הביקורת לציבור מצריך ניסוח ברור של ממצאים בלשון ברורה המשקפת בסיס עובדתי מוצק, התאמה לנומרה והמלצה לתיקון. המדריך לביקורת המדינה מורה כי "ממצא יתואר באופן אובייקטיבי ותמציתי ללא פרשנות של צוות הביקורת או הצגת עמדתו בנושא"; כי יש לתאר אותו "כהשוואה בין המצב הרצוי למצב המצוי" ולנסחו כך ש"יעמוד בפני עצמו"; וכי בתיאורו "ישמשו המילים 'בביקורת עלה כי...' או 'נמצא כי...'". תיאור ממצא הוא, אפוא, משימת כתיבה שדורשת חדות, בהירות ועמידה בכללים מתודולוגיים מדויקים⁵².

מודלי שפה גדולים מסוגלים לעבד נתונים גולמיים ולנסח על בסיסם ממצאים מובנים - בהתאם לפורמט, לסגנון ולרמת הפירוט של דוחות ביקורת קיימים. כלי בינה מלאכותית כגון Thomson Reuters CoCounsel Audit ו-Wolters Kluwer CCH Xcess Expert AI⁵³ מסוגלים אף ליצור טיוטה ראשונית של ממצא, לנמק אותו ביחס לסרגל הביקורת, להציע המלצות ואף להצליב את הנתונים עם ממצאים דומים מביקורות קודמות. כך נוצרת סינתזה בין דוחות ביקורת מזמנים שונים המעשירה את הממצאים ומעניקה ערך מוסף לציבור. הבינה המלאכותית יכולה לסייע גם בניסוח ובתרגום לשפות שונות, שהם אתגר מובהק בישראל לאור חובת הפרסום בעברית, ערבית ואנגלית.

הערכת סיכונים ותיעודם

תכנון ביקורת מצריך הערכה מדורגת של סיכונים: באילו תחומים עלולים להתרחש כישלונות, אסונות, תאונות, או הפרות חוק משמעותיות. המבקר צריך למפות את סדרי העדיפויות בתוך מנעד אפשרויות, לשקול את השיקולים

49 Artificial Intelligence (GAO), U.S. Government Accountability Office, לעיל ה"ש 18. GAO מתאר במסמך את פריסת מודל השפה הגדול ככלי לסינתזה של דוחות עבר, סיוע בסקירות עריכתיות וסריקת מסמכי קונגרס.

Deniz Appelbaum, Alexander Kogan & Miklos A. Vasarhelyi, *Big Data and Analytics in the Modern Audit Engagement*: 50 U.S. Government Accountability Office *Research Needs*, 36(4) Auditing: A Journal of Practice & Theory 1 (2017); (GAO), Artificial Intelligence: An Accountability Framework for Federal Agencies and Other Entities (GAO-21-519SP) (2021), <https://www.gao.gov/products/gao-21-519sp>.

Australian National Audit Office (ANAO), Governance of Artificial Intelligence at the Australian Taxation Office, Auditor-General Report 51 (No. 36 of 2023-24) (2025), www.anao.gov.au/work/performance-audit/governance-of-artificial-intelligence-the-australian-taxation-office.

52 משרד מבקר המדינה ונציב תלונות הציבור, *מדריך לביקורת המדינה*, לעיל ה"ש 26.

Thomson Reuters, *CoCounsel Audit*, thomsonreuters.com/en/products/cocounsel.html; Wolters Kluwer, *CCH Xcess Expert AI*, 53 wolterskluwer.com/en/solutions/cch-xcess-expert-ai.

ביקורת על גופים מבוקרים המשתמשים בבינה מלאכותית

ככל שמערכות הבינה המלאכותית הולכות וחודרות לממשל ולמגזר הציבורי - מקבלות החלטות על זכויות, מנהלות מכרזים, מזהות הונאות ומספקות שירותים לתושבים - מוסדות הביקורת נדרשים לפתח יותר ויותר יכולות חדשות לבדיקת מערכות אלה.

הקופסה השחורה

מערכות AI - ובפרט מודלים עמוקים ומודלי שפה גדולים - מתוארות לעיתים קרובות כ"קופסה שחורה": הן מקבלות קלט ומפיקות פלט, אך התהליך הפנימי אטום ולעיתים גם מפתחיו אינם מסוגלים להסבירו. עבור מבקר זהו אתגר מהותי. כיצד מבקרים מערכת שאפילו מפתחיה אינם מסוגלים תמיד להסביר את החלטותיה? כיצד מעריכים אם מערכת AI שמשמשת לקבלת החלטות בנוגע לתושבים - זכאות לגמלה, דירוג סיכון אשראי, אכיפה ברחבי וכו' - פועלת בצורה הוגנת, שוויונית ותקינה?

דוגמאות בין-לאומיות

ה-ANAO האוסטרלי ביצע ביקורת ביצועים על שימוש רשות המיסים (ATO) ב-AI ומצא כי לרשות לא הייתה מסגרת ממשל AI מספקת וכי פעלה ללא סיווג סיכונים ברור וללא מנגנוני בקרה הולמים.⁵⁸

ה-Bundesrechnungshof הגרמני הלך צעד נוסף קדימה ופיתח מתודולוגיה שלמה לביקורת מערכות אלגוריתמיות, כולל בדיקות שקיפות, הוגנות, אחריותיות ובטיחות. המתודולוגיה הזו היא אבן דרך חשובה עבור מוסדות ביקורת אחרים שמחפשים מסגרת עבודה לנושא שעדיין בחיתוליו.⁵⁹

ה-NAO הבריטי סקר את אופן הפריסה של AI ברשויות

והאיזונים הרלוונטיים בין הצרכים המשתנים על פי האינטרס הציבורי ולהחליט באילו נושאים על משרד להתמקד בתוכניות העבודה השנתיות והרב-שנתיות.

ה-NAO הבריטי דיווח על שימוש בכלי AI לטובת ניתוח נתוני המגזר הציבורי לצורך תיעודף נושאי הביקורת. במקום להסתמך על שיקול דעת בלבד - שמושפע מזמנות מידע, מהטיות קוגניטיביות ומנטייה לחזור לנושאים מוכרים - צוותי הביקורת משלבים ניתוח אלגוריתמי שמצביע על תחומי הסיכון הגבוהים ביותר ומאפשר הקצאת משאבים רציונלית ומבוססת נתונים.⁵⁴

סקירת רגולציה והשוואה בין-לאומית

לשם קיום ביקורת השוואתית ויצירת תשתית לגיבוש ממצאי ביקורת עדכניים, המבקר נדרש במסגרת עבודתו השגרתית להיות בקיא בחקיקה וברגולציה הקיימת הן במדינות מתקדמות והן בארגונים בין-לאומיים. מדובר בכמות בלתי רגילה של חומרים בשפות שונות, אשר מומחיות בתוכניהם מצריכה השקעת משאבים ניכרים. לעומת זאת, שימוש בכלי בינה מלאכותית לצרכים האמורים עשוי לחסוך זמן וממון רבים ולהניב תוצרים יעילים ומועילים.

בשנים האחרונות מתגבשת קהילה בין-לאומית סביב השימוש במערכות בינה מלאכותית לצורכי ביקורת, כפי שעולה משורת יוזמות של מוסדות ביקורת עליונים. ה-NAO הבריטי פרסם שני מסמכי הנחיה: 'How to audit 'Good practice guide-I artificial intelligence models' for organisations using AI';⁵⁵ וה-GAO האמריקאי משתמש במודל שפה פנימי לסיכום דוחות קודמים, סיוע בסקירות עריכתיות וסריקת מסמכי קונגרס;⁵⁶ וה-ANAO האוסטרלי ערך ב-2023 ביקורת על ממשל AI ברשות המיסים (ATO) שכללה שבע המלצות שכולן התקבלו.⁵⁷ במקביל, כלי AI מאפשרים למוסדות הביקורת עצמם לעקוב בזמן אמת אחר שינויי רגולציה בעולם, לבצע השוואות בין-לאומיות ולזהות פערים בין הנורמה המקומית לבין שיטות מיטביות.

54 National Audit Office (UK), Use of Artificial Intelligence, לעיל ה"ש 8.

55 Daniel Lambauer, How to Audit Artificial Intelligence Models, NAO Insight (2021), www.nao.org.uk/insights/how-to-audit-artificial-intelligence-models; National Audit Office (UK), Good Practice Guide for Organisations Using AI (2026), www.nao.org.uk/insights/good-practice-guide-for-organisations-using-ai.

56 U.S. Government Accountability Office, Artificial Intelligence, לעיל ה"ש 18.

57 Australian National Audit Office (ANAO), Auditor-General Report, לעיל ה"ש 52.

58 Governance of AI at ATO, הביקורת מצאה כי ל-ATO לא הייתה מסגרת ממשל AI מספקת.

59 SAs of Brazil, Finland, Germany, the Netherlands, Norway and the UK (incl. Bundesrechnungshof), Auditing Machine Learning Algorithms - A White Paper for Public Auditors (2020 updated 2023, 2025), auditingalgorithms.net

המבקר עדיין מקבל את כל ההחלטות, אך הכלי מבצע חלק ניכר מהעבודה הפרוצדורלית. האנלוגיה המתאימה היא נהיגה עם מערכת עזר: הנהג אחראי, ידיו על ההגה, אך המערכת עושה את רוב מלאכת הנהיגה.

רמה 3: סוכן אוטונומי (Agent) - מבקר דיגיטלי

זהו השדרוג המשמעותי ביותר של כלי הבינה המלאכותית בעת הזאת. סוכן AI שמקבל הנחיה כללית - "בצע ביקורת ציות על מחלקת הרכש ביחס לתקנות המכרזים" - ומבצע את כל התהליך באופן עצמאי: ניגש למערכות המידע, שואב נתונים ומסמכים, מנתח אותם ביחס לנורמות, מזהה פערים וחריגות, מנסח ממצאים ומפיק דוח. המבקר האנושי הופך ממבצע למפקח - קורא את הדוח, מאמת ממצאים מרכזיים ומפעיל שיקול דעת בסוגיות רגישות.

סוכן AI שיוצר קשר עם מערכות הגוף המבוקר, שואב חומרים, מנתח פערים, מגבש ממצאים ומפיק דוח - זהו תרחיש שהטכנולוגיה כבר מאפשרת היום, ובכל כמה חודשים הוא הופך לבשל ואמין יותר. השאלה המעשית אינה אם הדבר אפשרי, אלא מתי הוא יגיע לרמת אמינות שמאפשרת הסתמכות עליו, ומה תידרש המסגרת המבקרת לאפשר.

התהליך מול התוצר: האינטרס הציבורי בקיום ביקורת אנושית

האם לציבור יש אינטרס בכך שעבודת הביקורת תיעשה דווקא בידי אדם? בהקשר זה, יש להבחין בין שני ערכים שונים שהביקורת נועדה לשרת:

הערך שבתהליך מול הערך שבתוצר

מאז הקמתם של מוסדות ביקורת המדינה נתפסה עבודת הביקורת כתהליך אנושי מיסודו. המבקר אינו רק מנתח נתונים - הוא נוכח פיזית, מנהל שיחות, שואל שאלות שלא ניתן לנסח מראש, מזהה אווירה ותרבות ארגונית ומפעיל שיקול דעת מקצועי שמבוסס על ניסיון מצטבר⁶³. הנוכחות האנושית עצמה - עצם הידיעה שהארגון חשוף

הממשל ומצא שרק 30% מהגופים שהטמיעו בעבודתם AI דיווחו על מסגרת ממשל ייעודית. זהו ממצא מדאיג שמצביע על פער ניכר בין קצב האימוץ הנלהב של הטכנולוגיה לבין רמת הבקרה על יישומיה⁶⁰.

חוק ה-AI של האיחוד האירופי (EU AI Act), שנכנס לתוקף ב-2024, מחייב לסווג מערכות AI לפי רמות סיכון ומטיל חובות שקיפות, תיעוד ובקרה על מערכות בסיכון גבוה. גופי הביקורת בעולם יידרשו לציית לדרישות אלה ולפתח מומחיות טכנולוגית שתחייב שדרוג משמעותי של יכולות הצוותים⁶¹. בין היתר יחויבו הצוותים להקפיד על עקרון human-in-command (HIC) - אחת משלוש רמות הפיקוח האנושי שהגדיר ה-High-Level Expert Group on-AI של נציבות האיחוד האירופי, לצד human-in-the-loop (HITL) ו-human-on-the-loop (HOTL)⁶². כלומר, גם כאשר הביקורת נעזרת בכלי AI, שיקול הדעת המקצועי והאחריות לניסוח הסופי נשארים בידי המבקר האנושי.

מצ'ט לסוכן: רמות האוטונומיה של הכלי

לא כל השימושים בבינה מלאכותית שווים. יש ספקטרום רחב בין צ'טבוט פשוט שעונה על שאלות ובין סוכן אוטונומי שמבצע ביקורת מלאה. הבנת הספקטרום הזה חיונית לכל דיון רציני על עתיד הביקורת, שכן כל רמה מציבה דרישות שונות של פיקוח, אמון, תשתית ומסגרת נורמטיבית.

רמה 1: צ'ט - מנוע שאלות ותשובות

בינה מלאכותית ברמה בסיסית זו מתייחסת אך ורק לשאלות המסוימות ששואל המבקר. למשל: "מה קובע חוק חובת המכרזים לגבי פטורים?" או "סכם לי את עשר הנקודות המרכזיות בדוח הזה". זוהי הרמה הנפוצה ביותר היום, והיא כבר חוסכת שעות עבודה רבות של חיפוש ידני. המבקר נשאר בשליטה על התהליך. הכלי משמש מנוע חיפוש מתקדם וידידותי למשתמש.

רמה 2: עוזר (Copilot) - שותף לעבודה

הכלי לא רק עונה אלא גם יוזם: מציע ממצאים, מתריע על חריגות בנתונים ומנסח טיוטות ראשוניות של קטעים בדוח.

National Audit Office (UK), *Use of Artificial Intelligence*, 60 לעיל ה"ש 8. הסקר מצא כי רק 30% מהגופים שמתמשים ב-AI דיווחו על מסגרת ממשל ייעודית.

European Commission, EU AI Act, Regulation 2024/1689, August 2024. 61

European Commission, *Ethics Guidelines for Trustworthy AI*, European Commission (2019), digital-strategy.ec.europa.eu/en/library/ 62 ethics-guidelines-trustworthy-ai. ההנחיות מבחינות בין שלוש רמות פיקוח אנושי: התערבות אנושית בכל מחזור החלטה - HITL; ניטור פעילות המערכת בזמן ריצה ויכולת התערבות בה - HOTL; ושיקול דעת בבחינת המסגרת הכוללת וההחלטה אם ומתי להשתמש במערכת - HIC.

International Organization of Supreme Audit Institution (INTOSAI), *ISSAI 100* 63 לעיל ה"ש 1. העקרונות מדגישים את חשיבות העצמאות, האובייקטיביות ושיקול הדעת המקצועי.

מלאכותית יכול לבצע את אותה עבודה בתוך ימים או שבועות. יתרון היעילות אינו מתבטא רק בקיצור לוחות זמנים: הוא מאפשר ביקורת רציפה (continuous auditing) במקום ביקורת נקודתית, וכך מאפשר לאתר חריגה בזמן אמת ולא שנים אחריה. בגופי ביקורת מתקדמים כמו ה-GAO האמריקני וה-ANAO האוסטרלי הובהר כי היעילות היא יעד מוביל לאימוץ AI, ולא תוצר לוואי שלו⁶⁵. יתרה מכך, יעילות הביקורת האנושית מושפעת מהסתמכות על מדגמים מצומצמים, וכך מטבעה היא אפקטיבית במידה חלקית בלבד: היא מאתרת ליקויים בגורמים שנדגמו אך מחמיצה לעתים את השיטה הכוללת. בינה מלאכותית מאפשרת סריקה של 100% מהעסקאות, 100% מהמסמכים ו-100% מההחלטות, ובכך מעלה את יעילות הביקורת לרמה שלא הייתה אפשרית עד כה. בד בבד, היא מאפשרת זיהוי מדויק יותר של חריגות, לעומת מדגם אנושי הסובל מ"רעש" הנובע מהטיות אישיות (noise) ומשונות קבוצתית, כפי שהוכיחו כהנמן ועמיתיו⁶⁶.

חיסכון: צמצום עלויות והרחבת ממצאי הביקורת

עלות שעת עבודה של מבקר מקצועי במשרד מבקר המדינה נאמדת במאות שקלים, וכוללת משכורת, הכשרה והוצאות תפעול. עלות שעת עבודה של סוכן AI מתקדם, לאחר הפריסה ראשונית, נאמדת באגורות. אין מדובר רק בחיסכון תקציבי - זהו שינוי מהותי באופן הביקורת האפשרי: ניתן לבקר פי עשרה גופים באותו תקציב, לבחון פי מאה מסמכים ולכסות תחומים שלמים שכיום אינם נבדקים כלל בשל מחסור במשאבים. משמעותו של חיסכון כזה אינה ביצוע פחות פעולות, אלא השגת כיסוי נרחב יותר עם אותה השקעה, וזה התרגום המדויק של הערך הציבורי שבחיסכון.

יעילות וחיסכון אינם ערכים המנוגדים לערכי התהליך והתוצר שנדונו לעיל אלא משלימים אותם. ביקורת שאינה יעילה ואינה חסכונית אינה מממשת את תכליתה הציבורית גם אם היא אנושית לחלוטין. ולהפך, ביקורת ממוכנת שממזה את הערכים הללו אינה מוותרת על ערכים מוסריים אלא ממצה אותם באופן שלם יותר. האתגר אינו להכריע בין הבינה האנושית לבינה המלאכותית, אלא לוודא שהאינטרס הציבורי של יעילות וחיסכון לא ייוותר חסר מימוש בשל היצמדות נוסטלגית לדפוסי עבודה אנושיים בלבד.

באופן תמידי לנכחות של גורם חיצוני לו הבודק לעומק את פעילותו, היא מנגנון הרתעה רב עוצמה.

יתרה מכך, האינטראקציה בין המבקר לגורמים בגוף המבוקר היא חלק בלתי נפרד מתהליך ההשפעה: ההסברים, השאלות, חילוקי הדעות, תהליך ההערות והשיח המקצועי - כל אלה מעצבים לא רק את דוח הביקורת אלא את ההתנהגות של הגוף המבוקר עצמו. המבקר, במובן מסוים, הוא לא רק מפקח אלא גם מחנך. מנקודת מבט זו, החלפת המבקר האנושי במכונה היא לא רק שינוי טכני - היא ויתור על ערך מהותי שהציבור מעוניין בו.

מנגד, ניתן לטעון שבמבחן התוצאה, התוצר הסופי מממש את האינטרס הציבורי המובהק: דוח ביקורת מדויק ומקיף המופק בזמן קצר ומניע את הגופים המבוקרים לתקן ליקויים. אם מכונה מסוגלת להפיק דוח טוב יותר - מדויק יותר, מהיר יותר, עם כיסוי רחב ומלא יותר - מדוע צריך הציבור לשלם עבור תהליך אנושי יקר, איטי וחלקי? לפי גישה זו, לזהות הגורם המבצע את הביקורת אין השפעה מהותית כל עוד התוצר משרת את האינטרס הציבורי ומניע תיקון ליקויים.

האינטרס הציבורי בביקורת ממוכנת: יעילות וחיסכון

הדיון בערך שבתהליך ובערך שבתוצר התמקד בביקורת האנושית, אך הוא חסר טעם ללא הצד האחר של המשוואה: האינטרס הציבורי בביקורת הממומשת באמצעות מכונה. סעיף 2(ב) לחוק יסוד מבקר המדינה קובע את תכלית הביקורת ביחס לגופים המבוקרים: "מבקר המדינה יבחן את חוקיות הפעולות, טוהר המידות, הניהול התקין, היעילות והחסכון של הגופים המבוקרים, וכל ענין אחר שיראה בו צורך"⁶⁴. כאשר מדובר בהטמעת AI בעבודת הביקורת, השאלה המעשית הנגזרת היא אחרת: האם השימוש ב-AI מאפשר למבקר לזהות באופן יעיל ומקיף יותר ליקויים בחוקיות, בטוהר המידות, בניהול התקין, ביעילות ובחיסכון של הגופים המבוקרים. במישור זה הבינה המלאכותית מסוגלת להעצים מאוד את כושר הביקורת.

יעילות: הספק ומהירות כמענה על הפיגור הכרוני

בהסתמכות על ביקורת אנושית יחלפו חודשים ואף שנים מרגע התרחשות המעשה ועד פרסום הדוח. סוכן בינה

64 חוק-יסוד: מבקר המדינה, ס"ח התשמ"ח.

65 U.S. Government Accountability Office (GAO), Artificial Intelligence, לעיל ה'ש 51. המסמך מגדיר ארבעה עקרונות אחריות ל-AI ציבורי: ממשל, נתונים, ביצועים וניטור רציף.

66 Daniel Kahneman, Olivier Sibony & Cass R. Sunstein, Noise: A Flaw in Human Judgment, Little, Brown Spark (2021). 66 שיפוט אנושי בתחומי הרפואה, המשפט והביקורת סובל מ"רעש" (שונות לא רצויה בין מחליטים) ברמה שמסכנת את איכות ההחלטות שיתקבלו.

האם הבינה המלאכותית תייתר את עובדי הביקורת?

ההערכה המבוססת על הטכנולוגיה הקיימת היא שעד שנת 2030, וככל שיותר גופים מבוקרים יטמיעו במערכותיהם כלי בינה מלאכותית, יפחת הצורך בבינה אנושית לשם ביצוע עבודת הביקורת. הנתונים שהוצגו בפרקים הקודמים מדברים בעד עצמם: 57% משעות העבודה כבר ניתנות לאוטומציה, הפער בין הפוטנציאל (80% - 94% בתחומים רלוונטיים) לשימוש בפועל (15% - 33%) נובע בעיקר מהקצב האנושי ולא מהיכולת הטכנולוגית, ולפי מדידת METR העדכנית מינואר 2026 (Time Horizon 1.1) יכולותיהם של סוכני AI מוכפלות בקצב של כ-131 יום על פי כלל הנתונים משנת 2023 ואילך, וכ-89 יום בלבד על פי נתוני 2024 ואילך - קצב שמעיד כי ביצוע משימות של יום עבודה מלא יהיה בהישג יד בטווח הקרוב.

אולם חשוב להבחין בין שתי רמות פעילות של כלי הבינה המלאכותית: הרמה הטכנית והרמה החברתית-כלכלית. מבחינה טכנית, יש לבינה המלאכותית יכולת לבצע ביקורת בהיקף ובאיכות ההולמים את זו של הבינה האנושית. ברמה החברתית-כלכלית - מדובר באיגמה. כיצד ארגונים ינהלו מעבר תפעולי כה דרסטי? כיצד יתפרנסו עובדים שתפקידם יתייתר? כיצד שומרים על הידע האנושי המצטבר בתקופת המעבר? אלה שאלות חשובות ודחופות - אך הן לא צריכות לגרום לנו להתכחש לתשובה על השאלה הראשונה.

דפוס ההשפעה שכבר נצפה בשוק העבודה מחזק את ההערכה. אין עדיין פיטורים המוניים, אך יש האטה מובהקת בגיוס עובדים צעירים לתחומים עם חשיפה גבוהה ל-AI.⁶⁷ השוק מגיב בדרך שקטה אך משמעותית: לא בפיטורים המוניים, אלא בצמצום הדרגתי של גיוס העובדים החדשים. מי שעובד בביקורת היום אינו צריך לחשוש - אבל הוא צריך להיערך. על מי שמתכנן קריירה בתחום הביקורת לדעת שהמקצוע שאליו ייכנס ייראה שונה מהותית מהמקצוע המוכר לו כיום.

ראוי לציין שהפורום הכלכלי העולמי כבר זיהה את השליטה ב-AI כמיומנות הצומחת ביותר בשוק העבודה, עם הכפלה פי שבעה בשנתיים.⁶⁸

הגורמים המשפיעים על קביעת האינטרס הציבורי

ההכרעה בין שתי הגישות המתוארות לעיל בנוגע לקביעת האינטרס הציבורי אינה בינארית, והיא תלויה בכמה גורמים משפיעים.

הגורם הראשון נוגע לקצב האימוץ וההטמעה של מערכות AI בגופים המבוקרים עצמם. ככל שגופים מבוקרים עוברים לעבודה מבוססת AI - קבלת החלטות אלגוריתמית, מערכות אוטומטיות, שירותים דיגיטליים וכד' - כך נדרשת גם הביקורת עליהם להתאים את כליה. אם הגוף המבוקר הוא מכונה, אולי גם מי שמבקר אותו צריך להיות מכונה. ביקורת אנושית על מערכת אלגוריתמית שמקבלת אלפי החלטות בשנייה תהיה ככל הנראה חסרת משמעות מעשית.

הגורם השני תלוי באפיון הגוף המבוקר. יש גופים מבוקרים שבהם הגורם האנושי בביקורת עשוי להישמר גם בעידן ה-AI בשל אופיים החשאי, ובפרט מערכת הביטחון, גופי מודיעין או מערכות ליבה מדינתיות אחרות. בגופים אלה שיקול הדעת האנושי, הצורך במידור, הרגישות של המידע והמורכבות הפוליטית עשויים לדרוש נוכחות אנושית לצד כלי AI, ולא במקומם.

בגופי הביקורת על מערכות אלה יהיה ככל הנראה מקום להמשיך ולהעסיק מבקרים אנושיים - אך תפקידם ישתנה מהותית. משינוי תהליך העבודה (מכותב לעורך), דרך פיתוח מיומנויות חדשות (הנדסת פרומפטים, פיקוח על סוכנים, בדיקת מודלים ועוד) ועד להגדרת תפקידים היברידיים כגון "מבקר-טכנולוג" - ראו הרחבה בפרק המתווה האופרטיבי שלהלן.

טענת המאמר אפוא היא כי כאשר הביקורת האנושית הקיימת מכסה שבריר מהמציאות הציבורית, מגיעה באיחור ונסמכת על נתונים חלקיים, אזי הדרישה להישען על בינה אנושית גרידא אינה מובנת מאליה - היא עצמה דורשת הצדקה. הנטל אינו מוטל עוד על מי שמציע להחליף את המבקר האנושי, אלא על מי שטוען שאין להמירו.

מן הדיון הנורמטיבי, נעבור עתה לבחינה כמותית של אופק הזמן: מתי, על בסיס נתוני שוק העבודה והקצב המואץ של יכולות הבינה המלאכותית, עשויה ההחלפה להתממש בפועל.

67. LinkedIn Economic Graph Research Institute, *Labor Market Report* 67, לעיל ה"ש 44. הדוח מתעד צמצום בגיוס עובדים צעירים בתחומים חשופי AI.

68. World Economic Forum, *Future of Jobs Report* (2025), https://reports.weforum.org/docs/WEF_Future_of_Jobs_Report_2025.pdf. 68 הדוח מצביע על "AI fluency" כמיומנות הצומחת ביותר בשוק העבודה.

סיכונים ואתגרים

אין לזלזל בסיכונים שבשילוב בינה מלאכותית בעבודת ביקורת. דווקא מי שמכיר את היכולות המרשימות בתחום צריך להכיר גם את המגבלות הממשיות - ולנהל אותן באחריות.

הזיות - Hallucinations

מודלי שפה גדולים "ממציאים" לעיתים מידע - מצטטים מקורות שאינם קיימים, מציגים נתונים שגויים או מגבשים ממצאים שאין להם בסיס עובדתי. בתחום הביקורת, שבו דיוק עובדתי הוא היסוד לכל ממצא והמלצה, זהו סיכון חמור שאין להקל בו ראש. כל פלט של AI חייב לעבור אימות - אנושי או אוטומטי - ואין לקבלו כעובדה ללא בדיקה צולבת.

הטיות - Bias

מודלי AI נבנים על נתוני אימון שמשקפים הטיות קיימות בחברה ובארגונים. אם מודל אומן על דוחות ביקורת היסטוריים, הוא עלול לשחזר הטיות מוכרות - למשל מיקוד יתר בגופים מסוימים על חשבון אחרים, העדפת סוגי ממצאים מסוימים או התעלמות שיטתית מתחומים שלא קיבלו תשומת לב בעבר. מבקר שמסתמך על AI ללא מודעות להטיות אלה עלול, שלא במתכוון, להנציח בעיות במקום לחשוף אותן.

אבטחת מידע וסודיות

חומרי ביקורת כוללים לעיתים קרובות מידע רגיש ביותר - נתונים כספיים של גופים מבוקרים, מידע אישי של נבדקים, מידע מסווג, נתונים חסויים או כאלה שיש להותירם אנונימיים. השימוש בכלי AI חיצוניים מעלה שאלות מהותיות על הריבונות על מידע: היכן מעובד המידע, מי נחשף אליו, האם הוא משמש לאימון מודלים ועוד. מוסדות ביקורת יצטרכו לפתח מדיניות ברורה ומחייבת ובמקרים רבים להפעיל מודלים פנימיים בסביבה סגורה ומאובטחת.

אתיקה ואחריותיות

מי אחראי לממצא ביקורת שנוצר באמצעות AI? שאלת "חתימת המבקר" אינה פורמלית בלבד. יש לה השלכות משפטיות, מקצועיות ואתיות כבדות משקל. אם סוכן AI מפיק ממצא שגוי שמוביל לנזק לגוף מבוקר או לפגיעה באמון הציבור - מי נושא באחריות? המבקר שהפעיל את הכלי? מנהל יחידת הביקורת? ספק הטכנולוגיה? שאלות אלה דורשות מענה נורמטיבי מוסדר שעדיין אינו קיים ברוב המדינות.

תלות טכנולוגית

מוסד ביקורת שמסתמך על מערכות AI של ספקים חיצוניים עלול למצוא עצמו תלוי בטכנולוגיה, בעדכונים,

בתמחור ובזמינות. שאלת הריבונות הדיגיטלית רלוונטית במיוחד למוסדות ביקורת מדינה, שעצמאותם ואי-תלותם הם ערכי יסוד. בניית יכולות פנימיות, הכשרת צוותים טכנולוגיים ושמירה על אופציות מגוונות יכולים להיות חלק מהפתרון לאתגרים הללו.

ממתודולוגיה ליישום: מפת דרכים

ההכרה ביכולת הטכנולוגית אינה מספיקה - צריך גם דרך פעולה מעשית. להלן הצעה למפת דרכים בשלושה אפיקים, המותאמת הן למוסדות ביקורת מדינה והן ליחידות ביקורת פנימית בארגונים.

בטוח קצר: ההיכרות ושימוש בסיסי

- הכשרת צוותי ביקורת בשימוש בכלי AI קיימים - צ'אטבוטים מתקדמים, כלי סיכום מסמכים, מנועי חיפוש חכמים וכד'. הגדרת מדיניות שימוש ברורה: מה מותר, מה אסור ואילו סוגי מידע אין להעלות לכלים חיצוניים.
 - מינוי והכשרה של "שגריר AI" בכל יחידת ביקורת - גורם שתפקידו לזהות הזדמנויות, להנגיש כלים ולשמש כגשר בין הטכנולוגיה לצוות.
- המטרה המעשית בשלב זה: כל מבקר ישתמש ב-AI בעבודתו השוטפת.

בטוח הבינוני: אינטגרציה והתאמה

- פיתוח או רכישה של כלים ייעודיים לביקורת - מערכות שמחוברות למאגרי המידע הרלוונטיים ומותאמות לתהליכי העבודה הייחודיים של יחידת הביקורת.
- הרצת פיילוט של סוכני AI בנושאי ביקורת ספציפיים, לדוגמה ביקורת ציות, ביקורת רכש או ניתוח תקציבי.
- פיתוח מתודולוגיה לביקורת על מערכות AI בגופים מבוקרים והגדרת מסגרת ממשל AI פנימית.

בטוח הארוך: אוטונומיה ושינוי מבני

- פריסת סוכנים אוטונומיים לביקורת, עם פיקוח אנושי מדורג.
- שילוב מערכתי מלא של AI בתהליכי הביקורת, מתכנן ועד מעקב.
- פיתוח מודלים מותאמים אישית שמאומנים על דוחות ביקורת היסטוריים, חקיקה ונהלים רלוונטיים.

הטמעת בינה מלאכותית בתוך מוסדות הביקורת עצמם

הצעד הראשון הוא הטמעה מלאה של כלי AI בעבודתו השוטפת של מוסד הביקורת: ניתוח מסמכים, סריקת מאגרי נתונים, זיהוי חריגות ואיתור דפוסים. זאת לצד רכש או פיתוח של כלי ביקורת ייעודיים המחוברים למאגרי המידע של הגופים המבוקרים ולשירותי גילוי ראיות אלקטרוני⁷⁰ של מוסד הביקורת. על מוסד הביקורת להיות חלוץ בהטמעתה של הטכנולוגיה שהוא עתיד לבקר.

ביקורת על מערכות בינה מלאכותית בגופים המבוקרים

בד בבד, על מוסד הביקורת לפתח מתודולוגיה ייחודית לביקורת על מערכות AI בגופים מבוקרים. בעידן שבו משרדי ממשלה, רשויות רגולטוריות וחברות ממשלתיות משלבים אלגוריתמים בתהליכי קבלת החלטות - החל מקבלה למוסדות חינוך וכלה בהקצאת טיפול רפואי דחוף - על המבקר לבחון את טיב הנתונים, את הוגנות התוצאות, את קיומה של בקרה אנושית ואת עמידתם של המודלים בהוראות הדין ובכללי האתיקה. מסגרת ה-GAO למבקר בינה מלאכותית מגדירה את העקרונות הללו במפורש: ממשל, נתונים, ביצועים וניטור רציף.

ממשקים טכנולוגיים בין-ארגוניים

כדי שכלי הבינה המלאכותית המוטמעים במערכות המבקר יוכלו לקיים דיאלוג ישיר עם כלי הבינה המלאכותית שבמערכות הגוף המבוקר, יש לקיים תנאים אחדים: קיום תקני ממשק פתוחים (Open APIs), הסדרה חוזית של הרשאות גישה, יצירת שכבת אימות זהות אלגוריתמית והגדרת פרוטוקולי אבטחה הדדיים. כדי להבטיח עצמאות מקצועית ויכולת ביקורת רציפה שאינה תלויה ברצון הטוב של המבוקר, על מוסד הביקורת לקבוע את הדרישות, ולא רק לצרוך את המידע שהגוף המבוקר מספק לו⁷¹.

פיילוטים מבוקרים

המעבר ממצב של עבודה אנושית לעבודה הנתמכת בבינה מלאכותית מחייב הפעלת פיילוטים מבוקרים. אלה יערכו בצורה מובנית: בחירת נושא ביקורת שמאפיין תכנים דיגיטליים מובנים (כגון רכש, תשלומים או ציוד); הרצת

- התאמת המבנה הארגוני: פחות מבקרים "קלאסיים" ויותר מפקחים על מערכות AI, מנהלי נתונים ומומחי מתודולוגיה. גם תפקידי הפיקוח הללו נכונים לשלב המעבר ולא למתכונת הקבע של המקצוע⁶⁹.

יש לשקול גורם נוסף: בתנאים הישראליים, שבהם מוסדות ציבוריים פועלים עם משאבים מוגבלים ובהיקף קטן יחסית למדינות גדולות, השימוש ב-AI עשוי לא רק לשפר את איכות הביקורת אלא גם להניב חיסכון משמעותי בעלויות. התפתחות המבנה הארגוני הזה אינה היפותזה ספרותית. על פי דוח שוק העבודה של לינקדאין מינואר 2026 ("Building a Future of Work That Works"), בשנתיים האחרונות נוצרו ברחבי העולם כ-1.3 מיליון משרות חדשות הקשורות ל-AI, בהן תפקידי גשר חדשים בתחומים כמו Forward-Deployed Engineer, שצמח פי 42 מאז 2023, AI Engineer, שצמח פי 13. במקביל, בדוח Work Change Report של לינקדאין מינואר 2025 נמצא כי מספר החברות בארצות הברית המעסיקות עובד בתפקיד ראש תחום בינה מלאכותית גדל פי שלושה בחמש השנים האחרונות וב-50% נוספים בשנתיים האחרונות בלבד. התפקידים הללו הם המערכת החדשה של בקרה על מערכות אוטונומיות.

מתווה אופרטיבי לשילוב בינה מלאכותית בעבודת הביקורת

הפרק הקודם ארגן את הצעדים לפי לוח הזמנים שבו ייושמו. פרק זה מארגן אותם לפי הצירים התוכניים בעבודת הביקורת. פרק זה משלים את האמור עד כה במתווה אופרטיבי מפורט שנשען על מערכת שיקולים נורמטיביים ואופרטיביים. המתווה נבנה סביב תשעה צירי פעולה המכסים שני מעגלים משלימים: מעגל כנימי - הטמעת בינה מלאכותית בתוך מוסד הביקורת עצמו, ומעגל חיצוני - פיתוח יכולת הביקורת על בינה מלאכותית המשמשת את הגופים המבוקרים. נטען כי שני המעגלים ניזונים זה מזה: מוסד ביקורת שאינו משתמש בכלים מתקדמים אינו יכול לבקר ביעילות גוף שמשתמש בהם.

McKinsey Global Institute, *A New Future of Work* 69, לעיל ה"ש 33. הדוח מצביע על 12 מיליון שינויי תעסוקה נדרשים באירופה עד 2030.

70 eDiscovery הוא תהליך מובנה לאיתור, סינון, ניתוח ושימור מסמכים דיגיטליים רלוונטיים לחקירה או להליך משפטי. הוא מבוצע באמצעות פלטפורמות ייעודיות כגון Everlaw ו-Relativity וכפוף לתקני EDRM (Electronic Discovery Reference Model).

Organisation for Economic Co-operation and Development, OECD AI Principles (2019 updated 2024), <https://www.oecd.org/en/topics/71-ai-principles.html>. OECD מגדירה חמישה עקרונות לאימוץ AI אחראי במגזר הציבורי.

ההכשרה המקצועית. המודל המוצע כולל צוותים היברידיים שבהם למבקר אנושי אחד יש עשרות סוכני AI הפועלים תחתיו במקביל⁷³.

הדרך הפחות טראומטית למעבר היא הימנעות מצמצום כוח האדם הקיים והסתמכות על תחלופה טבעית. בהערכה שמרנית, כ-3% - 5% מכוח האדם של מוסד ביקורת ממוצע עוזבים בכל שנה בשל פרישה, מעבר לסקטור הפרטי או קידום בין-ארגוני, בדומה לשיעורי תחלופה טיפוסיים בשירות הציבורי הישראלי ובדמוקרטיות מפותחות⁷⁴. על פני חמש עד עשר שנים, עצירת גיוס חלקית, למשל גיוס עובד אחד במקום כל שני עובדים שעוזבים, עשויה לצמצם את הצוות האנושי ב-15% - 30% בלי לפגוע בעובדים הקיימים. מהלך זה מחייב תכנון אסטרטגי ארוך טווח.

סיכום

מאמר זה פתח בשאלה אם העין הפקוחה חייבת להיות עין אנושית. התשובה שעולה מן הנתונים, מתרחישי השימוש בעולם ומן הניתוח הנורמטיבי ברורה מכפי שרבים היו רוצים להודות - לא, היא אינה חייבת להיות אנושית, וייתכן שהגיעה העת לומר בכנות שיש צורך בהתאמה ובהכנה משמעותית ליום זה.

הטכנולוגיה הקיימת כבר היום - לא בעוד עשור ולא בעוד חמש שנים - מסוגלת לבצע את רובו המכריע של תהליך הביקורת: איסוף הנתונים, זיהוי החרגות, הצלבת המקורות, גיבוש הממצאים וניסוח הדוחות. הפער בין הפוטנציאל, הנאמד ב-80 עד 94 אחוזים בתחומים הרלוונטיים, ובין השימוש בפועל, שאינו עולה על 15 עד 33 אחוזים, אינו פער של יכולת אלא של מוכנות מוסדית. יכולות הסוכנים הדיגיטליים מוכסלות מדי חודשים ספורים, כך שהפער הזה יצטמצם בין אם מוסדות הביקורת יערכו לכך ובין אם לאו. היעדר היערכות אינו עיכוב של השינוי, הוא רק שינוי הזהות של מי שיוביל אותו.

אולם, חשיבותו של המאמר איננה בתיאור המצב הטכנולוגי - היא בהיפוך שאלת המוצא. כל עוד הדיון הציבורי נסוב סביב השאלה "כיצד לשלב את הבינה המלאכותית בעבודת

פיילוט במקביל לעבודה האנושית הרגילה כדי להשוות ממצאים; תיעוד מדויק של שיעור השגיאות, שיעור הזיהויים כוזבים ושיעורי ההחמצה; והכפפת כל שלב לביקורת איכות מקצועית בהתאם לתקני ISSAI ולמדריך לביקורת המדינה.

תחומי ליבה שיישמרו כולם או בחלקם בידי אדם

חרף הפוטנציאל הרחב של הבינה המלאכותית, קיימים תחומים שבהם שיקול הדעת האנושי הוא תנאי לקיום תקין של הביקורת. מדובר בביטחון, בענייני חוץ ומודיעין ובכל תחום שהכרעה בו מחייבת איזון בין זכויות יסוד (כגון פרטיות, הליך הגון או חופש ביטוי) לבין אינטרסים ציבוריים. בתחומים אלה ייקבע העיקרון "human-in-command", קרי: המבקר האנושי מגבש את המסקנה הסופית, אך נעזר ב-AI לשם עיבוד חומרי הגלם. תקנת ה-EU AI Act מציעה מסגרת תקדימית שניתן להתאים לצורך זה.

רגולציה על בינה מלאכותית הפועלת ללא בקרה אנושית

כאשר תחומי פעולה שלמים של הביקורת יעברו לידי מערכות אוטונומיות, תידרש רגולציה ייחודית להפעלתן: חובת תיעוד שלבי ההסקה, חובת הסבריות, חובת עדכון הגוף המבוקר בקיומה של ביקורת אוטונומית וחובת ההותרה של זכות ערעור אנושית במקרה של ממצא שלילי. מסגרת ה-NIST AI Risk Management Framework ותקנת ה-EU AI Act מספקות בסיס לעקרונות אלה. חקיקה ייעודית בישראל, שתבוא בהשראת חוק הביקורת הפנימית וחוק-יסוד: מבקר המדינה, תידרש לתקן את הדין כדי להתאימו לעידן של ביקורת נטולת יד אדם⁷².

אסטרטגיית משאבי אנוש חדשה למוסדות ביקורת

המעבר לעידן של ביקורת נתמכת AI מחייב עדכון מקיף של אסטרטגיית משאבי האנוש במוסדות ביקורת. נדרשות מיומנויות חדשות - הנדסת פרומפטים, פיקוח על סוכני AI, ניתוח נתוני עתק ובדיקת מודלים. יש לבנות מסלולי הסבה מקצועית לעובדי ביקורת, להגדיר תקנים ייעודיים של "מפקח על סוכני AI" או "מבקר מומחה לבינה מלאכותית" ולשנות את תוכניות הלימוד באוניברסיטאות ובמוסדות

72 National Institute of Standards and Technology (NIST), *AI Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1 (2023), <https://www.nist.gov/itl/ai-risk-management-framework>. מסגרת מתודולוגית לניהול סיכונים בינה מלאכותית המשמשת בסיס לרגולציית AI בארצות הברית ובגופי ממשל.

73 Claude, Anthropic, *Labor Market Impacts of AI* 73, לעיל ה"ש 6. המחקר ניתח מיליוני שיחות אנונימיות עם המודל Claude ומדד שימוש ב-AI בקטגוריות תעסוקה שונות של משרד העבודה האמריקני.

74 OECD, *Government at a Glance 2025* (2025), https://www.oecd.org/en/publications/government-at-a-glance-2025_0efdbcd-en.html. פרק 5 על תעסוקה במגזר הציבורי. הערה: שיעור התחלופה השנתי בשירות הציבורי המרכזי במדינות OECD נע בדרך כלל בטווח 3% - 7%, בהתאם להבדלים מבניים, מבחינת פנסיה והיצע שוק העבודה.



ואולי אף יתעצם דווקא משום שיופנו אליו כחות עבודות טכניות שמסיטות אותו מהמקום שבו הוא באמת יוצר ערך. הגדרת המבקר האנושי של העתיד דרך השכבות הללו בלבד איננה פשרה על המקצוע; היא מימוש המלא הראשון של הערך שהמקצוע התיימר לייצג מלכתחילה.

עבור ישראל - שמבקר המדינה שלה ניחן בסמכויות חוקתיות רחבות באופן יוצא דופן, ושהאקו-סיסטם הטכנולוגי שלה נמנה עם המתקדמים בעולם - השאלה אינה אם להיערך אלא כמה מהר לעשות זאת. המוסדות שיקדימו לא רק ישמרו על רלוונטיות - הם יגדירו מחדש את הסטנדרט שלפיו תימדד ביקורת המדינה בעשורים הבאים. המוסדות שימתינו יגלו שהציבור שהם אמורים לשרת כבר אינו מוכן להסתפק בדוח שמתפרסם חמש שנים אחרי המחדל, בשעה שכל אזרח מבין שהטכנולוגיה יכולה לזהות אותו מראש.

כותרת המאמר - המבקר האחרון? - הציגה שאלה. מי שממהר לענות עליה בביטחון - "לא, זה לא יקרה" או "התפקיד רק ישתנה" - אינו מגונן על מקצוע הביקורת, אלא על דמות מסוימת של מבקר. הביקורת תשרוד; היא אף עשויה לשגשג כפי שלא שגשה מעולם. אך תפקיד המבקר האנושי, במתכונת שאנו מכירים כיום, עשוי בהחלט להסתיים. השאלה שנותרת - המרתקת והמטרידה כאחת - איננה כיצד למנוע זאת, אלא האם ראוי לנסות למנוע.

המבקר, הוא מאשר בשתיקה את ההנחה שהעבודה הזו, בליבתה, חייבת להישאר אנושית. המאמר טוען שההנחה הזו טעונה הצדקה ואינה מובנת מאליה. המצב הקיים - שבו חלקים מן המגזר הציבורי אינם מבוקרים, הדגימה חלקית והדוחות מגיעים לעיתים שנים לאחר מעשה - אינו תוצר של בחירה מושכלת כי אם של מגבלת המשאבים האנושיים. כאשר המגבלה הזו מוסרת, ההיצמדות לסטטוס קוו חדלה להיות ברירת מחדל והופכת להחלטה שיש להצדיקה במפורש.

הציבור שהביקורת משרת אינו חפץ בעבודתו של המבקר; הוא חפץ בתוצריה - באמונה שהכסף הציבורי מנוהל כראוי, שההחלטות מתקבלות בשקיפות ושהמחדלים מזוהים ומתוקנים בזמן שבו עדיין אפשר לתקנם. אם ניתן לספק תוצרים אלה בכיסוי מלא במקום בדגימה, ברציפות במקום ברטרופסקטיבה ובתוך ימים במקום שנים - הרי שהזכות הציבורית לקבלם גוברת על הזכות המקצועית לבצעם באופן שבו התרגלנו שיבוצעו. ההיצמדות למבקר האנושי אינה נאמנות לערך הביקורת; במקרים רבים היא הפכה לחסם בפניו.

אין בכך קביעה שהגורם האנושי הופך מיותר לחלוטין. יש היבטים של שיקול דעת - הכרעות ערכיות, התמודדות עם לחצים פוליטיים ופרשנות של ממצאים בעלי השלכות ציבוריות כבדות - שבהן תפקידו של האדם יישאר מרכזי,