

בינה מלאכותית בתחום אבטחת הסייבר

נאדין וירקוויס והדס קליין

תחום ביטחון הסייבר יכול לצאת נשכר משמעותית משימושים בבינה מלאכותית (AI). טכניקות בינה מלאכותית יכולות לשפר את ביצועי האבטחה הכוללים במקומות שבהם מערכות אבטחה קונבנציונליות עלולות להיות איטיות ובלתי מספקות, ובכך לספק הגנה טובה יותר מפני המספר העולה של איומי סייבר מתוחכמים. לצד ההזדמנויות הרבות שניתן למצוא בשימוש בבינה מלאכותית לאבטחת סייבר, עולים גם סיכונים וחששות. כדי לחזק את בשלותה ומוכנותה של אבטחת הסייבר, יש צורך במבט כולל על סביבת הסייבר של הארגון, מבט שישלב בינה מלאכותית עם תובנות אנושיות, שכן רק בני אדם או רק בינה מלאכותית אינם מספיקים להשגת הצלחה מלאה בתחום זה. שימוש אחראי בטכניקות של בינה מלאכותית הינו חיוני כדי להמשיך ולמתן את הסיכונים והחששות הנלווים לטכנולוגיה זו.

מילות מפתח: אבטחת סייבר, בינה מלאכותית (AI), מודיעין ביטחון, גישת אבטחה משולבת (ISA), שרשרת התקיפה בסייבר

מבוא

מאז שפורסמה מתקפת מניעת השירות (denial-of-service) הראשונה ב־1988¹, הכמות, התחכום וההשפעה של מתקפות הסייבר עלו במידה ניכרת. ככל שמתקפות הסייבר הפכו לממוקדות וחזקות יותר, כך השתכללו אמצעי ההגנה. כלי האבטחה הראשון אמנם הוגבל לאיתור חתימות של וירוסים ולמניעת הפעלתם, אך כיום אנו עדים לפתרונות שתוכננו להעניק הגנה מקיפה נגד טווח רחב של סוגי מתקפות ולמגוון מערכות המהוות יעד למתקפה. יחד עם זאת, ההגנה על נכסי מידע בעולם הווירטואלי עדיין מהווה אתגר הולך וגדל.

נאדין וירקוויס, דוקטורנטית במכון האוניברסיטאי למדעים וטכנולוגיה של אוקינאוה ומתמחה לשעבר בתכנית ביטחון סייבר של המכון למחקרי ביטחון לאומי. הדס קליין, מנהלת תכנית ביטחון סייבר במכון למחקרי ביטחון לאומי.

כדי ליישם הגנה רציפה וחסונה, על מערכות האבטחה לדעת להתאים עצמן ללא הרף לסביבה, לאיומים ולגורמים המעורבים בשדה הסייבר, המשתנים בהתמדה. למעשה, המציאות בתחום אבטחת הסייבר שונה במקצת מהרצוי: גישות האבטחה מותאמות בדרך כלל לאיומים ידועים, וקשיחות והיעדר גמישות גורמים לכך שמערכות אבטחה מתקשות להסתגל באופן אוטומטי לשינויים החלים בסביבתן. גם כשקיימת אינטראקציה אנושית, תהליכי ההסתגלות צפויים להיות איטיים ובלתי מספקים.²

מאפייני הטכניקות של בינה מלאכותית, כמו גמישות, יכולת הסתגלות והתנהגות מערכת עמידה, יכולים לסייע בהתגברות על חסרונות שונים הקיימים בכלי אבטחת הסייבר של ימינו.³ עם זאת, ולמרות שבינה מלאכותית כבר שיפרה במידה ניכרת את אבטחת הסייבר,⁴ היא מעוררת חששות, ויש מי שרואים בה אפילו סיכון קיומי לאנושות.⁵ בהתאם לכך, מדענים ומומחי משפט הביעו את חששם מהתפקיד ההולך וגדל של ישויות בינה מלאכותית אוטונומיות במרחב הסייבר, והעלו ספק אם השימוש בהן מוצדק מבחינה אתית.⁶

מטרת מאמר זה היא להתמקד בחסרונות של אמצעי האבטחה המסורתיים בתחום הסייבר ולהדגיש את ההתקדמות שנעשתה עד היום ביישום טכניקות של בינה מלאכותית באבטחת באותם אזורים. בנוסף, המאמר מסכם את הסיכונים והחששות הנלווים להתפתחויות בתחום הבינה המלאכותית, וזאת באמצעות בחינת הסטטוס הנוכחי, התייחסות לבעיות הקיימות והצבעה על כיווני התפתחות לעתיד.

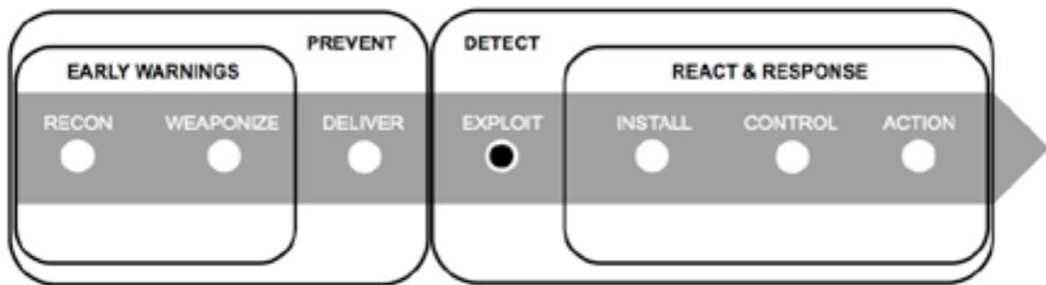
אתגרים באבטחת הסייבר

למרות הגברת המודעות לאיומי הסייבר והשקעת סכומים גדולים של כסף ומאמצים כבירים כדי להילחם בפשעי סייבר, היכולת של ארגונים לספק הגנה נאותה לנכסים הווירטואליים שלהם הינה עדיין סוגיה פתוחה.⁷ הגורמים המעורבים במרחב הסייבר נעים מיחידים וארגונים פרטיים, דרך גורמים לא מדינתיים ועד ארגונים ממשלתיים. אלה מבקשים להגן על נכסי הסייבר שלהם, לתקוף נכסי סייבר של אחרים, או לעשות את שני הדברים גם יחד. המקורות לאיומי הסייבר הם מגוונים ונובעים מפעולות זדוניות פוטנציאליות העשויות לבוא מסיבות פיננסיות, פוליטיות או צבאיות, ואפילו מסוגים מסוימים של תחביבים.⁸

למרות ההטרונגויות והדינמיות של מרחב הסייבר, ניתן לאתר קווי דמיון מסוימים של מתקפות סייבר ושל אמצעי-נגד להן, כדי שאפשר יהיה לבנות מסגרת אבטחה מקיפה והוליסטית לתחום זה. מרבית מתקפות הסייבר מתאפיינות בשלבי מתקפה, אותם ניתן לתאר כ"שרשרת התקיפה בסייבר" (cyber kill chain).⁹ מסגרת האבטחה מניחה שכל רצף של מתקפות מתחיל בשלב איסוף המידע, שבו

התוקף מנסה לאתר פרצות ופגיעויות ביעד המותקף (מערכת מטרה). השלב הבא במתקפה הוא התחמשות (weaponizing), שבו נקודות התורפה שהתגלו משמשות לפיתוח קוד זדוני ממוקד מטרה. לאחר מכן מגיע שלב החדירה (delivery), שבו התוכנה הזדונית מוחדרת למטרה הפוטנציאלית. לאחר שהתוכנה הזדונית הוחדרה בהצלחה, מגיע שלב הביצוע, שבו אותה תוכנה זדונית מפעילה את ההתקנה של הקוד הפולש. כתוצאה מכך, ניתן להתקין במערכת המארכת שנחשפה ערוץ פיקוד ושליטה, המאפשר לתוקף ליזום דרכו פעולות זדוניות. אמצעי-הנגד ייגזרו מהמקום שבו התגלתה הפעולה הזדונית בשרשרת התקיפה בסייבר.

גישת האבטחה המשולבת¹⁰ (ISA) מציעה רעיונות מרכזיים ומסגרת כוללת לטיפול בהגנת סייבר. המטרה המרכזית של גישה זאת היא להפיק התרעות או אזעקות, רצוי לפני שהמתקפה משוגרת (לפני שלב הביצוע). ההתרעה אמורה לחולל הודעת אזהרה שתדע לתרגם את הנתונים החדשים שנאספו על האיום למשימות יישומיות. ההודעה גם אמורה לתמוך בבחירת אמצעי-הנגד, או לכלול בתוכה אמצעי-נגד ייעודיים שימנעו מהארגון ליפול קורבן למתקפה. אם לא ניתן למנוע מראש את המתקפה, יש לנסות לאתר את היקפה ולאחר מכן להגיב בהתאם. אמצעי-הנגד צריכים לכלול פעולות לעצירת הפורץ או אף לביצוע מתקפת-נגד עליו, וזאת בנוסף להגדרת נוהלי שיקום שיאפשרו שחזור מהיר של המערכת והשבחה למצבה הקודם.



תרשים 1. השלבים בשרשרת תקיפה בסייבר מול אמצעי-הנגד במסגרת גישת האבטחה המשולבת

תרשים 1 מתאר את הקשר ההדדי המתקיים בין מתקפת סייבר על פי שרשרת התקיפה בסייבר ובין אמצעי-הנגד המופעלים על פי גישת האבטחה המשולבת. התרשים מפרט את שרשרת התקיפה בסייבר, כפי שהיא מומחשת על ידי החץ האפור, ואת גישת האבטחה המשולבת, המומחשת על ידי המסגרות.

שרשרת התקיפה בסייבר כוללת את שבעת השלבים של מתקפת הסייבר כמפורט בתרשים, ואילו גישת האבטחה המשולבת מורכבת מארבעה שלבי פעולת-נגד. כדי לזהות ולחסום מתקפות מוקדם ככל האפשר, יש לקחת בחשבון

במסגרת גישת האבטחה המשולבת הכוללת את כל שלבי מתקפת הסייבר בשרשרת התקיפה.¹¹ כאמור, הדגש מושם על מניעת מתקפה ואיתור פעילויות זדוניות במהלך שלושת השלבים הראשונים של הפריצה, המוצגים בתרשים – איסוף המידע (recon), ההתחמשות (weaponize) והחדירה (deliver). לאחר המתקפה המתוארת בתרשים כשלב הביצוע (exploit), יש להפעיל את אמצעי גישת האבטחה המשולבת כדי לשבש את פעילות התוכנה הזדונית. המדובר בזיהוי (detection), תגובה (reaction) ומענה (response).

טיבו המורכב והדינמי של מרחב הסייבר מוביל למגוון אסטרטגיות ואתגרים טכנולוגיים שמעכבים ומסבכים את יכולת הארגון להגן על עצמו בצורה מספקת בסביבה הווירטואלית. אתגרים אלה כוללים את הצורך בהשגת נתונים, סוגיות הקשורות בטכנולוגיה, וכן פגמים ופערים ברגולציה ובניהול תהליכים.

אתגרים באיסוף מודיעין סייבר

העובדה שהתוקפים מותירים אחריהם עקבות לאחר הניסיון לתקוף מערכת מסוימת היא המפתח המאפשר להכיר אותם טוב יותר. על רקע זה, גישת האבטחה המשולבת והכוללת מחייבת איסוף וניתוח של מידע רב כדי להפיק ממנו מודיעין סייבר.¹² יחד עם זאת, ישנם אתגרים הניצבים בפני השגת נתונים רלוונטיים, כמו גם עיבודם, ניתוחם ושימוש בהם. לפיכך, המאמצים הקשורים למניעה, לזיהוי ולתגובה יעילים לפריצות זדוניות נעזרים כיום תדיר בכלי אבטחה שמבקשים להכניס אוטומציה בתהליכי האבטחה.

האתגרים המרכזיים הניצבים בפני השגת נתונים רלוונטיים הם:¹³

- **כמות הנתונים:** כמות הנתונים עלתה באופן מעריכי (exponential) מאז שהתקנים אלקטרוניים ושימוש בהם הפכו לחלק בלתי נפרד מהעבודה ומחיי היומיום. כדי ליישם גישת אבטחה משולבת, יש לתת את הדעת לכל הנתונים מכל המערכות ובכל הארגונים.
- **הטרוגניות של נתונים ומקורותיהם:** השונות בנתונים ובמקורותיהם מקשה על זיהוי ואיסופם. יתרה מכך, הנתונים והמקורות חוצים גבולות ארגוניים, גאוגרפיים ותרבותיים. גם אם זוהתה ההטרוגניות הרלוונטית של סביבת הסייבר, הטופולוגיה וההתנהגות של מערכות ורשתות עשויות להשתנות, ולפיכך לחייב התאמה מתמדת.
- **מהירות נתונים גבוהה:** הקצב הגבוה שבו מופקים ומעובדים נתונים מציב אתגרים בתחומי האחסון והעיבוד שלהם. נתונים אלה חיוניים להמשך הניתוח. כאשר מדובר בעיבוד, ניתוח ושימוש בנתונים שהושגו, מערכות מניעה ואיתור של פריצות (IDPS) הן כלי רב ערך באבטחת סייבר – אחד מרבים בארסנל אבטחת הסייבר כיום.¹⁴ מערכת IDPS יכולה להיות חומרה או תוכנה, שהוגדרה כמיועדת

להגן על מערכות בודדות או על רשתות שלמות. מערכת זו מבוססת על שני עקרונות מרכזיים: זיהוי שימוש חריג (misuse detection), שמזהה פעילות זדונית באמצעות הגדרת דפוסי התנהגות חריגה של מערכת ו/או רשת; איתור אנומליות (anomaly detection), שמתבסס על הגדרת דפוסים של התנהגות נורמלית במערכת ו/או רשת. מומחי אבטחה מגדירים שני דפוסים אלה תוך הסתמכות בעיקר על ניסיונם וכן על ידע קיים בדבר איומי סייבר.¹⁵

לבטחון סייבר יש אופי מורכב ודינמי במיוחד. איומים חדשים מופיעים חדשות לבקרים, ומתקפות נתפרות באופן ספציפי, כך שיערימו על תרחישי הגנה מוכרים. אמנם, המאפיינים של מערכת IDPS הם ביצועים מיטביים, הגנה מרבית ומספר מזערי של שגיאות,¹⁶ אך מערכות האבטחה המסורתיות אינן מסוגלות עוד למלא דרישות אלו במלואן. להלן נקודות התורפה הטכנולוגיות הקריטיות ביותר של מערכת IDPS:¹⁷

- **שיעור זיהוי נמוך:** כל אי־דיוק בהגדרת דפוסים של נורמליות או חריגות בהתנהגות הרשת ו/או המערכת עלולים להשפיע על שיעור הזיהוי של מערכת IDPS. סביבת הרשת המשתנה תדיר הופכת משימה זו למאתגרת אף יותר. שגיאות בהגדרת דפוסים חריגים יכולות להוביל לשיעורים גבוהים של תוצאה שלילית שגויה (false negative). במקרה כזה, פעילות הרשת לא תזוהה מראש כניסיון תקיפה, משום שההנחה תהיה שאין מדובר בהתנהגות רשת זדונית. מנגד, הגדרה שגויה של דפוסים נורמליים תוכל לגרום לשיעורים גבוהים של תוצאה חיובית שגויה (false positive), ולהביא לכך שפעילות רשת שאינה זדונית תסווג ככזו בטעות.
- **קצב העברה איטי:** מערכות IDPS יכולות לסבול ממגבלות בעיבוד וניתוח ג'יגה־ביתים של נתונים לכל שנייה. מנגנונים שמתייחסים לסוגיה זו מתבססים בעיקר על התפלגות של עיבוד הנתונים, וכך הם יכולים להשפיע עוד יותר על פעולת המערכת, על תחזוקתה ועל העלויות הנלוות.
- **היעדר מדרגיות וחוסן:** סביבות סייבר הן דינמיות. התשתיות והתנועה ברשת משתנות ומתרחבות ללא הרף, וכמויות עצומות של נתונים הטרורגניים זקוקים לעיבוד וניתוח. דינמיקות כאלו תורמות לבעיות בתחום הביצוע ואף לאובדן יעילות, שכן מערכות IDPS לא תמיד מסוגלות לספק ולממש את הפונקציונליות שלהן כאשר הן מתמודדות עם דינמיקות מסוג זה.
- **היעדר אוטומציה:** מערכות IDPS אינן מסוגלות לפי שעה להסתגל אוטומטית לשינויים בסביבה. התוצאה יכולה להיות צורך בניתוח נתוני יומן על בסיס פרטני, התאמה ידנית של מערכות לשינויים בסביבת רשת, או מומחים שיקבעו מה התגובה המתאימה לכל הודעת אזהרה בנפרד. היעדר אוטומציה מביא לצורך מתמיד בפיקוח אנושי, גורם לעיכובים וכן מוסיף תקורות בעלויות ובמשאבים.

אתגרי הטכנולוגיה גורמים לארגונים להיתקל בשלב מסוים בליקויים באבטחה. ייתכן שאותם ארגונים ישתמשו אז בכמה מערכות אבטחה, או ירכשו מודיעין אבטחה בצורה של ייעוץ על ידי ספקים חיצוניים.¹⁸

אתגרים נוספים

מלבד הצורך בהשגת נתונים כוללת ובשימוש בטכנולוגיות אבטחה מוצקות להגנה קבועה על כל טווח המידע, יש לשקול גם את הצורך בתהליכים תומכים. הקמה ותחזוקה של תהליכים תומכים הן חשובות בדיוק כמו השגת הנתונים והשימוש בטכנולוגיות אבטחה הולמות. תהליכים פנים-ארגוניים ובין-ארגוניים יכולים לסייע לשפר ולשמר את גישת האבטחה המשולבת בארגונים, וזאת בנוסף לשיפור רמת המוכנות של אבטחת הסייבר שלהם.¹⁹ זאת ועוד, יצירת מה שמכונה "סביבת הסייבר העסקית"²⁰ מעודדת שותפויות בין גורמים מגוונים לכל אורך סביבת הסייבר, במטרה לטפל ולשתף באיומים ביטחוניים, בניסיון שנרכש ובמשאבים. ארגונים המגיעים מענפים שונים צפויים להציג דרישות שונות לאבטחת סייבר. הבדלים אלה יכולים להיות תוצאה של דרישות אבטחה הטרורגניות או של תגובות משתנות לנוכח מתקפות סייבר דומות.²¹ כאשר ארגונים צריכים להגן על תשתיות קריטיות, כמו מתקני אספקת מים או מתקני חשמל גרעיניים, הם יתמקדו באבטחה מוגברת במקום בהיבטים הכלכליים. ארגונים פרטיים נוטים להתמקד בהפסדים הכספיים ולא לייחס חשיבות רבה מדי לסכנה לביטחון הציבור.²² אלה רק מקצת האתגרים שמטרידים ארגונים בבואם לגבש את אסטרטגיית האבטחה שלהם. לנוכח התפקיד החשוב שממלאות מערכות בטחון הסייבר בהקשר זה, נתמקד להלן באמצעים הטכנולוגיים.

שכניקות מודיעין לסיוע לאמצעי אבטחה

מכונות חכמות הן פתרון המבטיח שיפור אמצעי האבטחה העכשוויים במסגרת ההתמודדות עם סוגיות של איסוף מודיעין לאבטחת סייבר. מכונות חכמות יכולות לבצע כמה מהיכולות הקוגניטיביות האנושיות (כמו ללמוד או לנמק) וכן לכלול פונקציות חישה (יכולת לשמוע או לראות). מכונות כאלו מפגינות את מה שאנו מכנים אינטליגנציה²³ או בינה. בינה מלאכותית כזו מאפשרת למכונות להתנהג בתבונה ולחקות בינה אנושית – אם כי במידה מוגבלת.

התפתחות מערכות הבינה, בין אם בחומרה ובין אם בתוכנה, אפשרה לפתח שיטות לפתרון בעיות מורכבות, כמו סיווג או הקבצה (clustering), שלא ניתן היה לפתור מבלי לעשות שימוש בסוג כלשהו של אינטליגנציה.²⁴ לעומת מערכות מחשב מסורתיות, המבוססות על אלגוריתמים²⁵ קבועים ומחייבות תבניות נתונים ידועות מראש לקבלת החלטות, מומחים בתחום הבינה המלאכותית במדעי המחשב

פיתחו טכניקות שונות וגמישות, כגון רשתות עצבים מלאכותיות או רשתות עצבים עמוקות, שמאפשרות למחשבים ללמוד²⁶ ולהתאים את עצמם אוטומטית לדינמיקות של הסביבה. במרחב הסייבר, אימוץ אוטומציה יכול לכלול תבניות נתונים הטרורגיים, מקורות נתונים משתנים, או "רעש"²⁷ בפעילויות הסייבר. אבטחת סייבר היא לכאורה התחום שיכול להשיג את הרווח הגדול ביותר מהכנסה לשימוש של הבינה המלאכותית הממוחשבת. מערכות בינה מלאכותית אוטונומיות אמורות להתגבר על האתגרים הניצבים בפני מערכות אבטחה קונבנציונליות.²⁸ התוצאה הצפויה משימוש בבינה מלאכותית היא מיתון בעיות הקשורות להשגת נתונים (כמות, הטרורגיות ומהירות הנתונים), וכן מיתון בעיות בכלים הקשורים לכך (שיעור זיהוי נמוך, קצב העברה איטי, היעדר מדרגיות וחוסן והיעדר אוטומציה). כך ניתן יהיה לשפר את היעילות והאפקטיביות של אבטחת הסייבר ושל הכלים הנלווים אליה.

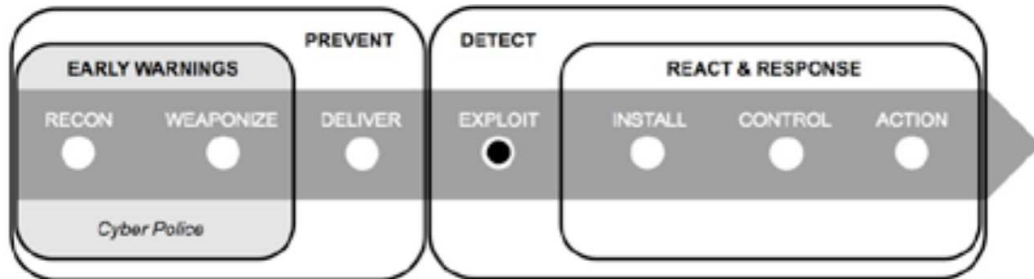
תחום הבינה המלאכותית פיתח ועדיין מפתח מספר רב של טכניקות שיתנו מענה להתנהגות מערכות חכמות. טכניקות רבות כאלו כבר הותקנו במערכות לאבטחת סייבר. אותן מערכות יכולות לעבד ולנתח מידע בהיקפים עצומים בתוך פרק זמן סביר, ובמקרה של ניסיון למתקפה הן יכולות לנתח את המידע ולבחור פעולות נגד ייעודיות. ישנם תרחישים אפשריים שונים ליישום טכניקות של בינה מלאכותית לצורכי אבטחה, וזאת בארבע הקטגוריות של גישת האבטחה המשולבת. באמצעות אותם תרחישים ניתן להמחיש את מגוון האפשרויות שמציעים ענפי הבינה המלאכותית השונים.

"סוכני משטרה" חכמים ואינטראקטיביים לניטור רשתות שלמות

הפרדיגמה של "סוכנים חכמים" (intelligent agents) היא ענף של בינה מלאכותית שצמח מתוך הרעיון לפיו ידע בכלל, וידע לפתרון בעיות בפרט, אמור להיות משותף לישויות שונות. "סוכן יחיד" הוא ישות קוגניטיבית אוטונומית²⁹ עם מערכת קבלת החלטות משלה ויעד אישי. כדי להשיג את היעד שלו, הסוכן היחיד פועל בצורה פרו־אקטיבית במסגרת סביבתו ומול סוכנים אחרים. לסוכנים היחידים יש גם התנהגות ריאקטיבית: הם מבינים שינויים בסביבתם ומגיבים עליהם, וכן יוצרים אינטראקציה עם הסביבה ועם סוכנים מבוזרים אחרים. בחלוף הזמן, הסוכנים מתאימים עצמם לשינויים הדינמיים בסביבתם על פי הניסיון שצברו.³⁰

עקב טבעם המבוזר והאינטראקטיבי, "סוכנים חכמים" נועדו לאסוף מידע על רשתות שלמות והמערכות ההיקפיות. נדמה כי מאפיין זה שימש לא רק כאמצעי הגנה, אלא גם לצורך איסוף מידע וביצוע פריצות (ראו לעיל שרשרת התקיפה בסייבר) במערכות המהוות יעד פוטנציאלי למתקפה.³¹ מכיוון שההתנהגות של כל

סוכן מתעצבת על פי ניסיונו בתוך סביבתו האישית, לא פשוט להגן מפני איומים פרטניים מסוג זה.



תרשים 2. "סוכנים חכמים" של משטרת הסייבר המפיקים התרעה מוקדמת על פי גישת האבטחה המשולבת

דרך מוצלחת להשתמש בסוכנים נגד מתקפות סייבר מבוזרות היא באמצעות בניית "משטרת סייבר" של "סוכנים חכמים". גישה כזו מתבססת על הרעיון של "סוכני משטרה" מלאכותיים בסביבת סייבר מוגדרת, שמטרתם היא לזהות פעילות זדונית הנעשית בדרך מבוזרת.³² כפי שהומחש בתרשים 2 לעיל, "סוכני משטרה" כאלה יכולים לסייע בהגנה כבר בשלבים הראשונים של מתקפת הסייבר.

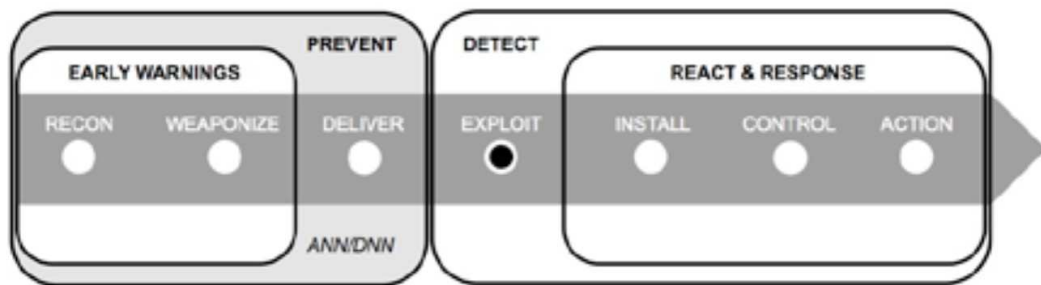
נוכחות של "סוכנים חכמים" ניתן למצוא גם במערכות חיסון מלאכותיות (AIS). שימוש בשני סוגים שונים של סוכנים – "סוכן זיהוי" ו"סוכן מתקפת נגד" – מאפשר לחקות את המאפיינים היעילים של מערכת החיסון האנושית. סוכני הזיהוי מנטרים את סביבות הסייבר ומנסים לאתר פעילות חריגה. כאשר סוכנים אלה מזהים פעילות זדונית, הם שולחים בצורה יזומה הוראות מבוזרות לסוכני מתקפת הנגד, שמופעלים כדי למנוע, למתן או אפילו להדוף את הפולשים לרשת.³³

רשתות עצבים מלאכותיות למניעת פריצות זדוניות

טכניקה נוספת שצמחה מהתחום של בינה מלאכותית היא רשת העצבים המלאכותית (ANN). רשתות עצבים מלאכותיות הן מודלים של למידה סטטיסטית, המחקים את המבנה והתפקוד של המוח האנושי. הן יכולות לסייע בלמידה ובפתרון בעיות, במיוחד בסביבות שבהן האלגוריתמים או הכללים לפתרון הבעיה קשים לניסוח או לא ידועים. מכיוון שהתנהגות מערכת של רשת עצבים מלאכותית הינה חמקמקה וקשה להגדרה, מערכות כאלו נחשבות ל"קופסה שחורה" בלתי מוגדרת.³⁴

רשתות עצבים מלאכותיות שימשו בהצלחה בכל השלבים של גישת האבטחה המשולבת לאבטחת סייבר, והן יכולות להכיל גם את כל השלבים של שרשרת התקיפה בסייבר. רשת עצבים מלאכותית המשולבת באבטחת סייבר יכולה

לשמש לניטור התנועה ברשת. כפי שמתואר בתרשים 3 להלן, ניתן לזהות פריצות זדוניות כבר בשלב החדירה ולפני שההתקפה עצמה יוצאת לפועל.³⁵ זוהי אחת המטרות שאבטחת הסייבר שואפת אליה, שכן היכולת למנוע מתקפות סייבר עוד קודם להתרחשותן היא הישג משמעותי שעשוי לקדם את הרעיון המרכזי של הגנה היקפית.³⁶ רשתות עצבים מלאכותיות גם יכולות לשמש בהצלחה ללמידה מפעילויות וממתקפות שהרשת עברה, כדי למנוע התממשות של מתקפות עתידיות.



תרשים 3. רשתות עצבים מלאכותיות למניעת מתקפות במסגרת גישת האבטחה המשולבת

היתרון הגדול של שימוש ברשת עצבים מלאכותית, בהשוואה לטכניקות קונבנציונליות המשמשות להגנת סייבר, הוא יכולת הלמידה שלה. כאמור, המצב הרגיל הוא שדפוסים המתארים פעילויות רשת נורמליות וחריגות כאחת מוגדרים באופן ידני על ידי מומחי אבטחה, בהתבסס על הידע המקצועי שלהם. ניתן לאמן את רשתות העצבים המלאכותיות כך שיזהו דפוסים כאלה באופן אוטומטי, באמצעות נתוני עבר שהועברו ברשת.

גישת IDPS המתבססת על אנומליות הוכיחה כי ניתן להשתמש בהצלחה ברשת עצבים מלאכותית להערכת מידע כותרת של חבילות נתונים ברשת,³⁷ וזאת במטרה ללמוד דפוסים של התנהגות רשת נורמלית.³⁸ בשלב ההכנה הראשון, רשת העצבים המלאכותית מאומנת לזהות ולהכיר דפוסים של תכונות כותרת השייכים לתנועת רשת נורמלית. משלב זה והלאה, כל חבילת נתונים עתידית העוברת ברשת המנוטרת מושווית לדפוסים שנלמדו. כאשר תכונות של כותרות חבילת נתונים תואמות לדפוס הרגיל, הן מועברות כרגיל. חריגות במידע של כותרת חבילת הנתונים, באופן שלא מתאים לדפוס שנלמד, יסווגו על ידי IDPS כאירוע זדוני ויידחו. נמצא כי גישה ייעודית זו משיגה שיפור בשיעור הזיהוי הכולל של ניסיונות פריצה, מבלי להפיק התרעות על תוצאות חיוביות שגויות או תוצאות שליליות שגויות. יתר על כן, בעוד שמערכות IDPS מסורתיות, המבוססות על חתימה או על אנומליות, פועלות בעיקר נגד פריצות מוכרות, גישת מערכת העצבים המלאכותית

מגינה בהצלחה מפני מקרים של פריצות שאינן מוכרות עדיין. בסיכומו של דבר, רשתות עצבים מלאכותיות אמורות לתמוך באופן מעשי בבניית מערכת IDPS איתנה, מסתגלת ומדויקת.³⁹

ניטור רשתות עצבים מלאכותיות אינו מוגבל לשימוש במערכות IDPS. ניתן לעשותו באמצעות כל מערכת שמנטרת פעילויות רשת. גם "חומות אש", מערכות לזיהוי פריצה או "מרכזיות רשת" עושות שימוש ברשתות עצבים מלאכותיות כדי לסרוק תנועה נכנסת ויוצאת ברשת. סימולציה ניסיונית המבוססת על רשת עצבים מלאכותית הראתה שגם מאמץ מחשובי קטן למדי יכול לזהות מראש תשעים אחוזים מתוכנת ריגול.⁴⁰

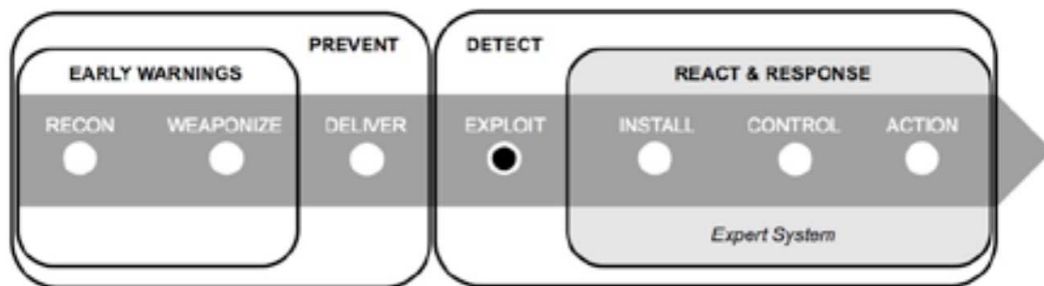
רשתות עצבים עמוקות (DNN), שהן צורה מפורטת, ממוחשבת ויקרה יותר של רשתות עצבים מלאכותיות,⁴¹ שימשו לאחרונה לא רק כדי להגן על ארגונים מפני מתקפות סייבר, אלא גם כדי לנבא מתקפות כאלו. השיפורים בחומרה הובילו להתקדמות בעיבוד נתונים במסגרת תשתיות רשת ולשיפור קיבולת אחסון, כך שטכנולוגיות רשתות עצבים עמוקות הפכו לנפוצות ולישימות יותר. פלטפורמת אבטחה ייעודית המבוססת על בינה מלאכותית, שעושה שימוש בגישה של רשת עצבים עמוקה, הוכיחה שהיא יכולה לנבא בהצלחה כ-85 אחוזים ממתקפות הסייבר.⁴² בעוד שגישות מסורתיות של אבטחת סייבר עברו מזהוי מתקפה למניעת מתקפה, טכניקות רשתות עצבים עמוקות יכולות אפוא להוביל לשלב חדש של אבטחת סייבר – חיזוי מתקפה.

מערכות מומחות כספקיות תמיכה בהחלטות של מומחי אבטחה

מערכות מומחות (Expert systems) הן תוכניות מחשב שנועדו לספק תמיכה בקבלת החלטות סביב בעיות מורכבות בתחום מסוים, ועושות שימוש רחב מאוד ביישומי בינה מלאכותית. הרעיון של מערכות מומחות מורכב מבסיס ידע שמאחסן את הידע המומחה, וממנוע הסקת מסקנות המשמש להפקת תובנות לגבי ידע שהוגדר מראש וכן לאיתור תשובות לבעיות שעל סדר היום.⁴³

מערכות מומחות מתייחסות לסוגים שונים של בעיות, וזאת בהתאם לאופן ההגעה לתובנות. גישה של גיבוש תובנות מבוססות מקרים (case-based reasoning) מאפשרת פתרון בעיות באמצעות שחזור של מקרים דומים מהעבר, בהנחה שהפתרון של מקרה העבר יכול להיות מאומץ ומיושם גם על המקרה החדש. בהמשך נשקלות הצעות לפתרונות חדשים, ואם נדרש הן מעודכנות, וכך מביאות לשיפור מתמשך ברמת הדיוק וביכולת להתמודד עם בעיות חדשות לאורך זמן. מערכות מבוססות כללים (rule-based systems) פותרות בעיות באמצעות כללים שהגדירו מומחים. הכללים מורכבים משני חלקים: תנאי ופעולה. הבעיות

מנותחות בשני שלבים: בשלב הראשון התנאי מוערך, ובשלב השני נקבעת הפעולה שיש לנקוט. שלא כמו במערכות מבוססות מקרים, מערכות מבוססות כללים אינן מסוגלות ללמוד כללים חדשים או לשנות אוטומטית כללים קיימים. עובדה זו קשורה ל"בעיית רכישת הידע", שהיא מכרעת לצורך התאמה לסביבות משתנות.⁴⁴ מומחי אבטחה עושים שימוש נרחב במערכות מומחות לצורך תמיכה בהחלטות בסביבות סייבר. הערכת נתוני הביקורת של מערכות אבטחה יכולה לקבוע באופן כללי אם פעילות של רשת או מערכת היא זדונית או תקינה. הכמות העצומה של הנתונים גורמת לכך שמומחי אבטחה עושים שימוש קבוע בדוחות סטטיסטיים כדי לסרוק ולנתח את כל מידע הביקורת בתוך פרק זמן סביר. מערכות מומחות המבוססת על בינה מלאכותית הוכיחו בהצלחה שהן יכולות לתמוך במאמצים אלה באמצעות ניטור בזמן אמת של סביבות הסייבר, אפילו במספר מערכות רב או במערכות הטרוגניות.⁴⁵ במקרים שבהם מזהה פריצה זדונית, מופקת הודעת התרעה. ההודעה מספקת מידע רלוונטי שמאפשר למומחי האבטחה לבחור את אמצעי האבטחה ביתר יעילות (react & respond – תגובה ומענה בתרשים 4 להלן). למרות כל הנאמר, חיוני לזכור כי אף שמערכות מומחות הצליחו לסייע למקבלי החלטות, הן אינן מסוגלות להחליף אותם.⁴⁶



תרשים 4. מערכות מומחות לתמיכה באמצעי תגובה ומענה על פי גישת האבטחה המשולבת

חסרונות הבינה המלאכותית בעולם אבטחת הסייבר

עד עתה דנו ביתרונות הבינה המלאכותית וכן בטכניקות השונות שיכולות לתת מענה לסוגיות טכנולוגיות חשובות בתחום אבטחת הסייבר. למרות היבטים חיוביים אלה, קיימים חששות וסיכונים המתלווים לשימוש בבינה מלאכותית בתחום אבטחת הסייבר:

- **חוסר יכולת לשמר אוטונומיה באבטחת הסייבר:** אף שנרשמה התקדמות עצומה באימוץ טכניקות של בינה מלאכותית באבטחת הסייבר, מערכות אבטחת

הסייבר עדיין אינן אוטונומיות במלואן ולא יכולות להוות תחליף מוחלט לקבלת החלטות אנושית. לכן, יש עדיין משימות שמחייבות התערבות של בני אדם.⁴⁷

- **פרטיות הנתונים:** טכניקות בינה מלאכותית, כמו רשתות עצבים מלאכותיות ועמוקות, הופכות למשוכללות יותר הודות להתקדמות בתחום החומרה. טכניקות חדשות כאלו צצות מדי יום. אולם לצורך הגובר בנתוני עתק (big data) יש גם צד שלילי כאשר הדברים מגיעים לפרטיות הנתונים: ניתוח כמויות עצומות של מידע עשוי לעורר בארגונים פרטיים וציבוריים גם יחד דאגה מפני פגיעה בפרטיות המידע האישי שהם מחזיקים, וחלקם אינם מוכנים על רקע זה לשתף את הנתונים שברשותם.⁴⁸ שאלות שיכולות להישאר ללא מענה עבור ארגונים שפרטיותם נפגעה הן: באילו נתונים אישיים נעשה שימוש? מדוע נעשה בהם שימוש? ואלו מסקנות מסיקים מהם במסגרת פתרונות מבוססי בינה מלאכותית?
- **היעדר רגולציה:** ישנם חששות משפטיים שונים בנוגע לבינה מלאכותית, אולם החשש הבולט ביותר הוא איבוד השליטה מצד בני האדם על ההשלכות הנובעות מבינה מלאכותית אוטונומית. טיבה הייחודי והקשה לחיזוי של הבינה המלאכותית גורם לכך שאין עדיין מסגרות משפטיות שיתנו מענה לתחום זה.⁴⁹
- **חששות מתחום האתיקה:** מערכות אבטחה מבוססות בינה מלאכותית מבצעות יותר ויותר החלטות עבור אנשים ספציפיים או מסייעות להם להגיע להחלטות (למשל, במקרה שנדון לעיל בנוגע למערכות מומחות). לנוכח התפתחות זו, מדאיג במיוחד שאין עדיין קוד אתי למערכות כאלו. התוצאה היא שהחלטות שמתקבלות עבורנו באמצעות בינה מלאכותית אינן בהכרח החלטות שאנו היינו מקבלים כבני אדם.⁵⁰

מסקנות

בינה מלאכותית נחשבת לאחת ההתפתחויות המבטיחות ביותר בעידן המידע, ונראה כי אבטחת סייבר הוא התחום שיכול להפיק מכך את המרב. אלגוריתמים, טכניקות, יוזמות וכלים חדשים המציעים שירותים המבוססים על בינה מלאכותית צצים ללא הרף בשוק האבטחה הגלובלי. בהשוואה לפתרונות אבטחת סייבר קונבנציונליים, מערכות מבוססות בינה מלאכותית אמורות להיות גמישות, מסתגלות ועמידות יותר, ובכך לסייע לשפר את ביצועי האבטחה ולספק הגנה טובה יותר למספר הולך וגובר של איומי סייבר מתוחכמים. נכון להיום, טכניקות של "למידה עמוקה" (deep learning) מסתמנות ככלים העוצמתיים והמבטיחים ביותר בתחום הבינה המלאכותית. רשתות עצבים עמוקות יכולות לחזות מתקפות סייבר מראש ולא רק למנוע אותן, וכך לפתוח שלב חדש באבטחת הסייבר. במקביל להבטחה הגלומה בבינה מלאכותית, היא גם מתגלה כסיכון גלובלי לאנושות. לסיכונים ולחששות המתלווים לבינה המלאכותית יש על מה להתבסס.

ניתן לזהות במסגרתם ארבע בעיות מרכזיות: היעדר אוטונומיה מלאה של הבינה המלאכותית; חשש לפגיעה בפרטיות הנתונים; היעדר מסגרת משפטית מספקת לטיפול בסוגיה; חששות אתיים שמקורם בהיעדר קוד מוסרי למערכות קבלת החלטות אוטונומיות. עקב המהירות בה מתפתח תחום הבינה המלאכותית, חיוני לפתור סיכונים וחששות אלה מהר ככל האפשר, אך נוכח האפשרות הסבירה שפתרונות בני קיימא לא יימצאו במהרה, מומלץ לעשות בשלב זה שימוש מושכל בבינה מלאכותית במסגרת אבטחת סייבר. דבר זה יסייע לפחות למתן חלק מהסיכונים והחששות.

עד עתה לא נרשמה הצלחה גורפת בתחום הגנת הסייבר הנעשית על ידי בני אדם בלבד או על ידי בינה מלאכותית בלבד. למרות השיפורים הרבים שהבינה המלאכותית הכניסה לתחום זה, מערכות קשורות לא מסוגלות עדיין להתאים את עצמן באופן מלא ואוטומטי לשינויים בסביבתן, ללמוד את כל סוגי האיומים והמתקפות ולבחור וליישם באופן אוטומטי אמצעי נגד ייעודיים שיגנו מפני מתקפות אלו. לכן, בשלב הטכנולוגי הנוכחי, כדי שניתן יהיה להגיע לבשלות באבטחת הסייבר, יש לשמור על התלות ההדדית החזקה בין בינה מלאכותית ובין הגורם האנושי, ויותר מכך, נדרש מבט כולל על סביבת הסייבר של ארגונים. אבטחת סייבר אינה רק סוגיה טכנולוגית. היא נוגעת גם לרגולציה ולאופן שבו סיכוני אבטחה מטופלים. כדי להשיג ביצועי אבטחה אופטימליים, חיוני לשלב בין כל הפתרונות הטכניים, התהליכים הרלוונטיים ובני האדם, וכך להגיע למסגרת אחת של אבטחה משולבת. בסופו של דבר, הגורם האנושי הוא זה שקובע עדיין, ולא (רק) הטכנולוגיה.

הערות

1. סטודנט בוגר מדעי המחשב בשם Robert Tappen Morris כתב ב־1988 את תוכנת המחשב הראשונה שהופצה דרך האינטרנט: התולעת מוריס. התוכנה לא נועדה לגרום נזק אלא רק לאמוד את ממדי האינטרנט, אך שגיאה קריטית יצרה שינוי בתוכנה וגרמה לה לשגר את מתקפת מניעת השירות הראשונה.
2. בנוגע לחסרונות של מערכת מניעה ואיתור של פריצות (IDPS), ראו: Amjad Rehman and Tanzila Saba, "Evaluation of Artificial Intelligent Techniques to Secure Information in Enterprises", *Artificial Intelligence Review* 42, no. 4 (2014): 1029-1044, במיוחד הפרק "Performance Issues: IDS".
3. Selma Dilek, Hüseyin Çakır and Mustafa Aydın, "Applications of Artificial Intelligence Techniques to Combating Cyber Crimes: A Review", *International Journal of Artificial Intelligence & Applications* 6, no. 1 (2015): 21-39.
4. Enn Tyugu, "Artificial Intelligence in Cyber Defense", in *Proceedings of 3rd International Conference on Cyber Conflict (ICCC)*, 7-10 June 2011, Tallinn Estonia, eds. C. Czosseck, E. Tyugu and T. Wingfield (Tallinn, Estonia: CCD COE, 2011), pp. 95-105; Xiao-bin Wang, Guang-yuan Yang, Yi-chao Li and Dan

- Liu, "Review on the Application of Artificial Intelligence in Antivirus Detection System", *Cybernetics and Intelligent Systems* (2008): 506-509.
5. הקרן לאתגרים גלובליים מצביעה על בינה מלאכותית כאחד הסיכונים הגוברים שעלול לאיים על האנושות בעתיד. עוד בנושא זה ראו: Dennis Pamlin and Stuart Armstrong, "Global Challenges – Twelve risks that threaten human civilisation" (Global Challenges Foundation: 2015), <http://globalchallenges.org/wp-content/uploads/12-Risks-with-infinite-impact.pdf>.
6. Stuart Russell, Tom Dietterich, Eric Horvitz, Bart Selman, Francesca Rossi, Demis Hassabis, Shane Legg, Mustafa Suleyman, Dileep George and Scott Phoenix, "Research Priorities for Robust and Beneficial Artificial Intelligence: An Open Letter", *AI Magazine* 36, no. 4 (2015): 105-114.
7. "A History of Cybersecurity: How Cybersecurity Has Changed in the Last 5 Years", My Digital Shield, October 5, 2015, <http://www.mydigitalshield.com/history-cyber-security-cyber-security-changed-last-5-years/>.
8. Rehman and Saba, "Evaluation of Artificial Intelligent Techniques".
9. ישנן גישות שונות לתיאור השלבים השונים של מתקפת סייבר. מאמר זה מתייחס אל "שרשרת התקיפה בסייבר" (Cyber Kill Chain®), כפי שהוגדרה על ידי חברת "לוקהיד מרטין" ושימשה בהרחבה את קהילת הביטחון האמריקאית. עוד בנושא זה ראו: www.lockheedmartin.com.
10. Gabi Siboni, "An Integrated Security Approach: The Key to Cyber Defense", *Georgetown Journal of International Affairs*, May 7, 2015, <http://journal.georgetown.edu/an-integrated-security-approach-the-key-to-cyber-defense/>.
11. Tony Sager, "Killing Advanced Threats in Their Tracks: An Intelligent Approach to Attack Prevention", A SANS Analyst Whitepaper (2014), <https://www.sans.org/reading-room/whitepapers/analyst/killing-advanced-threats-tracks-intelligent-approach-attack-prevention-35302>.
12. המושג מודיעין סייבר הוא יותר מאשר נתונים גולמיים זמינים. הוא עוסק בהשגת מידע יישומי של נתונים ממוינים, מעובדים ומוערכים. עוד בנושא מודיעין על איומי סייבר וההגדרה של מודיעין סייבר ומידע סייבר, ראו: "What Is Cyber Threat Intelligence and Why Do I Need It?" iSIGHT Partners Inc. (2014), http://www.isightpartners.com/wp-content/uploads/2014/07/iSIGHT_Partners_What_Is_20-20_Clarify_Brief1.pdf.
13. Siboni, "An Integrated Security Approach".
14. Ahmed Patel, Mona Taghavi, Kaveh Bakhtiyari and Joaquim Celestino Junior, "An Intrusion Detection and Prevention System in Cloud Computing: A Systematic Review", *Journal of Network and Computer Applications*, 36, Issue 1, January 2013, pp. 25-41, <http://dx.doi.org/10.1016/j.jnca.2012.08.007>.
15. Susan M. Bridges and Rayford B. Vaughn, "Fuzzy Data Mining and Genetic Algorithms Applied to Intrusion Detection", 12th Annual Canadian Information Technology Security Symposium (2000): 109-122.
16. Patel and others, "An Intrusion Detection and Prevention System in Cloud Computing: A Systematic Review".
17. Rehman and Saba, "Evaluation of Artificial Intelligent Techniques".
18. Sager, "Killing Advanced Threats in Their Tracks: An Intelligent Approach to Attack Prevention".

19. המועצה הפדרלית לבדיקת מוסדות פיננסיים (FFIEC) פיתחה את הכלי להערכת אבטחת סייבר כדי שיסייע לארגונים לזהות סיכונים ולקבוע היערכות נכונה לאבטחה. רמת המוכנות מסייעת לארגונים לקבוע האם ההתנהגויות, הנהלים והתהליכים שלהם תומכים בצורה נאותה בהיערכותם לאבטחת סייבר. היא גם מזהה פעולות אפשרויות שיגבירו את יעילות ההיערכות. עוד בנושא זה, ראו: <https://www.ffiec.gov/cyberassessmenttool.htm>.
20. Abdul Wahab Amirudin, "Facing Cyberattacks in 2016 and Beyond", The Star, January 28, 2016, <http://www.thestar.com.my/tech/tech-opinion/2016/01/28/facing-cyber-attacks-in-2016-and-beyond/>.
21. Rehman and Saba, "Evaluation of Artificial Intelligent Techniques".
22. Linda Ondrej, Todd Vollmer and Milos Manic, "Neural Network Based Intrusion Detection System for Critical Infrastructures", 2009 International Joint Conference on Neural Networks (Atlanta, GA, 2009): 1827-1834.
23. Yoshua Bengio, "Learning Deep Architectures for AI", Foundations and Trends® in Machine Learning, 2, no. 1 (2009): 1-127.
24. Tyugu, "Artificial Intelligence in Cyber Defense".
25. לאלגוריתמים קבועים חיווט לוגי קשיח, ברמת קבלת החלטות, לטובת ניתוח נתונים. ראה שם.
26. Olin Hyde, "Machine Learning for Cybersecurity at Network Speed & Scale", an Invitation to Collaborate on the Use of Artificial Intelligence against Adaptive Adversaries, ai-one (2011), www.ai-one.com/.
27. "רעש" – מידע לא מדויק או לא רלוונטי בנתונים שנאספו.
28. "Cybersecurity and Artificial Intelligence: A Dangerous Mix", INFOSEC Institute, February 24, 2015, <http://resources.infosecinstitute.com/cybersecurity-artificial-intelligence-dangerous-mix>.
29. אפשר לפרש ישות סייבר קוגניטיבית כתוכנית יחידה, בין אם חומרה או תוכנה, שיש לה יכולות קוגניטיביות בדומה לאדם. בתחום הבינה המלאכותית, יכולות קוגניטיביות של "סוכן חכם" יכללו תפיסה של סביבת הסייבר, השגה וניתוח של נתונים שנאספו במרחב הסייבר, והסקת מסקנות מנתונים אלה.
30. Stan Franklin and Art Graesser, "Is It an Agent, or Just a Program? A Taxonomy for Autonomous Agents", Third International Workshop on Agent Theories, Architectures and Languages (London: Springer-Verlag, 1997): 21-35.
31. Alessandro Guarino, "Autonomous Intelligent Agents in Cyber Offence", in 5th International Conference on Cyber Conflict, eds. K. Podins, J. Stinissen and M. Maybaum (Tallinn, Estonia: NATO CCD COE, 2013): 377-388.
32. Tyugu, "Artificial Intelligence in Cyber Defense".
33. Xia Ye and Junshan Li, "A Security Architecture Based on Immune Agents for MANET", International Conference on Wireless Communication and Sensor Computing (2010): 1-5.
34. Christian Bitter, David A. Elizondo and Tim Watson, "Application of Artificial Neural Networks and Related Techniques to Intrusion Detection", World Congress on Computational Intelligence (2010): 949-954.
35. שם.
36. Tyugu, "Artificial Intelligence in Cyber Defense".

37. כותרת של חבילת נתונים מכילה תכונות כמו אורך נתונים שהועברו, סוג פרוטוקול הרשת או כתובות המקור והיעד של חבילת הנתונים. לכן, כותרת החבילה נושאת מידע שיכול לשמש כדי להבדיל בין התנהגות רשת רגילה ובין ניסיונות חדירה.
38. Ondrej and others, "Neural Network Based Intrusion Detection System".
39. Bitter and others, "Application of Artificial Neural Networks and Related Techniques to Intrusion Detection".
40. ההדמיות הניסיוניות של זיהוי תוכנת הריגול שמו דגש על זיהוי תולעים וספאם. עוד בנושא זה, ראו: Dima Stopel, Robert Moskovitch, Zvi Boger, Yuval Shahr and Yuval Elovici, "Using Artificial Neural Networks to Detect Unknown Computer Worms", Neural Computing and Applications, 18, no. 7 (2009): 663-674.
41. Geoffrey E. Hinton, Simon Osindero and Yee-Whye Teh, "A Fast Learning Algorithm for Deep Belief Nets", Neural Computation, 18, no. 7 (2006): 1527-1554.
42. Victor Thomson, "Cyber Attacks Could Be Predicted With Artificial Intelligence", iTechPost, April 21, 2016, www.itechpost.com/articles/17347/20160421/cyber-attacks-predicted-artificial-intelligence-help.htm.
43. Tyugu, "Artificial Intelligence in Cyber Defense".
44. Serena H. Chen, Anthony J. Jakeman and John P. Norton, "Artificial Intelligence Techniques: An Introduction to Their Use for Modelling Environmental Systems", Mathematics and Computers in Simulation, 78, no. 2-3 (2008): 379-400.
45. שם.
46. Debra Anderson, Thane Frivold and Alfonso Valdes, "Next-Generation Intrusion Detection Expert System (NIDES): A Summary", SRI International, Computer Science Laboratory, May 1995.
47. Katherine Noyes, "A.I. + Humans = Serious Cybersecurity", Computerworld, April 18, 2016, www.computerworld.com/article/3057590/security/ai-humans-serious-cybersecurity.html.
48. Tom Simonit, "Microsoft and Google Want to Let Artificial Intelligence Loose on Our Most Private Data", MIT Technology Review, April 19, 2016, <https://www.technologyreview.com/s/601294/microsoft-and-google-want-to-let-artificial-intelligence-loose-on-our-most-private-data/>.
49. Bernd Stahl, David Elizondo, Moira Carroll-Mayer, Yingqin Zheng and Kutoma Wakunuma, "Ethical and Legal Issues of the Use of Computational Intelligence Techniques in Computer Security and Computer Forensics", The 2010 International Joint Conference on Neural Networks (IJCNN) (2010): 1-8.
50. Nick Bostrom, "Ethical Issues in Advanced Artificial Intelligence", in Cognitive, Emotive and Ethical Aspects of Decision Making in Humans and in Artificial Intelligence, eds. Iva Smit and George E. Lasker (Windsor, ON: International Institute for Advanced Studies in Systems Research/Cybernetics, 2003), 2, pp. 12-17.