

Cybersecurity and Law

2025 Nr 1(13)

DOI: 10.35467/cal/214606



AI usage in Cybersecurity

Zastosowanie sztucznej inteligencji w rozwiązaniach cyberbezpieczeństwa

Piotr ZABOROWSKI

Uczelnia Techniczno-Handlowa im. Heleny Chodkowskiej

ORCID: 0009-0008-1353-2176

E-mail: piotr.zaborowski@protonmail.com

Mateusz KOZŁOWSKI

Akademia Sztuki Wojennej

ORCID: 0009-0005-6899-0244

E-mail: m.kozlowski@doktorant.akademia.mil.pl

Abstract

The primary objective of this article is to discuss the potential applications of solutions based on artificial intelligence and machine learning in the process of detecting and countering threats in cyberspace. It describes the processes through which AI can enhance the operation of security systems by improving threat identification, automation, and the prediction of attack vectors. Both technical aspects and methods for implementing artificial intelligence solutions will be addressed. Additionally, legal and ethical concerns, as well as the risks that may arise from the large-scale deployment of such tools, will be presented.

An analysis of the technical aspects of AI/ML was conducted, including supervised and unsupervised learning, intrusion detection systems (IDS), EDR solutions, and user behavior analytics (UBA). A literature review and analysis of practical case studies (e.g., automated firewalls, phishing detection) were employed to evaluate the effectiveness of these solutions. Additionally, a critical approach was used to identify limitations, such as the quality of training data and the demand for computational power.

AI has been shown to significantly accelerate threat identification and enable the prediction of attack vectors, but its effectiveness depends on data quality and is susceptible to model hallucinations and false positives. Hybrid human-AI teams have been

identified as a key direction for development, ensuring a balance between automation and human intuition in security management.

The integration of human and algorithmic systems is essential for effectively countering modern cybersecurity threats. There is an urgent need to create coherent ethical and legal frameworks, as well as to invest in developing the competencies of technical staff and computational infrastructure.

Keywords

cybersecurity, machine learning, artificial intelligence, computational infrastructure

Streszczenie

Zasadniczym celem artykułu jest omówienie możliwości zastosowania rozwiązań opartych na sztucznej inteligencji i uczeniu maszynowym w procesie wykrywania i przeciwdziałania zagrożeniom występującym w cyberprzestrzeni. Opisuje on procesy w których AI może usprawniać działanie systemów bezpieczeństwa poprzez poprawę identyfikacji zagrożeń, automatyzację oraz przewidywanie wektorów ataków. Poruszone zostaną zarówno aspekty techniczne, jak i sposoby implementacji rozwiązań sztucznej inteligencji. Dodatkowo przedstawione zostaną wątpliwości natury prawnej i etycznej oraz zagrożenia, które mogą wynikać z zastosowania tego typu narzędzi na szeroką skalę.

Przeprowadzono analizę technicznych aspektów AI/ML, w tym uczenia nadzorowanego i nienadzorowanego, systemów wykrywania włamań (IDS), rozwiązań EDR oraz analizy zachowań użytkowników (UBA). Zastosowano przegląd literatury i analizę przypadków praktycznych (np. automatyczne zapory sieciowe, wykrywanie phishingu) w celu oceny skuteczności rozwiązań. Dodatkowo, wykorzystano podejście krytyczne do identyfikacji ograniczeń, takich jak jakość danych treningowych i zapotrzebowanie na moc obliczeniową.

Wykazano, że AI znacząco przyspiesza identyfikację zagrożeń i umożliwia przewidywanie wektorów ataków, ale jej skuteczność zależy od jakości danych oraz jest narażona na halucynacje modeli i fałszywe alarmy. Hybrydowe zespoły ludzi i AI zostały określone jako kluczowy kierunek rozwoju, zapewniający równowagę między automatyzacją a ludzką intuicją w zarządzaniu zabezpieczeniami.

Integracja systemów ludzkich i algorytmicznych jest niezbędna do skutecznego przeciwdziałania nowoczesnym zagrożeniom cybernetycznym. Pilna konieczność tworzenia spójnych ram etyczno-prawnych oraz inwestycji w rozwój kompetencji kadry technicznej i infrastruktury obliczeniowej.

Słowa kluczowe

cyberbezpieczeństwo, uczenie maszynowe, sztuczna inteligencja, infrastruktura obliczeniowa

Uwagi ogólne

Choć sam pomysł sztucznej inteligencji sięga połowy ubiegłego wieku, sporo czasu minęło, zanim technologia ta znalazła zastosowanie w cyberbezpieczeństwie. W pierwszych czasach istnienia komputerów nikt nie zakładał możliwości istnienia wirusów oraz luk bezpieczeństwa. Komputery służyły obliczeniom matematycznym oraz nie były ze sobą połączone. W przypadku braku bezpośredniego połączenia między urządzeniami, wcześniejsi adepci technologii nie przejmowali się bezpieczeństwem urządzeń¹.

Wszystko to zaczęło zmieniać się w latach 90-tych XX wieku. W tym momencie urządzenia IT przestały być dostępne wyłącznie dla badaczy i bogatych profesjonalistów oraz dostępna moc obliczeniowa zrobiła się wystarczająca do treningu bardziej złożonych modeli uczenia maszynowego. Te dwie rzeczy przesądziły o rewolucji AI, która choć zaczęła się bardzo powoli, obecnie nabiera rozpędu i wkracza w coraz to nowsze dziedziny codziennego życia².

Wśród niezaprzeczalnych zalet w zastosowaniu AI w cyberbezpieczeństwie należy wyróżnić szybsze wykrywanie zagrożeń, poprawę skuteczności w ich wykrywaniu oraz łatwość w predykcji popularnych zagrożeń w przyszłości. Mimo szeregu usprawnień, jakie sztuczna inteligencja przynosi w cybersferze, nie można zapominać o problemach, jakie pojawiają się w przypadku jej użycia. Wątpliwości budzą aspekty etyczne pozyskiwania danych - część z nich może pochodzić ze źródeł udostępnionych bez zgody użytkownika oraz pojawiają się głosy nad wykrywaniem nietypowych zagrożeń - modele AI korzystają często ze starych danych, wobec czego zupełnie nowe wektory ataku mogą być pomijane.

Sama koncepcja zabezpieczania się przed zagrożeniami w cyberprzestrzeni pojawiła i rozpowszechniła się dopiero w latach 90-tych. Wcześniej komputery były dostępne wyłącznie dla profesjonalistów, oraz co najważniejsze, nie były ze sobą bezpośrednio połączone. Dynamiczny rozwój internetu dał możliwość przestępcom różnego rodzaju ataków na urządzenia IT, co poskutkowało powstaniu pierwszych wirusów komputerowych. Te rozwijały się stopniowo, zaczynając od prostych, złośliwych aplikacji pisanych dla rozgłosu lub uwagi. Obecnie znaczna część zagrożeń opiera się na inżynierii społecznej³, co utrudnia wykrycie przez tradycyjne metody (antywirusy, firewalle). Zmiana dotknęła również popularności samych zagrożeń. Jeszcze dwie dekady temu wiadomości phishingowe były zaledwie ułamkiem zagrożeń, które czekały na nieświadomych

¹ J. Dzierżyńska-Mielczarek, Rynek mediów w Polsce. Zmiany pod wpływem nowych technologii cyfrowych. Warszawa 2018, s. 76.

² Computer and Internet Use in the United States: 2021, <https://www.census.gov/newsroom/press-releases/2024/computer-internet-use-2021.html> [dostęp: 28.10.2025].

³ H. Jahankhani, A. Al.-Nemrat, A. Hosseinian-Far, Cybercrime classification and characteristics [w:] B. Akhgar, A. Staniforth, F. Bosco (red.), Cyber Crime and Cyber Terrorism. Investigator's Handbook, New York 2014, s. 150; K. Markiewicz, Rozwój społeczeństwa informacyjnego [w:] E. Szczepaniuk, M. Gawlik-Kobylińska, Werner, J. (red.), Bezpieczny rozwój społeczeństwa informacyjnego, Warszawa 2016, s. 180.

użytkowników⁴. O bezpieczeństwie myślało się jak o jednorazowym produkcie - użytkownik kupował oprogramowanie antywirusowe i czuł się bezpiecznie. Obecnie świat IT uległ gwałtownej zmianie⁵. Najpopularniejszymi urządzeniami w internecie nie są komputery a smartfony, co znacząco zmieniło schematy zagrożeń. Najpopularniejszym rodzajem przestępstw są obecnie phishingi, czyli strony wyludzające dane osobiste lub bankowe poprzez imitowanie wyglądu innych stron z dobrą reputacją⁶. Malware traci na znaczeniu, gdyż w systemach mobilnych instalacja oprogramowania spoza oficjalnych źródeł jest utrudniona. W przypadku systemu Android, użytkownik jest zmuszony do przejścia skomplikowanej procedury. W pierwszej kolejności musi włączyć tryb deweloperski, następnie zezwolić na instalację aplikacji spoza zaufanych źródeł, potwierdzić swoją decyzję, a następnie przed instalacją aplikacji, zaakceptować jej wymagania. W przypadku systemu iOS, dodawanie aplikacji spoza oficjalnego sklepu jest praktycznie niemożliwe.

Sztuczna inteligencja to koncept znany od dziesięcioleci. Pierwsze próby zaprojektowania tzw. „cyfrowego mózgu” przypadają na połowę XX wieku. Mimo dziesięcioleci prób wykorzystywania sztucznej inteligencji, AI zaczęło zyskiwać popularność dopiero w nowym mileniu. Pierwsze próby przynosiły bardzo mizerne rezultaty, z uwagi na niską moc obliczeniową urządzeń IT w tamtym czasie⁷. Dynamiczny rozwój technologii umożliwił powstanie narzędzi, na których można polegać. Mimo szeregu niezaprzeczalnych zalet tej technologii należy pamiętać o jej ograniczeniach.

Pierwszym problemem jest jakość wejściowych danych. Każdy model uczenia maszynowego jest tak dobry jak dane, na których został nauczony. W przypadku niskiej jakości wykorzystanych informacji, AI nie będzie potrafiło wygenerować dobrych wyników. Kolejną niedogodnością może być tworzenie nieistniejących rzeczy tzw. halucynacje. Niektórzy nazywają tą przypadłość również „objawami wytwórczymi”. W przypadku prostych zadań sztuczna, inteligencja wygeneruje dobre rezultaty, a kiedy napotka przypadek, który jest zbyt trudny - odmienny od danych pierwotnych - wynik będzie niemiarodajny. W przypadku rozwiązań typu client-side (operacje wykonywane po stronie klienta), wynik fałszywie pozytywny przeważnie nie prowadzi do poważnych konsekwencji. Większym problemem są rozwiązania server-side (operacje, które dzieją się na serwerze). Dla przykładu, serwer DNS korzystający z rozwiązań AI

⁴ Zob. szerzej: F. Radoniewicz, Phishing [w:] K. Chałubińska-Jentkiewicz (red.), *Leksykon cyberbezpieczeństwa*, Warszawa 2024, s. 198-199; M. Nowikowska, Identity Theft. Protection of Personal Data in Cyberspace [w:] L. Tafaro, I. Laki, I. Florek (red.), *Digital well-being – a concern for the quality of life* Warsaw-Budapest 2023, s. 148-164.

⁵ T. C. Helmus, Artificial Intelligence, Deepfakes, and Disinformation: A Primer, „Perspective. Expert insights on a timely policy issue” 2022, No. 6, s. 2.

⁶ M. Nowikowska, M. Kozłowski, Analiza wybranych zagrożeń dla tożsamości cyfrowej wynikających z postępu technologicznego, „Journal of Modern Science” 2025, nr. 3, vol. 63, s. 787.

⁷ T. Mucci, The history of AI, <https://www.ibm.com/think/topics/history-of-artificial-intelligence> [dostęp: 28.10.2025]

może zablokować użytkownikom dostęp do strony na podstawie błędnej predykcji, w wyniku czego użytkownicy tracą do niej dostęp.

Uczenie maszynowe

Uczenie maszynowe (ang. *machine learning*)⁸ jest jedną z najbardziej rewolucyjnych technologii w tworzeniu nowoczesnych strategii walki z cyberzagrożeniami, szczególnie w aspekcie ich wykrywania. W przeciwieństwie do klasycznych narzędzi, które korzystają z predefiniowanych reguł lub znanych sygnatur wirusów, AI może dostosowywać się do nowych ataków, reagować i dostosowywać wzorce, czy wykrywać anomalie w gigantycznych zbiorach danych. Tego typu możliwości czynią oprogramowanie oparte na uczeniu maszynowym perfekcyjnie dostosowane do ciągle zmieniającej się cyberprzestrzeni, szczególnie w obliczu pojawiania się exploitów typu zero-day, czy nowych typów malware.

Wydaje się, że poprzez proaktywne i defensywne działania, AI jest wręcz stworzone do wykorzystania w celu poprawienia poziomu bezpieczeństwa IT, zarówno w dużych organizacjach, jak i pośród użytkowników indywidualnych. Głównym elementem narzędzi opartych na uczeniu maszynowym jest rozpoznawanie wzorców (ang. *pattern recognition*). Modele, które zostały wytrenowane na dużych zbiorach danych metodą uczenia nadzorowanego⁹ mogą zaklasyfikować działanie plików i użytkowników jako szkodliwe lub neutralne w sposób prawie idealny. Wiele rodzajów algorytmów, takich jak drzewa decyzyjne (ang. *decision tree*), czy sieci neuronowe, są używane do wykrywania sygnatur malware, ataków typu phishing czy wykrywania podejrzanych prób logowania¹⁰.

Modele wyuczone metodą uczenia nienadzorowanego są szczególnie użyteczne w wykrywaniu nowych, nieznanych typów zagrożeń. Choć efektywność tego typu predykcji jest mniejsza, tego typu narzędzia wciąż mogą dostarczyć wartościowych informacji, szczególnie w przypadku analizy ruchu sieciowego, czy monitorowania zachowania urządzenia.

Mechanizmy sztucznej inteligencji są obecnie wykorzystywane w wielu dziedzinach cyberbezpieczeństwa. Systemy wykrywania włamań (ang. *intrusion detection system*) korzystają z tego typu rozwiązań w celu analizowania ruchu sieciowego w czasie rzeczywistym, identyfikując podejrzane działania, które w innym przypadku nie zostałyby zauważone. Liczne oprogramowania EDR wykorzystują AI w celu poprawy jakości działania, poprzez monitorowanie

⁸ K. Liszczyk, Artificial intelligence i machine learning – wykorzystanie algorytmów i analiz do profilowania i personalizacji przekazów dystrybuowanych w sieci [w:] Dylematy współczesnej informatyki ekonomicznej. Teoria i praktyczne zastosowania, Politechnika Częstochowska, Częstochowa 2022, s. 42; A. Król-Nowak, K. Kotarba, Wprowadzenie [w:] A. Król-Nowak (red.), Podstawy uczenia maszynowego, Kraków 2022, s. 9.

⁹ A. Król-Nowak, K. Kotarba, Reprezentacja danych [w:] A. Król-Nowak (red.), Podstawy uczenia maszynowego, Kraków 2022, s. 15-16.

¹⁰ Zob. M. Nowikowska, M. Kozłowski, op.cit., s. 791.

działania plików w systemie operacyjnym, zapobiegając atakom malware czy ransomware. Dostawcy usług chmurowych również korzystają z ML w celu znalezienia przejętych kont, czy analizy ruchu sieciowego¹¹.

Pomimo widocznych zalet użycia uczenia maszynowego, eksperci zauważają problemy z wykorzystaniem tylko tej technologii. Największą z nich jest dostosowanie cyberprzestępców do nowych rozwiązań, gdzie mogą oni manipulować danymi tak, aby AI nie wykrywało zagrożeń. Wykrywanie fałszywie pozytywnych (ang. false positive) i fałszywie negatywnych (ang. false negative) zagrożeń także stanowi widoczny problem. Dodatkowo, użycie AI często zwiększa zapotrzebowanie na moc obliczeniową, co przekłada się na wymierne straty finansowe i środowiskowe.

Rosnąca liczba zagrożeń cyberbezpieczeństwa oraz ich coraz bardziej zróżnicowany typ powoduje potrzebę rozwijania sposobu obsługi incydentów.¹² W tradycyjnym modelu obsługi zgłoszeń operator obsługujący incydent potrzebuje wykonać wiele czasochłonnych prac, od poprawnego wykrycia, przez skategoryzowanie zagrożenia, do podjęcia odpowiednich działań mających na celu jego wyeliminowanie. z uwagi na coraz większą liczbę zgłoszeń tego typu, działania przestają być efektywne, z uwagi na opóźnienia wynikającego z konieczności ręcznej obsługi zgłoszenia. Brak podjęcia natychmiastowej próby neutralizacji ataku może spowodować zwiększone straty wizerunkowe i finansowe¹³. Atak może zostać zreplikowany na inne systemy, do czego nie doszłoby w przypadku szybszego zneutralizowania zagrożenia. Systemy automatycznej obsługi incydentów korzystają z narzędzi AI w celu przyspieszenia procesu oraz mogą być użyte do sprawdzenia poprawności danych. Wiele akcji można wykonać automatycznie, z ingerencją operatora w określonych, nietypowych przypadkach.

Narzędzia służące do automatyzacji procesu obsługi zgłoszeń opierają się na uczeniu maszynowym w celu wykrywania zagrożeń, oceniania ich ważności oraz zapobieganiu im. Wiele z tych procesów odbywa się w czasie rzeczywistym, przez co znacząco obniżają czas, w którym podejmowane są działania. Idealnym przykładem automatycznego systemu chroniącego użytkowników jest firewall, który pod wykryciu podejrzanego ruchu sieciowego, dokona analizy i zablokuje niebezpieczne połączenie.¹⁴ Dzięki użyciu sztucznej inteligencji proces ten może odbywać się szybciej i co najważniejsze, mieć wyższą wiarygodność. Kolejnym przykładem może być automatyczne wykrywanie stron typu phishing. AI potrafi zinterpretować kod strony, porównać wygląd, szatę graficzną czy odnośniki

¹¹ B. Dellot, B. Brhmie Balaram, Machine learning, „RSA Journal” 2018, vol. 164, no. 3 (5575), s. 44.

¹² Zob. szerzej: Zespół CERT Polska, Raport roczny 2024 z działalności CERT Polska. Krajobraz bezpieczeństwa polskiego internetu, Warszawa 2025.

¹³ C. Troy, K.G. Smith, M.A. Domino, CEO demographics and accounting fraud: Who is more likely to rationalize illegal acts?, „Strategic Organization” 2011, no. 4(9), s. 259–282.

¹⁴ J. Castro, J. Paine, R. Truong, How Cloudflare is using automation to tackle phishing head on, <https://blog.cloudflare.com/how-cloudflare-is-using-automation-to-tackle-phishing/> [dostęp 25.10.2025]

z bazą wyuczonych wzorców i stwierdzić, czy strona podszywa się oraz jaka firma bądź marka jest atakowana.

User Behavior Analytics

UBA (ang. *User Behavior Analytics*) to kolejne narzędzie, które może zostać użyte do zwiększenia poziomu cyberbezpieczeństwa. Przeważnie jest ono stosowane w większych korporacjach, w celu identyfikacji nietypowych działań podejmowanych przez użytkowników systemu. Tradycyjne oprogramowanie nie wykryje przejętego konta, czy próby kradzieży danych przez pracownika, w przeciwieństwie do UBA. Tego typu software stale monitoruje działanie użytkowników, rejestruje nie tylko odwiedzone strony, ale także czas wpisywanego hasła, jakie pliki są przeglądane, jaki jest sposób obsługi myszki i klawiatury, z jakich aplikacji korzysta użytkownik. Tego typu oprogramowanie potrafi także śledzić, skąd użytkownik uzyskuje dostęp do systemu, z jakiego adresu IP korzysta, jakie urządzenia zostało użyte¹⁵.

Wskazuje się, że cyberataki stają się coraz bardziej zaawansowane, a tradycyjne metody zabezpieczeń okazują się niewystarczające. UBA oferuje dodatkową warstwę ochrony, umożliwiając wczesne wykrycie i reagowanie na zagrożenia, zanim wyrządzą one szkodę. Główną różnicą między UBA a tradycyjnymi metodami bezpieczeństwa jest skupienie na zachowaniach użytkowników, a nie tylko na zabezpieczeniach opartych na sygnaturach, czy znanych zagrożeniach. Rozwiązania UBA analizują zachowania użytkowników w kontekście bezpieczeństwa, wykrywają anomalie i identyfikują zagrożenia w sposób proaktywny, nawet te wcześniej nieznane. Dlatego analiza behawioralna w cyberbezpieczeństwie jest tak bardzo ważna.

Oprogramowanie tego typu może stwarzać problemy, m. in. poprzez zbyt rygorystyczne ustawienia czułości prób wykrycia ataku, co może spowodować dużą ilość zgłoszeń fałszywie pozytywnych, pozbawiając użytkowników dostępu do systemu. Systemy tego typu przetwarzają ogromne ilości danych dotyczących zachowań i zwyczajów użytkowników, generując problemy natury etycznej.

Wnioski

Narzędzia oparte na sztucznej inteligencji stały się integralną częścią procesu ochrony przed zagrożeniami cyberbezpieczeństwa. Oferują one szybszą i często bardziej dokładną detekcję zagrożeń w porównaniu do tradycyjnych rozwiązań bezpieczeństwa. Możliwość przeprowadzania analizy na wielkich zbiorach danych, wykrywania anomalii ruchu sieciowego i przewidywanie nowych ataków czyni je niezwykle atrakcyjnym, zarówno dla użytkowników biznesowych, jak i indywidualnych. W centrach cyberbezpieczeństwa (SOC/CERT)

¹⁵ M. Kosiński, What is user behavior analytics (UBA)?, <https://www.ibm.com/think/topics/user-behavior-analytics> [dostęp: 14.11.2025].

przyspieszają obsługę zgłoszeń, ułatwiają klasyfikację i skracają czas reakcji poprzez podejmowanie automatycznych działań. Należy jednak pamiętać o niedoskonałościach, które posiada tego typu software, takich jak ryzyka związane z przetwarzaniem dużej ilości danych, zwiększone zapotrzebowanie na moc obliczeniową oraz możliwość prób oszustwa sztucznej inteligencji przez cyberprzestępców. Mimo tych obaw niemal każdy nie ma wątpliwości, że tego typu oprogramowanie będzie cały czas rozwijane i coraz szerzej wykorzystywane w sferze cyberbezpieczeństwa.

Dynamiczny rozwój technologii sztucznej inteligencji w sektorze bezpieczeństwa cyfrowego wskazuje na konieczność kompleksowego podejścia do integracji systemów opartych na uczeniu maszynowym z istniejącą infrastrukturą bezpieczeństwa. Przeprowadzone rozważania pozwalają uznać, że przyszłość ochrony systemów informatycznych będzie w znacznej mierze zależna od zdolności organizacji do efektywnego łączenia tradycyjnych rozwiązań bezpieczeństwa z zaawansowanymi algorytmami predykcyjnymi.

Kluczowym wyzwaniem staje się zapewnienie interoperacyjności między różnorodnymi systemami AI, a już wdrożoną infrastrukturą bezpieczeństwa. Organizacje muszą inwestować w platformy, które umożliwią skuteczną integrację narzędzi rozpoznawania zagrożeń. Dodatkowo, wraz z rosnącym zaangażowaniem AI w procesy bezpieczeństwa, pojawia się konieczność dostosowania do rosnącej siatki regulacji prawnych. Unijne rozporządzenie o sztucznej inteligencji wprowadza rygorystyczne wymagania dotyczące przejrzystości systemów decyzyjnych. Organizacje wdrażające rozwiązania oparte o uczenie maszynowe w cyberbezpieczeństwie muszą zapewnić, że decyzje podejmowane przez algorytmy są nadzorowane i nie naruszają praw użytkowników systemów.

Podsumowując powyższe można wskazać, że przyszłość cyberbezpieczeństwa nie polega na całkowitym zastąpieniu eksperta człowieka systemami sztucznej inteligencji, ale na budowaniu hybrydowych zespołów, w których algorytmy zwiększają zdolności analityczne specjalistów ds. bezpieczeństwa. Inwestycja w infrastrukturę AI musi być połączona z równoczesnym rozwojem kompetencji kadr, standardów interoperacyjności oraz ram etyczno-prawnych, które zapewnią zaufanie do systemów i maksymalizują ich potencjał obronny.

Bibliografia

- Araya D., Cybersecurity and AI. Artificial Intelligence for Defence and Security, „Centre for International Governance Innovation” 2022.
- Bingley R., Combatting Cyber Terrorism: A Guide to Understanding the Cyber Threat Landscape and Incident Response Planning, „IT Governance Publishing” 2024.

- Castro J., Paine J., Truong R., How Cloudflare is using automation to tackle phishing head on, <https://blog.cloudflare.com/how-cloudflare-is-using-automation-to-tackle-phishing/>.
- Computer and Internet Use in the United States: 2021, <https://www.census.gov/newsroom/press-releases/2024/computer-internet-use-2021.html>.
- Dellot B., Brhmie B., Machine learning, "RSA Journal" 2018, vol. 164, no. 3 (5575).
- Dzierżyńska-Mielczarek J., Rynek mediów w Polsce. Zmiany pod wpływem nowych technologii cyfrowych. Warszawa 2018.
- Helmus T.C., Artificial Intelligence, Deepfakes, and Disinformation: A Primer, „Perspective. Expert insights on a timely policy issue” 2022, No. 6.
- Jahankhani H., Al-Nemrat A., Hosseinian-Far A., Cybercrime classification and characteristics [w:] B. Akhgar, A. Staniforth, F. Bosco (red.), Cyber Crime and Cyber Terrorism. Investigator's Handbook, New York 2014.
- Kosiński M., What is user behavior analytics (UBA)?, <https://www.ibm.com/think/topics/user-behavior-analytics>.
- Król-Nowak A., Kotarba K., Reprezentacja danych [w:] A. Król-Nowak (red.), Podstawy uczenia maszynowego, Kraków 2022.
- Król-Nowak A., Kotarba K., Wprowadzenie [w:] A. Król-Nowak (red.), Podstawy uczenia maszynowego, Kraków 2022.
- Liszczyk K., Artificial intelligence i machine learning – wykorzystanie algorytmów i analiz do profilowania i personalizacji przekazów dystrybuowanych w sieci [w:] Dylematy współczesnej informatyki ekonomicznej. Teoria i praktyczne zastosowania, Politechnika Częstochowska, Częstochowa 2022.
- Markiewicz, K., Rozwój społeczeństwa informacyjnego [w:] E. Szczepaniuk, M. Gawlik-Kobylińska, Werner, J. (red.), Bezpieczny rozwój społeczeństwa informacyjnego, Warszawa 2016.
- Mucci T., The history of AI, <https://www.ibm.com/think/topics/history-of-artificial-intelligence>.
- Nowikowska M., Identity Theft. Protection of Personal Data in Cyberspace [w:] L. Tafaro, I. Laki, I. Florek (red.), Digital well-being – a concern for the quality of life Warsaw-Budapest 2023.
- Nowikowska M., Kozłowski M., Analiza wybranych zagrożeń dla tożsamości cyfrowej wynikających z postępu technologicznego, „Journal of Modern Science” 2025, nr. 3, vol. 63.
- Radoniewicz F., Phishing [w:] K. Chałubińska-Jentkiewicz (red.), Leksykon cyberbezpieczeństwa, Warszawa 2024.
- Troy C., Smith K.G., Domino M.A., CEO demographics and accounting fraud: Who is more likely to rationalize illegal acts?, "Strategic Organization" 2011, no. 4(9).
- Zespół CERT Polska, Raport roczny 2024 z działalności CERT Polska. Krajobraz bezpieczeństwa polskiego internetu, Warszawa 2025.