

Zdzisław Rus

Szkoła Doktorska

Akademia WSB w Dąbrowie Górniczej

DOI: <https://doi.org/10.35464/1642-672X.PS.2025.3.03>

Sztuczna inteligencja: współtowarzysz czy zagrożenie dla ludzkości?

Artificial Intelligence: Companion or Threat to Humanity?

ABSTRACT: The article analyzes the development, applications, and social consequences of artificial intelligence, focusing on the balance between its potential and threats. Using literature analysis and case studies from medicine, education, environmental protection, and industry, the author identifies key risk factors such as job automation, data privacy loss, algorithmic bias, and military uses. The study highlights the lack of sufficient legal and ethical regulations to mitigate AI's negative effects. The conclusions emphasize the need for international frameworks ensuring responsible and human-centered AI development that integrates innovation with ethical accountability.

KEYWORDS: computer science, automation, ethics of technology, legal regulations, artificial intelligence.

STRESZCZENIE: Autor analizuje rozwój, zastosowania i konsekwencje społeczne sztucznej inteligencji (AI), koncentrując się na relacji między jej potencjałem a zagrożeniami. Stosując analizę źródeł i przypadków użycia AI w medycynie, edukacji, ochronie środowiska oraz przemyśle, autor identyfikuje kluczowe czynniki ryzyka, takie jak automatyzacja pracy, utrata prywatności, uprzedzenia algorytmiczne czy militarne zastosowania technologii. Wskazuje także brak wystarczających regulacji prawnych i etycznych, które mogłyby ograniczyć negatywne skutki rozwoju AI. We wnioskach podkreśla konieczność tworzenia międzynarodowych ram odpowiedzialnego wykorzystania sztucznej inteligencji, łączących innowacyjność z ochroną wartości humanistycznych.

SŁOWA KLUCZOWE: informatyka, automatyzacja, etyka technologii, regulacje prawne, sztuczna inteligencja.

Rozwój sztucznej inteligencji (Artificial Intelligence, AI) reprezentuje jedno z najbardziej znaczących osiągnięć w dziedzinie technologii naszej ery. Ta dynamiczna ewolucja AI, od jej początkowych form do obecnych zaawansowanych systemów, odzwierciedla ogromny potencjał tej technologii. AI, rozumiana jako zdolność maszyn do wykonywania zadań wymagających ludzkiej inteligencji, obejmuje szeroki zakres aplikacji, od prostych algorytmów do złożonych systemów uczenia maszynowego. Jej zdolność do analizowania dużych zbiorów danych, uczenia się z doświadczeń oraz adaptowania się do nowych sytuacji sprawia, że jest ona nie tylko narzędziem o niezliczonych zastosowaniach, ale także polem dla innowacji i postępu w wielu dziedzinach.

Jednakże, z szybkim postępowaniem AI wiążą się również obawy o potencjalne zagrożenia i wyzwania etyczne. Te obawy koncentrują się nie tylko na technicznych aspektach AI, ale również na społecznych i moralnych implikacjach jej rozwoju. Z jednej strony AI może przynieść ogromne korzyści, takie jak poprawa efektywności, dokładności i szybkości w wykonywaniu zadań, które do tej pory wymagały ludzkiej interwencji, z drugiej strony nasuwa się pytanie o wpływ AI na rynek pracy, prywatność, bezpieczeństwo i moralność decyzji podejmowanych przez maszyny. Ta dwustronna natura AI czyni ją tematem intensywnych debat, zarówno w środowisku naukowym, jak i społecznym.

Niniejszy esej ma na celu zbadanie różnorodnych aspektów AI, biorąc pod uwagę zarówno jej potencjalne korzyści, jak i ryzyka. Zagłębiając się w historię AI, jej obecne zastosowania oraz przewidywane trajektorie rozwoju, można lepiej zrozumieć, jak ta rewolucyjna technologia kształtuje teraźniejszość i jak może wpłynąć na przyszłość. Przyjrano się, w jaki sposób AI wpływa na różne sfery życia, od biznesu i przemysłu po edukację i zdrowie, a także ocenimy jej potencjalne skutki dla społeczeństwa jako całości.

Ostatecznie, celem jest przedstawienie zbilansowanego obrazu sztucznej inteligencji, uwzględniającego zarówno pozytywne, jak i negatywne konsekwencje jej rozwoju dla przyszłości ludzkości. Poprzez analizę faktów, trendów i różnych perspektyw, esej ten dąży do głębszego zrozumienia złożoności i wielowymiarowości AI, podkreślając jej znaczenie w kształtowaniu świata, w którym żyjemy. W ten sposób, starano się dostarczyć przemyślaną i informacyjną podstawę do dalszej dyskusji i refleksji na temat roli AI we współczesnej rzeczywistości i jej przyszłych implikacji.

Historia AI ma swoje korzenie w latach 50. XX w., kiedy termin ten zaczął być postrzegany nie tylko jako fascynujący koncept naukowy, ale również jako realne pole badawcze. W tym okresie naukowcy i inżynierowie zaczęli eksperymentować z pierwszymi programami komputerowymi, które miały naśladować niektóre aspekty ludzkiej inteligencji. Jednym z takich wczesnych

programów była ELIZA, stworzona w środowisku Massachusetts Institute of Technology (MIT) przez Josepha Weizenbauma [Bostrom, 2016, s. 56–57]. Program ten był zdolny do prowadzenia prostych konwersacji w języku naturalnym, choć jego możliwości były ograniczone. ELIZA używała sztuczki zwanej „odbijaniem”, polegającej na przekształcaniu wypowiedzi użytkownika i zadawaniu pytań, co sprawiało wrażenie zrozumienia i empatii. Innym przełomowym osiągnięciem tamtych czasów był perceptron, wczesna forma sieci neuronowej, opracowany przez Franka Rosenblatta. Perceptron, choć prosty w konstrukcji, zapoczątkował rozwój algorytmów uczenia maszynowego, demonstrując, że maszyny mogą uczyć się i dostosowywać na podstawie danych wejściowych [Bostrom, 2016, s. 58–59].

Te wczesne eksperymenty stanowiły kamień milowy w rozwoju AI, demonstrując potencjał maszyn do wykonywania zadań, które wcześniej uważano za domenę wyłącznie ludzkich umysłów. Mimo że te pierwsze programy były proste w porównaniu do dzisiejszych standardów, ich znaczenie dla rozwoju AI jest nie do przecenienia. Ustanowiły one podstawy, na których oparto dalsze badania i rozwój w tej dziedzinie. Zapoczątkowały nową erę w informatyce, otwierając drzwi do badań nad bardziej zaawansowanymi algorytmami i zastosowaniami AI [CyberDefence24, 2020]. W tym okresie rozwinęły się również kluczowe koncepcje, takie jak sztuczne sieci neuronowe, algorytmy uczenia maszynowego i przetwarzanie języka naturalnego, które do dziś są podstawą dla wielu zaawansowanych aplikacji sztucznej inteligencji. Te wczesne osiągnięcia pokazały, że maszyny mogą nie tylko wykonywać skomplikowane obliczenia, ale także w pewnym stopniu naśladować ludzką zdolność do uczenia się, przetwarzania języka i nawet prowadzenia dialogu [CyberDefence24, 2020].

W latach 70. i 80. XX w. obszar AI doświadczył znaczącego postępu, który był w dużej mierze napędzany przez rozwój w dziedzinie algorytmów i wzrost mocy obliczeniowej komputerów. Te dekady były świadkami początków przełomu w uczeniu maszynowym i sieciach neuronowych. Zanim nastąpiły te zmiany, AI skupiała się głównie na regułach bazujących na logice i programowaniu symbolicznym. Jednak z czasem, w miarę jak naukowcy zaczęli lepiej rozumieć potencjał sieci neuronowych i uczenia maszynowego, zaczęto eksplorować nowe metody, które umożliwiały maszynom uczenie się z danych, zamiast polegania na sztywno zakodowanych instrukcjach. Wprowadzenie i rozwój algorytmów takich jak backpropagation w latach 80. przyczyniło się do przełomu w trenowaniu sieci neuronowych, umożliwiając im adaptowanie się i uczenie z doświadczenia, co było krokiem milowym w kierunku stworzenia bardziej zaawansowanych form AI [CyberDefence24, 2020].

W tym okresie, szczególnie ważne było zastosowanie sieci neuronowych w różnych dziedzinach i problemach. Przykładem może być rozwój systemów ekspertowych, które wykorzystywały AI do symulowania rozumowania ekspertów w określonych dziedzinach, jak medycyna czy geologia. Te systemy, choć nadal ograniczone, stanowiły ważny krok w kierunku automatyzacji procesów decyzyjnych i analizy złożonych danych [Fehler, 2017, s. 133–134]. Ponadto, w latach 80. XX w. zaczęto dostrzegać potencjał AI w takich obszarach jak przetwarzanie języka naturalnego i rozpoznawanie wzorców, co później zaowocowało rozwojem technologii, które dzisiaj są powszechnie stosowane w rozpoznawaniu mowy i analizie obrazów. Te postępy, choć wydawały się wtedy ograniczone, stały się fundamentem dla późniejszych, bardziej zaawansowanych aplikacji AI i otworzyły drogę dla przyszłych innowacji w tej dziedzinie [Fehler, 2017, s. 134–135].

Na przełomie XX i XXI w. AI wkroczyła w nową erę, napędzaną przez rozwój potężnych komputerów i dostęp do ogromnych zbiorów danych. To właśnie w tym czasie kluczową rolę zaczęły odgrywać techniki głębokiego uczenia, które wykorzystują wielowarstwowe sieci neuronowe do modelowania złożonych wzorców i zależności w danych [Fehler, 2017, s. 135–136]. Dzięki temu AI mogła podejmować się zadań dotychczas uznawanych za nieosiągalne dla maszyn. Przełomowe stało się wykorzystanie głębokich sieci neuronowych, które, dzięki swojej zdolności do uczenia się z nieustrukturyzowanych danych, zrewolucjonizowały wiele dziedzin. W szczególności, zaobserwowano znaczący postęp w dziedzinie rozpoznawania mowy i przetwarzania języka naturalnego. Systemy takie jak Siri czy Alexa, oparte na zaawansowanych algorytmach AI, stały się codziennością, oferując interakcje z użytkownikami, które były niegdyś domeną *science fiction*. Ich zdolność do zrozumienia i reagowania na ludzką mowę w naturalny sposób wyznaczała nowe standardy w interakcjach człowiek–maszyna [McKinsey&Company, 2018].

Równocześnie AI zaczęła odgrywać kluczową rolę w medycynie, otwierając nowe możliwości w diagnostyce i personalizowanej medycynie. Zastosowanie algorytmów uczenia maszynowego w analizie obrazów medycznych, takich jak skany MRI czy tomografie komputerowe, umożliwiło szybszą i dokładniejszą diagnozę, co jest kluczowe w leczeniu wielu schorzeń. Ponadto, integracja AI z genetyką i biologią molekularną zaczęła przyczyniać się do rozwoju spersonalizowanych terapii lekowych, dostosowanych do indywidualnych cech genetycznych pacjentów. Ta synergia między AI a medycyną nie tylko przyspieszyła postęp w leczeniu chorób, ale także otworzyła nowe horyzonty dla badań nad przedłużaniem życia i poprawą jego jakości. W ten sposób, na przełomie wieków, AI wyraźnie pokazała swój potencjał jako narzędzie

przynoszące znaczące korzyści ludzkości, zarówno w codziennym życiu, jak i w aspektach związanych ze zdrowiem i dobrym samopoczuciem [McKinsey&Company, 2018].

Obecnie AI nie ogranicza się jedynie do zadaniowych aplikacji. Sztuczna inteligencja wkroczyła również w świat sztuki i kreatywności, tworząc niezwykle dzieła i asystując artystom w ich pracy. Algorytmy, takie jak te stosowane w generatywnych sieciach przeciwnych (GANs), pozwalają na tworzenie realistycznych obrazów i muzyki, co otwiera nowe horyzonty dla artystycznej ekspresji [McKinsey&Company, 2018].

Rozwój AI niesie ze sobą potencjalne zagrożenia, które mogą mieć znaczący wpływ na rynek pracy i strukturę społeczeństwa. Jednym z najbardziej palących zagadnień jest rosnąca automatyzacja, która, choć zwiększa efektywność i produktywność, równocześnie rodzi obawy o przyszłość zatrudnienia. Automatyzacja, wzmocniona przez AI, może prowadzić do eliminacji wielu tradycyjnych miejsc pracy, szczególnie tych, które są rutynowe i wymagają małej elastyczności lub kreatywności. Przykładem mogą być prace w sektorze produkcji, gdzie roboty i algorytmy AI już zastępują ludzką pracę. Ten trend może się nasilać, co wiąże się z ryzykiem wzrostu bezrobocia, szczególnie wśród osób o niższym wykształceniu lub umiejętnościach, które trudniej dostosować do zmieniających się wymagań rynku pracy [Pastwa, Siudak, Sztokfisz, 2019, s. 117–118].

Ponadto, automatyzacja może przyczynić się do pogłębiania nierówności społecznych. W miarę jak AI i robotyka przejmują więcej zadań, dochodzi do zwiększenia popytu na wysoko wykwalifikowanych pracowników, którzy mogą projektować, utrzymywać i nadzorować te systemy. Tymczasem pracownicy o niższych kwalifikacjach, wykonujący zadania łatwe do zautomatyzowania, mogą być zmuszeni do poszukiwania pracy w coraz bardziej konkurencyjnym i ograniczonym rynku. Ten podział rynku pracy może prowadzić do społeczno-ekonomicznej segregacji, gdzie grupa wysoko wykwalifikowanych pracowników korzysta z benefitów płynących z technologicznego postępu, podczas gdy inni doświadczają niestabilności zatrudnienia i związanych z tym konsekwencji finansowych i społecznych [Pastwa, Siudak, Sztokfisz, 2019, s. 117–118].

W obliczu tych wyzwań ważne jest, aby społeczeństwa i rządy rozważyły strategie adaptacyjne. Możliwe rozwiązania obejmują inwestycje w edukację i przekwalifikowanie pracowników, aby lepiej przygotować ich do pracy w coraz bardziej zautomatyzowanym świecie. Ponadto, istnieje potrzeba opracowania nowych modeli zatrudnienia i systemów zabezpieczenia społecznego, które mogą lepiej radzić sobie z niestabilnym rynkiem pracy. Te kroki są kluczowe

w minimalizowaniu negatywnych skutków automatyzacji, a jednocześnie maksymalizowaniu korzyści płynących z postępu technologicznego. Wprowadzenie takich zmian może pomóc stworzyć bardziej zrównoważone i sprawiedliwe społeczeństwo, które efektywnie wykorzystuje potencjał AI, jednocześnie chroniąc i wspierając swoich członków w obliczu tych bezprecedensowych zmian [Raport NASK, 2019].

Użycie AI w celach wojskowych stanowi kolejny obszar potencjalnych zagrożeń, budząc obawy o nowe formy konfliktów i zmiany w charakterze wojny. Rozwój autonomicznych systemów zbrojeniowych, które wykorzystuje AI do podejmowania decyzji w czasie rzeczywistym na polu walki, stawia przed nami etyczne i strategiczne dylematy. Technologie te, od dronów bojowych po zaawansowane systemy rozpoznawania i namierzania, zwiększają efektywność działań wojennych, ale jednocześnie rodzą pytania o odpowiedzialność za decyzje podejmowane przez maszyny. Problem ten jest szczególnie palący, gdyż systemy oparte na AI mogą działać z większą prędkością i na większą skalę niż człowiek, co może prowadzić do eskalacji konfliktów i zwiększenia ryzyka nieprzewidzianych skutków [Różanowski, 2007, s. 211–212].

Ponadto, rozwój autonomicznych systemów zbrojeniowych może doprowadzić do nowego wyścigu zbrojeń, w którym kluczową rolę będą odgrywać zaawansowane technologie AI. Może to nie tylko zdestabilizować globalny ład bezpieczeństwa, ale także skupić znaczne zasoby na rozwoju militarnym kosztem innych potrzebnych inwestycji, takich jak edukacja czy opieka zdrowotna. Istnieje również obawa, że zwiększone poleganie na systemach AI w strategii wojskowej może prowadzić do zmniejszenia kontroli ludzkiej nad decyzjami wojennymi, co z kolei stwarza ryzyko nadużyć i błędów. Te wyzwania wymagają międzynarodowej współpracy i prawnych regulacji ograniczających rozwój i użycie autonomicznych systemów zbrojeniowych, zapewniając, że technologie AI będą stosowane w sposób etyczny i zgodny z międzynarodowym prawem humanitarnym.

Kwestie prywatności i bezpieczeństwa danych w erze AI stanowią istotne wyzwanie, zwłaszcza w obliczu rosnącego zastosowania algorytmów AI do masowego zbierania i analizy danych. Współczesne systemy AI, wymagające dostępu do ogromnych ilości danych do efektywnego działania, mogą prowadzić do naruszeń prywatności, gdy informacje osobowe są zbierane, przetwarzane i wykorzystywane bez wystarczającej zgody lub świadomości użytkowników. Te dane, często zawierające wrażliwe informacje, mogą być używane do różnych celów, począwszy od personalizacji reklam, aż po bardziej kontrowersyjne zastosowania, takie jak profilowanie społeczne czy nadzór. Takie praktyki rodzą obawy o to, jak i przez kogo te dane są wykorzystywane, a także

o potencjalne skutki nadużyć, włącznie z możliwością manipulacji i kontroli społecznej [Rózanowski, 2007, s. 211–212].

Kolejne znaczące zagrożenie to bezpieczeństwo tych danych. W miarę jak AI staje się coraz bardziej zintegrowana z codziennym życiem, rośnie ryzyko cyberataków i wycieków danych, co może mieć poważne konsekwencje zarówno dla jednostek, jak i dla organizacji. Problematiczne staje się również wykorzystanie AI do zaawansowanych ataków hakerskich i tworzenia bardziej wyrafinowanych metod cyberprzestępczości. Odpowiednie zabezpieczenia i regulacje dotyczące gromadzenia i wykorzystania danych stają się zatem kluczowe, aby zapewnić ochronę prywatności i bezpieczeństwo w cyfrowym świecie. Ważne jest również to, aby użytkownicy byli świadomi potencjalnych ryzyk i mieli kontrolę nad swoimi danymi, co wymaga przejrzystości i odpowiedzialności ze strony twórców i operatorów systemów AI [Sima i in., 2020, s. 350–351].

Etyczne dylematy związane z decyzjami podejmowanymi przez maszyny stanowią jedno z najbardziej skomplikowanych zagadnień w obszarze AI. Szczególnie istotne stają się te kwestie w kontekście autonomicznych pojazdów, które muszą podejmować decyzje w sytuacjach wymagających oceny moralnej. Przykładem takiego dylematu jest problem „wagonika”, gdzie autonomiczny pojazd musi zdecydować, jak zareagować w sytuacji nieuchronnego wypadku – czy ma chronić życie pasażerów kosztem pieszych, czy odwrotnie. Te decyzje nie są tylko kwestią algorytmów i kodowania, ale także wpływają na fundamentalne kwestie moralne i etyczne. Problem ten wykracza poza tradycyjne inżynierie i programowanie, wymaga zaangażowania filozofów, etyków, a także społeczeństwa w szerszą debatę o tym, jakie wartości powinny być zakodowane w systemach AI [Sima i in., 2020, s. 350–351].

Dylematy etyczne dotyczące AI nie ograniczają się do autonomicznych pojazdów. Wiele innych zastosowań AI, takich jak systemy rekomendacyjne, rozpoznawanie twarzy czy algorytmy wykorzystywane w wymiarze sprawiedliwości, również rodzi pytania o sprawiedliwość, uprzedzenia i transparentność. Istnieje ryzyko, że systemy AI mogą nieświadomie propagować uprzedzenia lub dyskryminację, na przykład poprzez wykorzystanie danych, które odzwierciedlają istniejące stereotypy społeczne. Wyzwaniem staje się zatem zapewnienie, że AI podejmuje decyzje w sposób sprawiedliwy i bezstronny, co wymaga stałego monitorowania, oceny i dostosowywania algorytmów. Tylko poprzez takie podejście możliwe będzie stworzenie systemów AI, które działają na korzyść całego społeczeństwa, uwzględniając różnorodność i złożoność ludzkich wartości i potrzeb [Bostrom, 2003].

Sztuczna inteligencja przynosi znaczące korzyści w dziedzinie medycyny, otwierając nowe możliwości dla lepszej i bardziej spersonalizowanej opieki

zdrowotnej. Dzięki AI możliwe stało się przetwarzanie i analiza ogromnych zbiorów danych medycznych, co przyczynia się do szybszej i dokładniejszej diagnostyki. Algorytmy uczenia maszynowego są wykorzystywane do analizy obrazów medycznych, takich jak skany MRI czy tomografie komputerowe, co pozwala na wczesne wykrywanie chorób, w tym nowotworów, z większą precyzją niż kiedykolwiek wcześniej. Ponadto AI ma kluczowe znaczenie w rozwijaniu personalizowanej medycyny, pozwalając na dostosowanie leczenia do indywidualnych cech genetycznych i zdrowotnych pacjenta. Takie podejście nie tylko zwiększa skuteczność terapii, ale także minimalizuje ryzyko skutków ubocznych, co ma ogromne znaczenie dla jakości życia pacjentów [Bostrom, 2003, s. 4–5].

Innym aspektem, w którym AI przynosi korzyści medycynie, jest zdolność do przewidywania epidemii i zarządzania informacjami zdrowotnymi w czasie rzeczywistym. Systemy oparte na AI mogą analizować dane z różnych źródeł, włączając w to dane demograficzne, środowiskowe i kliniczne, aby przewidzieć rozprzestrzenianie się chorób i pomóc w planowaniu odpowiedzi zdrowotnych. Dodatkowo AI odgrywa ważną rolę w badaniach klinicznych, przyspieszając proces odkrywania nowych leków i terapii. Dzięki zdolności do analizy złożonych wzorców i symulacji, AI pomaga naukowcom w identyfikowaniu potencjalnych nowych leków i w zrozumieniu ich interakcji z ludzkim organizmem, co przyspiesza rozwój nowych metod leczenia. Wszystkie te postępy wskazują na ogromny potencjał AI w przekształcaniu i ulepszaniu opieki zdrowotnej, oferując nadzieję na zdrowszą i bardziej bezpieczną przyszłość [Bostrom, 2003, s. 5].

Sztuczna inteligencja znajduje coraz szersze zastosowanie w ochronie środowiska jako ważne narzędzie w walce ze zmianami klimatycznymi i zarządzaniu zasobami naturalnymi. Dzięki zdolności do analizy dużych zbiorów danych, AI umożliwia naukowcom i ekologom lepsze zrozumienie i modelowanie skomplikowanych systemów środowiskowych. W dziedzinie zmian klimatycznych algorytmy AI są wykorzystywane do przewidywania przyszłych trendów klimatycznych, analizy wpływu działań ludzkich na środowisko oraz oceny skuteczności różnych strategii łagodzenia zmian klimatycznych. Na przykład, poprzez modelowanie scenariuszy emisji gazów cieplarnianych i ich wpływu na temperaturę globalną AI pomaga w opracowywaniu bardziej precyzyjnych planów redukcji emisji i adaptacji do zmieniającego się klimatu [Różanowski, 2017, s. 109].

Ponadto AI odgrywa kluczową rolę w zarządzaniu zasobami naturalnymi, wspomagając monitorowanie i ochronę różnorodnych ekosystemów. Systemy oparte na AI mogą na przykład analizować dane z satelitów i czujników

terenowych, aby monitorować degradację lasów, zanieczyszczenie wód czy migracje dzikich zwierząt. Ta zdolność do zbierania i przetwarzania danych w czasie rzeczywistym pozwala na szybkie reagowanie na zagrożenia ekologiczne, takie jak wylesianie, kłusownictwo czy nadmierne wykorzystanie zasobów wodnych. AI wspiera także rozwój zrównoważonych praktyk rolniczych, pomagając rolnikom w optymalizacji wykorzystania wody, nawozów i pestycydów, co przyczynia się do zmniejszenia śladu ekologicznego rolnictwa. Wszystkie te zastosowania pokazują, że AI ma ogromny potencjał do przyczyniania się do ochrony środowiska, wspierając globalne wysiłki na rzecz zrównoważonego rozwoju i ochrony naszej planety dla przyszłych pokoleń [Rózanowski, 2017, s. 110].

Sztuczna inteligencja wnosi rewolucyjne zmiany do edukacji, oferując możliwości dostosowania procesu nauczania do indywidualnych potrzeb uczniów. Dzięki niej nauczyciele i edukatorzy mogą korzystać z personalizowanych systemów edukacyjnych, które analizują styl nauki i postępy każdego ucznia, przygotowując materiały i tempo nauczania do ich unikalnych potrzeb. Takie systemy mogą identyfikować obszary, w których uczniowie mają trudności, i dostarczać dodatkowe zasoby lub ćwiczenia, aby wspierać ich rozwój. To indywidualne podejście umożliwia uczniom efektywniejszą naukę, jednocześnie zachęcając ich do aktywnego uczestnictwa w procesie edukacyjnym.

Sztuczna inteligencja ma potencjał do znacznego wzbogacenia doświadczeń edukacyjnych poprzez dostarczanie interaktywnych i angażujących materiałów dydaktycznych. Na przykład, zaawansowane narzędzia AI mogą tworzyć symulacje i wirtualne środowiska, które umożliwiają uczniom eksplorowanie złożonych koncepcji naukowych lub historycznych w bardziej interaktywny sposób. Dzięki temu uczniowie mogą lepiej zrozumieć i zapamiętać materiał, a także rozwijać kluczowe umiejętności, takie jak krytyczne myślenie i rozwiązywanie problemów. Wprowadzenie AI do edukacji otwiera także nowe możliwości dla uczniów z różnorodnymi potrzebami edukacyjnymi, w tym dla osób z niepełnosprawnościami, oferując im dostosowane narzędzia i wsparcie, które umożliwiają pełniejszy udział w procesie nauczania. W ten sposób AI staje się kluczowym narzędziem w tworzeniu bardziej inkluzywnego i efektywnego środowiska edukacyjnego [Rózanowski, 2017, s. 111].

Omówienie ram regulacyjnych dotyczących sztucznej inteligencji jest kluczowym elementem w kontekście zapewnienia, że jej rozwój przebiega w sposób zrównoważony i bezpieczny dla społeczeństwa. Obecnie, zarówno na poziomie międzynarodowym, jak i narodowym, podejmowane są inicjatywy mające na celu stworzenie ram prawnych i etycznych dla AI. Międzynarodowe organizacje, takie jak Unia Europejska czy Organizacja Narodów

Zjednoczonych, pracują nad wytycznymi, które mają na celu zarówno promowanie innowacji w dziedzinie AI, jak i adresowanie kluczowych kwestii, takich jak prywatność, bezpieczeństwo danych, odpowiedzialność i etyka. Te wytyczne mają na celu stworzenie standardów, które będą wspierać przejrzystość, uczciwość i odpowiedzialność w projektowaniu i wdrażaniu systemów AI, jednocześnie promując ich pozytywne zastosowania [Chłopecki, 2018, s. 78–80].

Na poziomie narodowym wiele krajów opracowuje własne przepisy regulujące rozwój i zastosowanie AI. Takie inicjatywy obejmują na przykład prawodawstwo dotyczące autonomii pojazdów, wykorzystania algorytmów w procesach decyzyjnych czy też kwestii związanych z robotyką i automatyką. Ważnym elementem tych regulacji jest ustalenie, kto ponosi odpowiedzialność za działania i decyzje podejmowane przez AI, zwłaszcza w przypadku, gdy mogą one prowadzić do szkód lub naruszeń prawnych. Ponadto, istotne są również etyczne wytyczne, które pomagają w kształtowaniu odpowiedzialnego rozwoju AI, zwracając uwagę na ważne kwestie, takie jak niesprawiedliwość, dyskryminacja czy wpływ AI na rynek pracy. Włączenie tych aspektów do procesu regulacyjnego jest niezbędne, aby zapewnić, że AI będzie rozwijać się w sposób, który służy dobru całego społeczeństwa a jednocześnie minimalizuje potencjalne negatywne skutki [Chłopecki, 2018, s. 82–83].

Sztuczna inteligencja będąca jednym z najbardziej dynamicznie rozwijających się obszarów technologicznych, ma potencjał do znaczącego wpływu na różne aspekty życia społecznego i indywidualnego. Z jednej strony, stwarza ona niespotykane dotąd możliwości w zakresie medycyny, edukacji, ochrony środowiska i wielu innych dziedzin, oferując rozwiązania dla niektórych z najpilniejszych wyzwań naszych czasów. Z drugiej strony AI rodzi również poważne wyzwania, w tym zagrożenia dla prywatności, bezpieczeństwa danych, rynku pracy, a także kwestie etyczne związane z autonomią maszyn i ich wpływem na społeczne struktury.

Podsumowując, istotne jest odpowiedzialne kształtowanie rozwoju AI, z uwzględnieniem zarówno jej pozytywnych, jak i negatywnych aspektów. Wymaga to zintegrowanego podejścia, które obejmuje rozwój odpowiednich ram regulacyjnych, stałe monitorowanie wpływu AI na społeczeństwo, a także promowanie etycznych zasad w projektowaniu i wdrażaniu systemów AI. Kluczowe znaczenie ma tu współpraca między naukowcami, inżynierami, decydentami, a także społeczeństwami, aby zapewnić, że AI będzie rozwijać się w sposób zrównoważony i korzystny dla wszystkich. Ostatecznie, przyszłość AI i jej wpływ na ludzkość zależy od tego, jak podejmiemy do tych wyzwań i możliwości, i jak będziemy kształtować tę technologię w nadchodzących latach.

Bibliografia

- Bostrom N. (2003). Transhumanist Values. W: F. Adams (red.), *Ethical Issues for the 21st Century*. Philosophical Documentation Center Press, 3–4.
- Bostrom N. (2016). *Superinteligencja. Scenariusze, strategie, zagrożenia*. Helion, 56–57.
- Chłopecki A. (2018). *Sztuczna inteligencja – szkice prawnicze i futurologiczne*. Warszawa: Wydawnictwo C.H. Beck, s. 78–80.
- Fehler W. (2017). Sztuczna inteligencja – szansa czy zagrożenie? *Studia Bobolanum*, 28, 3, 133–134.
- Pastwa A., Siudak R., Sztokfisz B. (2019). *Sztuczna inteligencja. Made in Asia*. Kraków: Instytut Kościuszki, 117–118.
- Raport NASK (2019). *Sztuczna inteligencja w społeczeństwie i gospodarce*. Warszawa.
- Rózanowski K. (2007). Sztuczna inteligencja rozwój, szanse i zagrożenia. *Zeszyty Naukowe Warszawskiej Wyższej Szkoły Informatyki*, 2, 211–212.
- Rózanowski K. (2017). Sztuczna inteligencja: rozwój, szanse i zagrożenia. *Zeszyty Naukowe Warszawskiej Wyższej Szkoły Informatyki*, 2, s. 109.
- Sima V., Gheorghe I.G., Subić J., Nancu D. (2020). Influences of the Industry 4.0 Revolution on the Human Capital Development and Consumer Behavior: A Systematic Review. *Sustainability*, 12(10), 350–351.

Netografia

- CyberDefence24 (2020). Sztuczna Inteligencja – błogosławieństwo czy zagrożenie dla ludzkości? <https://www.cyberdefence24.pl/armiai-slugby/sztuczna-inteligencjablogoslawienstwo-czy-zagrozenie-dlaludzkości> [data pobrania: 20.11.2023].
- McKinsey&Company (2018). Ramię w ramię z robotem. Jak wykorzystać potencjał automatyzacji w Polsce, <https://www.mckinsey.com/pl/our-insights/ramie-w-ramie-z-robotem> [data pobrania: 20.11.2023].