

SZTUCZNA INTELIGENCJA W MEDYCYNIE: MOŻLIWOŚCI ZASTOSOWANIA A WYZWANIA ETYCZNE I PRAWNE

Beata Piekło

*Uniwersytet Ekonomiczny w Krakowie
e-mail: pieklobeata01@gmail.com, ORCID: 0009-0000-0241-9516*

Streszczenie: Zasadniczym celem niniejszej publikacji jest przedstawienie rosnącej roli sztucznej inteligencji w medycynie, analizując zarówno jej potencjał jak i wyzwania związane z regulacjami prawnymi i aspektami etycznymi. W artykule omówiono kluczowe obszary w jakich sztuczna inteligencja może usprawnić procesy diagnostyczne i terapeutyczne. Wskazane zostały zasady, według których powinno się wykorzystywać wspomnianą technologię w obszarach związanych ze zdrowiem. Zwrócono również uwagę na problem odpowiedzialności za błędy popełnione w związku z użyciem algorytmów sztucznej inteligencji. Przedstawiono zarówno wyzwania etyczne jak i kwestie prawne związane z rozwojem tej technologii w dziedzinach medycyny.

Słowa kluczowe: sztuczna inteligencja, medycyna, leczenie, ryzyko, algorytmy.

Wprowadzenie

W obecnych czasach niemal codziennie mamy do czynienia z nowoczesną technologią, w tym sztuczną inteligencją w różnych formach np. asystent Google, Voiceboty czy gry komputerowe. Od niedawna zaczęto również wdrażać systemy sztucznej inteligencji do medycyny i opieki zdrowotnej co stwarza wiele nowych możliwości w leczeniu i diagnostyce, ale jednocześnie niesie to za sobą znaczące zagrożenia. Dlatego, aby zminimalizować ryzyko wystąpienia tych zagrożeń, konieczne jest jak najszybsze wprowadzenie właściwych regulacji prawnych, które zapewnią odpowiedzialne korzystanie z tych rozwiązań. Niekiedy zdarza się, że innowacyjne technologie wyprzedzają obowiązujące przepisy i tak też dzieje się w przypadku stosowania sztucznej inteligencji w różnych obszarach medycyny w Polsce.

Nie ma obecnie jednej powszechnie przyjętej definicji sztucznej inteligencji, zarówno jeśli chodzi o terminologię specjalistyczną i fachową, jak i w doktrynie prawa. Potocznie sztuczną inteligencję (ang. *artificial intelligence*, *AI*) można określić jako zdolność systemu komputerowego do naśladowania ludzkich procesów poznawczych takich jak myślenie czy uczenie się. Zgodnie z definicją podaną przez Słownik języka polskiego PWN sztuczną inteligencją nazywamy „dział informatyki badający reguły rządzące zachowaniami umysłowymi człowieka i tworzący programy lub systemy komputerowe symulujące ludzkie myślenie”¹. Definicja ta ma w istocie charakter socjologiczny. Opisuje rozwiązanie o charakterze technicznym poprzez odniesienie do opinii ludzi na temat podobieństwa tego rozwiązania do inteligentnego działania. Należy przypisać temu charakter subiektywny, zależny od osoby oceniającej dane rozwiązanie i w związku z tym nie nadaje się do zastosowania np. w celu zbudowania spójnej regulacji prawnej sztucznej inteligencji². Jedne z pierwszych działań związanych ze sztuczną inteligencją miały miejsce w latach 50 dwudziestego wieku. W 1950r. Alan Turing opracował tzw. Test Turinga, który miał na celu odróżnienie ludzi od maszyn, a tym samym należy uznać, że już wtedy maszyny były w stanie symulować zachowanie człowieka. Na pierwszym seminarium poświęconym tematyce sztucznej inteligencji (tzw. Warsztaty w Dartmouth) jednym z gospodarzy był John McCarthy, któremu przypisuje się stworzenie pojęcia sztucznej inteligencji w 1955r. Zdefiniował on ją jako „konstruowanie maszyn, o których działaniu dałoby się powiedzieć, że są podobne do ludzkich przejawów inteligencji”. Warto wspomnieć również o definicji, którą zaproponowali Andreas Kaplan i Michael Haenlein. Wskazują oni, że sztuczna inteligencja to „zdolność systemu do prawidłowego interpretowania danych pochodzących z zewnętrznych źródeł, nauki na ich podstawie oraz wykorzystywania tej wiedzy, aby wykonywać określone zadania i osiągać cele poprzez elastyczne dostosowanie”. Ostatnia z podanych definicji jest w mojej ocenie najbardziej odpowiednia, gdyż pozostałe bazują na podobieństwie do inteligencji ludzkiej co jest bezcelowe i niepotrzebnie zwraca uwagę na kwestie osobowości sztucznej inteligencji. Sztuczna inteligencja może być zaprogramowana do symulowania ludzkich zachowań, ale nie ma własnej osobowości ani świadomości i nie jest w stanie odczuwać tak jak człowiek. W związku z tym i z powodu braku autonomiczności systemu brak jest podstaw do przypisania sztucznej inteligencji osobowości prawnej. Wolna wola i rozum są atrybutami człowieka, które determinują przyznanie mu zdolności prawnej i zdolności do czynności prawnych. Obecnie nie można ich przyznać sztucznej inteligencji, gdyż nie ma ona swojego własnego interesu prawnego w osiągnięciu określonego celu, ani nie ma zdolności do samodzielnego kierowania swoim postępowaniem. W Polsce na razie nic nie wskazuje na to żeby miało ulec to zmianie. Koncepcja wyrażona w „Polityce Rozwoju Sztucznej Inteligencji w Polsce na lata 2019-

¹ W. Doroszewski (red.), *Słownik języka polskiego*, <https://sjp.pwn.pl/sjp/sztuczna-inteligencja;2466532.html> [07.11.2014].

² T. Zalewski, *Prawo sztucznej inteligencji*, CH Beck, Warszawa 2020, s. 2.

2027” opiera się na supremacji człowieka nad systemami sztucznej inteligencji. Powinna być ona zorientowana na człowieku i jego środowisku (*the Human-Centric Approach on AI*). Jak zostało wskazane w załączniku nr 3 dokumentu należy przeciwstawiać się działaniom zmierzającym do nadania osobowości prawnej sztucznej inteligencji. Niektóre państwa wyrażają odmienne stanowisko i dokonały już czynności zmierzających do przyznania sztucznej inteligencji osobowości³. W miejscowości Haaselt w Belgii burmistrz w 2017 roku podpisał akt urodzenia dla robota o imieniu Fran, a tym samym nadała mu obywatelstwo belgijskie. Podobnie robot humanoidalny o imieniu Sophii po wystąpieniu na forum „Future Investment Initiative” w Arabii Saudyjskiej otrzymała obywatelstwo tego kraju.

Sztuczna inteligencja ma ogromny potencjał i szerokie zastosowanie w medycynie. AI już teraz jest pomocne w wykonywaniu różnorodnych zadań, takich jak np. diagnostyka obrazowa, przewidywanie chorób, prowadzenie badań i obsługa danych pacjentów. Algorytmy AI mogą analizować ogromne ilości danych dużo szybciej i dokładniej niż ludzie, co może przynieść wymierne korzyści w świecie medycyny. Jednak jej wdrożenie napotyka istotne wyzwania prawne, takie jak ochrona danych, odpowiedzialność za decyzje algorytmów czy regulacje dotyczące wyrobów medycznych. Dodatkowo pojawiają się kwestie etyczne, takie jak prywatność pacjenta, równość dostępu, autonomia w podejmowaniu decyzji czy ryzyko uprzedzeń. Kluczowe dla skutecznego i bezpiecznego wdrożenia systemów sztucznej inteligencji w medycynie będzie stworzenie odpowiednich ram prawnych i standardów etycznych, które zapewnią ochronę pacjentów i umożliwią pełne wykorzystanie możliwości, jakie oferuje ta technologia.

Rodzaje sztucznej inteligencji i przykłady jej zastosowania w medycynie

Sztuczną inteligencję można podzielić na kilka kluczowych kategorii. Najbardziej ogólny podział ze względu na poziom zaawansowania składa się z Narrow Artificial Intelligence (NAI), czyli Wąskiej Sztucznej Inteligencji, Ogólnej Sztucznej Inteligencji (General Artificial Intelligence, AGI) oraz Nadludzkiej Sztucznej Inteligencji (Artificial Superintelligence, ASI). Pierwszy z nich to najbardziej rozwinięty i najczęściej wykorzystywany typ AI, który specjalizuje się w realizacji konkretnych zadań. Działa w wąskim zakresie wyznaczonych dla siebie ram i nie potrafi wykroczyć poza granice swoich algorytmów. Przykładem NAI są systemy do analizy dużych zbiorów danych, algorytmy rekomendacji czy wirtualni asystenci jak np. Siri, a także programy rozpoznające obrazy. W medycynie najczęściej wykorzystywane są podrodzaje AI, które należą do kategorii Wąskiej Sztucznej Inteligencji. Należy do nich m. in. uczenie maszynowe (Machine Learning, ML), który wykorzystuje algorytmy do identyfikowania wzorców

³ Polityka Rozwoju Sztucznej Inteligencji w Polsce na lata 2019–2027, Warszawa 2019, s. 40–43 i s. 103; file:///C:/Users/48693/Downloads/Polityka_Rozwoju_Sztucznej_w_Polsce_na_lata_2019_2027.pdf [29.11.2024].

w zbiorach danych, a następnie tworzy za ich pomocą modele danych, które umożliwiają przewidywanie. W tym miejscu można wyróżnić 3 metody uczenia się AI, mianowicie uczenie nadzorowane, uczenie nienadzorowane, uczenie przez wzmacnianie. Na tym etapie należy podkreślić, że zarówno uczenie nienadzorowane jak i uczenie przez wzmacnianie są na innym poziomie rozwoju niż uczenie nadzorowane, gdyż wymagają bardziej zaawansowanych algorytmów i dużych zasobów oraz są trudniejsze do interpretacji. Obie te formy uczenia, mimo że wciąż pozostają w niektórych obszarach na etapie badań i rozwoju, znajdują swoje praktyczne zastosowanie w medycynie. W metodzie uczenia nadzorowanego AI uczy się z pomocą prawidłowego klucza odpowiedzi metodą „uczenia się na błędach”. System otrzymuje pewien zbiór danych i próbuje przewidzieć wynik, a następnie dostosowuje swoje procesy bazując na różnicy pomiędzy swoją odpowiedzią a odpowiedzią referencyjną. Ma to swoje zastosowanie w diagnozowaniu chorób na podstawie oznaczonych danych np. w analizie obrazów medycznych (jak wykrywanie guzów na tomografii komputerowej). W uczeniu nienadzorowanym nie ma podanej prawidłowej odpowiedzi tak jak w przypadku uczenia nadzorowanego, a algorytm ma za zadanie odkryć wzorce w zbiorze danych, w którym dane uczące są nieoznakowane. Algorytmy nienadzorowane przydają się przede wszystkim przy analizie danych, odkrywaniu nowych wzorców w dużych zbiorach danych czy segmentacji obrazów medycznych. Jest to przydatne np. przy grupowaniu pacjentów o podobnych profilach zdrowotnych lub przy analizie genomów. Ostatnia z metod, w przeciwieństwie do dwóch poprzednich nie przygotowuje się za pomocą danych uczących, a podejmuje decyzje poprzez interakcje z otoczeniem. Opiera się to na mechanizmie prób i błędów, gdzie AI uzyskuje „nagrody” za prawidłowe decyzje, a za złe podlega „karze”. Dzięki temu agent testuje różne działania i uczy się, które z nich prowadzą do lepszego rezultatu. Ta forma AI może być stosowana w celu np. optymalizacji planów leczenia lub pomóc w rozwoju narzędzi diagnostycznych. Jako podkategorie uczenia maszynowego wyróżniamy tzw. uczenie głębokie (*Deep Learning, DL*), które wykorzystuje głębokie sieci neuronowe inspirowane strukturą ludzkiego mózgu. Dzięki nim możliwa jest analiza dużych zbiorów bardziej skomplikowanych danych, takich jak obrazy i dźwięki, co jest szczególnie przydatne w radiologii, patomorfologii i rozpoznawaniu obrazów medycznych. Modele deep learning są wykorzystywane np. do analizy danych medycznych pacjentów, aby sugerować odpowiednio dobrane terapie lub do rozpoznawania zmian nowotworowych za pomocą tomografii komputerowej, rezonansu magnetycznego czy mammografii. Ma to swoje zastosowanie również w patologii cyfrowej, w diagnozie chorób neurologicznych, oftalmologii oraz robotyce medycznej. Ze względu na funkcjonalność czy podział techniczny wśród kategorii AI oprócz uczenia maszynowego wyróżniamy systemy eksperckie, planowanie, robotykę, przetwarzanie obrazu oraz mowy, a także przetwarzanie języka naturalnego, które z kolei dzieli się na klasyfikowanie tekstu, tłumaczenie maszynowe, odpowiadanie na pytania i generowanie mowy. Systemy eksperckie działają na podstawie baz wiedzy

zdefiniowanych przez ekspertów co jest przydatne m.in. w diagnostyce medycznej (są to np. systemy, które sugerują leczenie na podstawie objawów pacjenta, jego historii choroby i danych epidemiologicznych.) Planowanie odnosi się do optymalnych strategii działań w celu poprawy efektywności i wyników leczenia. Może być pomocne przy planowaniu zabiegów chirurgicznych, terapii czy w zarządzaniu opieką zdrowotną. Robotyka posługuje się robotami medycznymi, które wprowadzają powtarzalność i precyzję w stopniu ciężko osiągalnym przez ręce ludzkie. Jako przykład można przywołać tutaj model robota chirurgicznego stosowanego w wielu dziedzinach medycyny jak urologia, ginekologia i kardiochirurgia, który nazywany jest Systemem da Vinci. Szacuje się, że do 2022 roku przeprowadzono przy jego pomocy ponad 12 mln operacji. System da Vinci należy do urządzeń typu *master–slave*, gdzie w części sterowniczej (*master*) znajduje się konsola, natomiast część wykonawczą (*slave*) stanowi robot o czterech interaktywnych ramionach, które mają bezpośredni kontakt z pacjentem. Trzy ramiona wyposażone są w narzędzia chirurgiczne, a dwa z nich odpowiadają prawej i lewej dłoni chirurga. Czwarte ramię służy do sterowania kamerą endoskopową w ciele pacjenta. Przełomem w dostępie do robota da Vinci w Polsce był rok 2017, kiedy system otrzymał pozytywną opinię Agencji Oceny Technologii Medycznych i Taryfikacji. Prezes AOTMiT zarekomendował system chirurgiczny da Vinci jako świadczenie gwarantowane. W Polsce szpitale, które zostały wyposażone w system operacyjny da Vinci znajdują się m.in. w Białymstoku, w Warszawie czy we Wrocławiu. W Stanach Zjednoczonych system ten jest standardowo używany w leczeniu operacyjnym raka prostaty. Około 80% operacji prostatektomii wykonywanych jest z pomocą robota da Vinci. W ginekologii znajduje zastosowanie w wykonywaniu takich zabiegów jak np. histerektomia czy miomektomia. W kardiochirurgii umożliwia wykonanie wielu operacji zastawkowych czy też zabiegów na tętnicach wieńcowych. Stosowanie systemu da Vinci przynosi wiele korzyści- zapewnia komfort pracy chirurgom oraz wysoką precyzyjność, a także zmniejsza ryzyko wystąpienia niepożądanych działań podczas zabiegu. Dzięki wykorzystaniu kamery, która generuje obraz 3D w wysokiej rozdzielczości (HD) możliwe jest uzyskanie powiększonego obrazu pola operacyjnego z wysoką szczegółowością. Największą wadą tego systemu są wysokie koszty zakupu i utrzymania⁴. Przetwarzanie mowy polega na zastosowaniu algorytmów AI, które rozpoznają, analizują, interpretują i generują ludzką mowę w formie dźwiękowej lub tekstowej. Przykładem użycia jest np. telemedycyna i zdalne konsultacje albo asystenci głosowi dla lekarzy. Przetwarzanie obrazu opiera się na wykorzystaniu algorytmów analizy obrazu, które rozpoznają, klasyfikują i interpretują dane wizualne, takie jak obrazy z badań diagnostycznych (np. RTG) czy zdjęcia. Jest to przydatne np. w segmentacji obrazu, monitorowaniu postępu choroby czy w wykrywaniu zmian patologicznych jak guzy itp. Przetwarzanie języka naturalnego (*ang. Natural Language Processing, NLP*) to dziedzina sztucznej

⁴ Zob. Robot da Vinci w Polsce – zastosowanie, dostępność, NFZ; <https://www.zwrotnikraka.pl/robot-da-vinci-w-polsce-zastosowanie/>; [19.11.2024].

inteligencji, która łączy ze sobą elementy programowania i lingwistyki, dzięki czemu umożliwia komputerom rozumienie, przetwarzanie i generowanie języka ludzkiego (tekstu lub mowy). NLP stosowane jest do analizy dokumentacji medycznej, co wspomaga szybsze diagnozowanie i pomaga dobrać odpowiedni sposób leczenia pacjenta. Systemy oparte na tej technologii mogą wspierać podejmowanie decyzji klinicznej, a także wspomagać interakcje pomiędzy pacjentem a systemem medycznym za pośrednictwem chatbotów medycznych i wirtualnych asystentów.

W wielu miejscach na świecie przeprowadzane są liczne badania i eksperymenty, które podkreślają ogromne znaczenie sztucznej inteligencji w medycynie. Badacze z Karolinska Institutet opracowali system, który na podstawie przesiewowych badań mammograficznych ukazuje, które kobiety są szczególnie zagrożone pojawieniem się raka piersi w przyszłości i wymagają bardziej szczegółowych badań.

Na wspomnianych zdjęciach istnieje wiele czynników, a AI znajduje w nich różne wzorce i potrafi łączyć te czynniki wskazując prognozy postępowania choroby. Badanie przeprowadzono w szpitalu Capio Sankt Göran w Sztokholmie w Szwecji. Tradycyjnie każdy obraz mammograficzny podlega analizie przeprowadzanej przez dwóch radiologów. We wspomnianym badaniu taki sposób oceny został porównany do analizy dokonywanej przez jednego radiologa, ale ze wsparciem AI. Populację badawczą stanowiły kobiety w wieku od 40 do 74 lat regularnie uczestniczące w badaniach mammograficznych. Niniejsze badanie dowiodło, że w przypadku jednego radiologa wspieranego przez sztuczną inteligencję w procesie odczytu mammografii przesiewowej zwiększyły się wskaźniki wykrywania raka piersi o 4% i zmniejszył się o połowę czas odczytu obrazów. Dodatkowo w przypadku jednego radiologa działającego wspólnie z AI było o 6% mniej wyników fałszywie pozytywnych niż przy pracy dwóch radiologów. Główny autor badania - Fredrik Strand stwierdził, że „nasze badanie pokazuje, że sztuczna inteligencja jest gotowa do kontrolowanej implementacji w przesiewowych badaniach mammograficznych. Jednak należy wybrać system AI, który został odpowiednio przetestowany na obrazach uzyskanych z tego samego rodzaju sprzętu do mammografii oraz zapewnić ciągle monitorowanie systemu po wprowadzeniu go do kliniki. W dłuższej perspektywie AI ma potencjał, aby przejąć większość pracy związanych z przesiewowymi badaniami mammograficznymi”. Jak podkreśliła autorka pracy opublikowanej w piśmie „The Lancet Digital Health”- Karin Dembrower „sztuczna inteligencja i ludzie postrzegają obrazy nieco inaczej. Tworzy to synergę, która zwiększa nasze szanse na wykrycie raka”⁵. W innym badaniu wykorzystano możliwości sztucznych sieci neuronowych do segmentacji i klasyfikacji obrazów pochodzących z optycznej tomografii koherencyjnej (ang. *optical coherence tomography*, OCT). Sieć segmentacyjna umożliwia wydzielenie 15 różnych cech morfologicznych siatkówki oraz artefaktów rejestrowanych podczas zbierania

⁵ K. Dembrower i in., *Sztuczna inteligencja w wykrywaniu raka piersi w mammografii przesiewowej w Szwecji: prospektywne, populacyjne, badanie typu non-inferiority z udziałem czytających w parach*, „The Lancet Digital Health” tom 5, nr 10, s. 703 – 711.

danych. Wyniki tej sieci są następnie przekazywane do sieci klasyfikacyjnej, której zadaniem jest przeprowadzenie decyzyjnego triage'u (badanie pilne, umiarkowanie pilne, rutynowe oraz obserwacja) oraz rozpoznanie 10 różnych schorzeń, takich jak neowaskularyzacja naczyńówki, obrzęk płamki bez CNV, druzy, zanik geograficzny, błona nasiatkówkowa, trakcja szklistkowo-siatkówkowa, pełnościenny otwór w płamce, otwór w płamce o niepełnej grubości, centralna retinopatia surowicza oraz stan „prawidłowy”. W ramach tego podejścia osiągnięte wyniki klasyfikacyjne uznano za porównywalne do rezultatów uzyskiwanych przez specjalistów. W przyszłości system ten mógłby być wykorzystywany do selekcji pacjentów poza szpitalem, a nawet poza gabinetem lekarskim, zwłaszcza że systemy OCT są coraz częściej używane przez optometrystów⁶.

Jednym z obszarów, w którym prowadzone są obecnie testy są teleoperacje. W październiku 2023 roku naukowcy z Singapuru i Japonii pomyślnie sprawdzili możliwość przeprowadzenia takiego zabiegu. Zespół chirurgów z Singapuru wykonał zdalnie operację gastrektomii na symulowanym żołądku znajdującym się w zrobotyzowanej sali operacyjnej w Japonii. Był to pierwszy tego rodzaju zabieg w obu krajach, który został przeprowadzony dzięki systemowi robotycznemu zdolnemu do precyzyjnego odwzorowania ruchów chirurgów z Singapuru, przesyłanych za pomocą dedykowanej międzynarodowej sieci światłowodowej. W tym celu wykorzystano system chirurgiczny hinotori, opracowany przez japońską firmę Mediaroid zajmującą się technologiami medycznymi, będący pierwszym rodzimym systemem chirurgii robotycznej w Japonii. Mimo odległości przekraczającej 5000 km operacja zakończyła się sukcesem. Rozwój takiej technologii otwiera nowe możliwości na przyszłość, choć rodzi również istotne pytania natury etycznej i prawnej⁷. Z kolei 11 września 2023 roku w Trzeciej Katedrze Chirurgii Czwartego Szpitala Uniwersytetu Medycznego w Hebei zdalnie przeprowadzono radykalną dystalną gastrektomię u 51-letniego pacjenta z nowotworem żołądka. Zabieg został wykonany przy użyciu robota obsługującego sieć 5G i zakończył się powodzeniem, przy minimalnym opóźnieniu. Ten przypadek potwierdził wstępnie, że zdalne przeprowadzanie radykalnej dystalnej gastrektomii z wykorzystaniem robota 5G jest możliwe i bezpieczne w przypadku pacjentów z rakiem żołądka. Badanie to stanowi solidną podstawę do dalszych prac naukowych⁸. Innym przykładem wykorzystania AI w medycynie są chatboty oparte na LLM (large language models), które mogą

⁶ D. Ting, L. Pasquale, L. Peng, J. Campbell, A. Lee, R. Raman, G. Tan, L. Schmetterer, P. Keane, T.Y. Wong, *Artificial intelligence and deep learning*, „British Journal of Ophthalmology” 2018, 103(2), s. 167–175; zob. też A. Woźniacka, S. Patrzyk, *Sztuczna inteligencja w medycynie*, Wydawnictwo Uniwersytetu Medycznego w Łodzi, Łódź 2022, s.20-22.

⁷ A. Ang, *Is Asia-Pacific ready for robotic telesurgery? A pioneering trial between Singapore and Japan has proven that it is feasible*; <https://www.healthcareitnews.com/news/asia/asia-pacific-ready-robotic-telesurgery> [29.11.2024].

⁸ H. Guo, Y. Tian, J. Shi, P. Yang, J. Yang, P. Ding, X. Zhao, Z. Zhang, Q. Zhao, *Pioneering case: Robot-assisted remote radical distal gastrectomy for gastric cancer based on 5G communication technology*, „Intelligent Surgery” 2024, tom 7, s. 22-26, <https://www.sciencedirect.com/science/article/pii/S2666676624000048?via%3Dihub> [29.11.2024].

być ogromnym wsparciem w diagnostyce. Jako przykład można przytoczyć przypadek 4-letniego Alexa, który pokazuje jak zastosowanie sztucznej inteligencji może przynieść realne korzyści w diagnostyce medycznej. Problemy zdrowotne chłopca zaczęły się w trakcie pandemii COVID-19, kiedy rodzice zauważyli u niego zmiany w zachowaniu - stał się rozdrażniony, często zgrzytał zębami, gryzł różne przedmioty, a później zaczął także powłóczyć nogą. Mimo licznych wizyt u specjalistów - dentysty, ortodonta, neurologa i pediatry - nikt nie potrafił znaleźć przyczyny jego stanu. Lekarze badali go jedynie pod kątem swoich specjalizacji, co utrudniło postawienie właściwej diagnozy. W ciągu trzech lat rodzice odwiedzili 17 różnych specjalistów bez rezultatów. Matka chłopca postanowiła skorzystać z ChatGPT, wpisując wszystkie objawy syna. Sztuczna inteligencja zasugerowała, że przyczyną może być zespół spętanego rdzenia kręgowego, czyli rzadkie schorzenie, które ogranicza ruchomość rdzenia kręgowego przez nieprawidłowe połączenia tkankowe. Po konsultacji z neurochirurgiem diagnoza została potwierdzona na podstawie wyników rezonansu magnetycznego. Choroba ta, choć często diagnozowana po urodzeniu, w przypadku Alexa ujawniła się dopiero w późniejszym etapie rozwoju. Kluczową rolę w znalezieniu diagnozy odegrała AI, która połączyła wszystkie objawy w całość, czego nie udało się wcześniej osiągnąć lekarzom⁹. Takich przypadków, w których wykorzystanie systemów AI w medycynie przyniosło namacalne korzyści jest mnóstwo. Wszystkie podane przykłady ukazują jak ważną rolę może odgrywać sztuczna inteligencja w diagnostyce i leczeniu, mimo towarzyszącym temu obaw i kontrowersji.

Aspekty prawne wykorzystywania sztucznej inteligencji w medycynie

Wykorzystanie sztucznej inteligencji w medycynie jest ściśle powiązane z przestrzeganiem różnych regulacji prawnych, zarówno na poziomie krajowym, jak i unijnym. Obejmują one przepisy dotyczące ochrony danych osobowych, bezpieczeństwa wyrobów medycznych czy standardów. W krajach Unii Europejskiej, w tym w Polsce, jednym z kluczowych aktów prawnych jest rozporządzenie o ochronie danych osobowych, czyli RODO (Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2016/679 z dnia 27 kwietnia 2016 r. w sprawie ochrony osób fizycznych w związku z przetwarzaniem danych osobowych i w sprawie swobodnego przepływu takich danych oraz uchylenia dyrektywy 95/46/WE)¹⁰. RODO wymaga, aby dane zbierane były w sposób transparentny i zrozumiały, co w kontekście medycyny ma szczególne znaczenie. AI w medycynie może przetwarzać dane osobowe pacjentów jedynie na podstawie wyraźnej i dobrowolnej zgody pacjenta lub innej podstawy prawnej, np. w celu ochrony życia lub zdrowia (art.

⁹ Zob. Przez trzy lata nie pomogło mu nawet 17 lekarzy. To chatGPT uwolnił chłopca od cierpienia; <https://www.poradnikzdrowie.pl/aktualnosci/przez-trzy-lata-nie-pomoglo-mu-nawet-17-lekarzy-to-chatgpt-uwolnił-chłopca-od-cierpienia-aa-PRa4-sJsy-HDgM.html> [29.11.2024].

¹⁰ T.j. Dz. U. UE. L. z 2016 r. Nr 119, str. 1 z późn. zm.

6 i art. 9 RODO). Systemy AI muszą przetwarzać tylko te dane osobowe, które są niezbędne do osiągnięcia określonych celów medycznych (art. 5 RODO). Przetwarzanie nadmiarowych danych może naruszać przepisy RODO. Pacjent ma prawo do uzyskania informacji o tym, jakie dane osobowe są przetwarzane przez system AI, w jakim celu oraz przez jaki czas będą one przechowywane (art. 15 RODO). Podmiot stosujący algorytmy lub zautomatyzowane decyzje powinien, na podstawie art. 22, ujawnić kluczowe zasady podejmowania tych decyzji oraz ich znaczenie i potencjalne skutki dla osoby, której dane dotyczą. Zgodnie z art. 17 RODO pacjent może żądać usunięcia swoich danych osobowych z systemu AI, pod warunkiem że nie istnieje inna podstawa prawna do ich dalszego przetwarzania (np. obowiązek archiwizacji dokumentacji medycznej). Wprowadzenie systemu AI w medycynie, szczególnie jeśli wiąże się z przetwarzaniem dużej ilości danych wrażliwych, wymaga przeprowadzenia oceny skutków dla ochrony danych (art. 35 RODO). Systemy sztucznej inteligencji wykorzystywane do celów diagnostycznych i terapeutycznych są często klasyfikowane jako wyroby medyczne, co oznacza, że podlegają regulacjom unijnym takim jak Rozporządzenie Parlamentu Europejskiego i Rady (UE) 2017/745 z dnia 5 kwietnia 2017 r. w sprawie wyrobów medycznych¹¹. Dokument ten wprowadza szczegółowe wymagania dotyczące bezpieczeństwa, jakości oraz certyfikacji systemów AI. Systemy AI wykorzystywane do diagnozy, monitorowania lub leczenia chorób mogą być klasyfikowane jako wyroby medyczne i muszą spełniać odpowiednie wymagania techniczne. Przed wprowadzeniem sztucznej inteligencji do obrotu konieczne jest przeprowadzenie procedury oceny zgodności, w ramach której system AI zostanie poddany audytowi pod kątem bezpieczeństwa i skuteczności. Systemy powinny być ciągle monitorowane oraz zapewniony powinien być odpowiedni poziom ich kontroli w praktyce klinicznej. Kolejnym aktem jest Dyrektywa NIS2 (Dyrektywa Parlamentu Europejskiego i Rady (UE) 2022/2555 z dnia 14 grudnia 2022 r. w sprawie środków na rzecz wysokiego wspólnego poziomu cyberbezpieczeństwa na terytorium Unii, zmieniająca rozporządzenie (UE) nr 910/2014 i dyrektywę (UE) 2018/1972 oraz uchylająca dyrektywę (UE) 2016/1148)¹². Dotyczy ona bezpieczeństwa sieci i systemów informacyjnych, w tym systemów AI wykorzystywanych w medycynie. Placówki medyczne wykorzystujące AI muszą wdrożyć odpowiednie środki techniczne i organizacyjne w celu zabezpieczenia systemów przed atakami cybernetycznymi. W przypadku naruszenia bezpieczeństwa danych pacjentów lub innych krytycznych incydentów związanych z AI, placówka medyczna jest zobowiązana do zgłoszenia tego odpowiednim organom nadzoru. 12 lipca 2024 r. w Dzienniku Urzędowym Unii Europejskiej została opublikowana nowa unijna regulacja dotycząca sztucznej inteligencji, mianowicie Akt w sprawie sztucznej inteligencji (AI Act)¹³. Celem AI Act jest zapewnienie, że rozwój i stosowanie sztucznej inteligencji w Europie

¹¹ T.j. Dz. U. UE. L. z 2017 r. Nr 117, str. 1 z późn. zm.

¹² T.j. Dz. U. UE. L. z 2022 r. Nr 333, str. 80.

¹³ T.j. Dz.U. L, 2024/1689 z 12.7.2024.

będzie odbywać się w sposób etyczny, przejrzysty i zgodny z prawami człowieka. W ramach tego aktu AI stosowana w medycynie będzie uznawana za technologię wysokiego ryzyka i objęta szeregiem rygorystycznych wymogów. AI Act oczekuje od dostawców systemów AI przeprowadzania audytów oraz ocen zgodności, aby zapewnić bezpieczeństwo pacjentów i skuteczność stosowanych technologii. AI wymaga, aby każda technologia przeszła audyt techniczny i uzyskała certyfikat potwierdzający jej bezpieczeństwo oraz skuteczność. Konieczne jest monitorowanie działania systemu oraz aktualizacja w odpowiedzi na nowe zagrożenia. Zarówno lekarze, jak i pacjenci muszą być informowani o tym, że AI wspiera diagnozę lub leczenie, oraz mieć dostęp do wyjaśnień dotyczących sposobu podejmowania decyzji. Ponadto AI Act wskazuje, że systemy muszą być zgodne z przepisami o ochronie danych, zapewniając prywatność pacjentów.

Granice etyczne stosowania sztucznej inteligencji w medycynie

Rozwój sztucznej inteligencji w medycynie rodzi istotne dylematy etyczne. W 2019 r. Organizacja Współpracy Gospodarczej i Rozwoju (OECD) opublikowała zalecenia dotyczące sztucznej inteligencji w dokumencie pt. *„Recommendation of the Council on Artificial Intelligence”*. OECD wskazuje, że sztuczna inteligencja ma potencjał, aby poprawić dobrobyt i dobrostan ludzi, przyczynić się do pozytywnej, zrównoważonej globalnej aktywności gospodarczej, zwiększyć innowacyjność i produktywność oraz pomóc w reagowaniu na kluczowe globalne wyzwania, podkreślając że jednocześnie transformacje te mogą mieć różne skutki zwłaszcza w odniesieniu do zmian gospodarczych, przekształceń na rynku pracy, nierówności oraz skutków dla demokracji i praw człowieka, a także prywatności, ochrony danych i bezpieczeństwa cyfrowego. OECD posługuje się określeniem „interesariuszy”, co obejmuje wszystkie organizacje i osoby zaangażowane w systemy AI lub dotknięte nimi, bezpośrednio lub pośrednio, a podmioty AI stanowią podzbiór interesariuszy. Podkreślono, że podmioty AI powinny szanować praworządność, prawa człowieka, wartości demokratyczne i zorientowane na człowieka w całym cyklu życia systemu. Ponadto powinny zobowiązać się do przejrzystości i odpowiedzialnego ujawniania informacji dotyczących systemów, w tym takich informacji, które umożliwią kwestionowanie wyników osobom, na które negatywnie wpływa system sztucznej inteligencji. Kolejnym zaleceniem jest wdrożenie w stosownych przypadkach mechanizmów zapewniających możliwość bezpiecznego wyłączenia, naprawienia, bądź wycofania z eksploatacji (według potrzeb) systemów sztucznej inteligencji, jeśli stwarzają ryzyko spowodowania nadmiernych szkód lub wykazują niepożądane zachowanie. Rada stwierdziła również, że podmioty działające w obszarze sztucznej inteligencji powinny ponosić odpowiedzialność za prawidłowe funkcjonowanie systemów sztucznej inteligencji¹⁴. Innym ważnym dokumentem

¹⁴ Zalecenie Rady w sprawie sztucznej inteligencji; <https://legalinstruments.oecd.org/en/instruments/>

wydanym w 2019 r. są Wytyczne Etyczne dla Godnej Zaufania AI (Ethics Guidelines for Trustworthy AI) opublikowane przez Komisję Europejską, a przygotowane przez grupę ekspertów wysokiego szczebla ds. Sztucznej Inteligencji. Kierowali się oni podejściem Human-Centric AI, w którym człowiek stoi w centrum refleksji etycznej. Wskazane zostały zasady etyczne, których należy przestrzegać przy opracowywaniu, wdrażaniu i wykorzystywaniu systemów AI. Jedną z nich jest stosowanie AI zgodnie z takimi dewizami jak autonomia człowieka, zapobieganie szkodom, sprawiedliwość i możliwość wyjaśnienia, zważając jednocześnie że korzystanie z tych systemów może wiązać się z określonym ryzykiem i wywoływać negatywne skutki, również takie które mogą okazać się trudne do przewidzenia, zidentyfikowania lub zmierzenia (np. wpływ na demokrację, praworządność i sprawiedliwość dystrybutywną lub wpływ na sam umysł ludzki). Z tego powodu w stosownych przypadkach należałoby przyjąć odpowiednie środki ograniczające to ryzyko, które będą proporcjonalne do jego skali. Bazując na tym eksperci sformułowali siedem wymogów, które AI powinna spełniać. Są to: przewodnia i nadzorczą rola człowieka, solidność techniczna i bezpieczeństwo, ochrona prywatności i zarządzanie danymi, przejrzystość, różnorodność, niedyskryminacja i sprawiedliwość, dobrostan społeczny i środowiskowy oraz odpowiedzialność. W dokumencie wskazano na potrzebę przyjęcia listy kontrolnej oceny godnej zaufania sztucznej inteligencji, która ma na celu zapewnienie odpowiedniego stosowania ww. wymogów. Należy pamiętać o tym, że taka lista kontrolna oceny nigdy nie będzie miała wyczerpującego charakteru¹⁵. W czerwcu 2021r. WHO opublikowało pierwszy globalny raport na temat sztucznej inteligencji (AI) w ochronie zdrowia i wskazało sześć wytycznych dotyczących jej projektowania i stosowania. W wyniku dwuletnich konsultacji prowadzonych przez powołany przez WHO panel międzynarodowych ekspertów powstał raport pt. *„Etyka i zarządzanie sztuczną inteligencją w ochronie zdrowia”*. Wskazano w nim, że co prawda wykorzystywanie AI w ochronie zdrowia jest bardzo obiecujące i przyczyniło się do ważnych postępów w takich dziedzinach medycyny jak radiologia czy profilaktyka, ale niesie to za sobą jednocześnie wiele transnarodowych obaw etycznych, prawnych handlowych i społecznych. Te wyzwania powinny zostać jak najszybciej rozwiązane poprzez zastosowanie odpowiednich regulacji. Nieuregulowane wykorzystanie sztucznej inteligencji może skutkować podporządkowaniem praw pacjentów i społeczności potężnym interesom komercyjnym firm technologicznych lub interesom rządów w zakresie nadzoru i kontroli społecznej. Kolejną przeszkodą, która może wpłynąć na wykorzystanie AI jest tzw. „podział cyfrowy”, który odnosi się do nierównomiernego podziału dostępu do technologii informacyjnych i komunikacyjnych pomiędzy krajami. W raporcie zwrócono uwagę także na problem systemów, które szkolone są głównie w oparciu o dane zbierane od osób z krajów o wysokich dochodach

oecd-legal-0449 [27.11.2024].

¹⁵ Wytyczne Etyczne dla Godnej Zaufania Sztucznej Inteligencji, file:///C:/Users/48693/Downloads/ethics_guidelines_for_trustworthy_ai-pl_880095D1-02DC-2AB8-EF31B15CEBBECA79_60436.pdf [28.11.2024].

w wyniku czego mogą nie działać prawidłowo w przypadku osób z krajów o niskich i średnich dochodach. Istnieje ryzyko, że bogatsze kraje i regiony będą miały większe możliwości wdrażania zaawansowanych technologii, co pogłębi przepaść między różnymi grupami społecznymi. Może to skutkować tym, że wykorzystanie AI będzie utrzymywać istniejące nierówności. Jeśli dane treningowe nie uwzględniają różnorodności populacji (np. płci, wieku, grup etnicznych), decyzje AI mogą dyskryminować mniejszości. Przykładem może być niższa skuteczność algorytmów diagnostycznych dla ciemniejszej karnacji skóry. Systemy AI powinny być zaprojektowane w taki sposób, aby odzwierciedlały różnorodność środowisk społeczno-ekonomicznych i opieki zdrowotnej. Konieczne w związku z tym są szkolenia dla ludzi zatrudnionych w branży zdrowotnej na temat zasad etycznych, praktycznych aspektów korzystania z tej technologii, a także potencjalnych korzyści i zagrożeń z tego płynących oraz praw pacjenta związanych z prywatnością danych medycznych. Dane zdrowotne są jednymi z najbardziej wrażliwych informacji osobistych. Ich wykorzystywanie w systemach AI niesie ryzyko naruszenia prywatności pacjentów np. w wyniku wycieku danych lub ich nieautoryzowanego wykorzystania. Ponadto WHO podkreśla, że AI powinna pełnić funkcję narzędzia wspierającego decyzje, a nie je zastępować. Autonomia pacjenta wymaga, aby decyzje medyczne były podejmowane w oparciu o jego zgodę i zrozumienie procesu, co może być trudne, gdy technologia jest nieprzejrzysta. Zaawansowane algorytmy, zwłaszcza te oparte na głębokim uczeniu maszynowym, działają jak „czarne skrzynki”, których procesy decyzyjne są trudne do zrozumienia nawet dla ekspertów. Brak transparentności budzi obawy o ich etyczne zastosowanie i zaufanie użytkowników. WHO zwraca uwagę, że algorytmy medyczne muszą być zrozumiałe dla personelu medycznego, aby mogły być stosowane w praktyce klinicznej.

W celu ograniczenia ryzyka i zmaksymalizowania możliwości wynikających z wykorzystania sztucznej inteligencji w ochronie zdrowia, WHO zebrała najważniejsze zasady, stanowiące podstawę regulacji i zarządzania sztuczną inteligencją. Zasady te stanowią podstawę dla rządów, twórców technologii, firm, społeczeństwa obywatelskiego i organizacji międzyrządowych do przyjęcia etycznego podejścia do właściwego wykorzystania sztucznej inteligencji w dziedzinie zdrowia. Pierwszą z zasad jest ochrona autonomii człowieka, co oznacza że ludzie powinni zachować kontrolę nad systemami opieki zdrowotnej i decyzjami medycznymi, a więc ostateczna decyzja medyczna zawsze powinna należeć do członka personelu medycznego. AI muszą działać w zgodzie z przepisami dotyczącymi ochrony danych (np. RODO), aby zapewnić pacjentom poczucie bezpieczeństwa. Pacjenci muszą mieć prawo do wyrażenia świadomej zgody na stosowanie AI w ich leczeniu czy diagnozie. Kolejną zasadą jest promowanie ludzkiego dobrostanu i bezpieczeństwa oraz interesu publicznego, przy czym systemy AI powinny być gruntownie testowane pod kątem bezpieczeństwa i skuteczności, aby uniknąć niezamierzonych skutków ubocznych. Technologia ta powinna być rozwijana w sposób zrównoważony, z uwzględnieniem długoterminowych korzyści, jednocześnie starając się minimalizować negatywny

wpływ na środowisko. Mając na uwadze bezpieczeństwo pacjentów regularne testy i walidacja systemów powinny być prowadzone również po ich wdrożeniu, aby upewnić się że ich skuteczność nie ulega zmianie w czasie. WHO wskazało również na potrzebę transparentności oraz możliwości wyjaśnienia, co oznacza że przed zaprojektowaniem lub wdrożeniem technologii AI powinny zostać opublikowane lub udokumentowane dostępne powszechnie i wystarczające informacje na temat tych systemów. Algorytmy AI muszą działać w sposób przejrzysty i zrozumiały dla użytkowników, co pomoże w budowaniu zaufania i umożliwi audyt. Czwartą regułą jest wspieranie odpowiedzialności i rozliczalności. Mimo że technologie AI wykonują określone zadania, to interesariusze ponoszą odpowiedzialność za zapewnienie, że są one wykorzystywane w odpowiednich warunkach i przez odpowiednio przeszkolone osoby. W razie błędów lub szkód wywołanych w związku z korzystaniem z AI muszą istnieć jasne mechanizmy identyfikacji winy oraz dochodzenia roszczeń dla osób, na które negatywnie wpłynęły decyzje oparte na algorytmach. Jako kolejny cel WHO wzięło zapewnienie inkluzywności i równości, co wymaga aby AI w obszarze ochrony zdrowia była projektowana w sposób zachęcający do możliwie najszerzego sprawiedliwego korzystania i dostępu, niezależnie od wieku, płci, dochodu, rasy, pochodzenia etnicznego, orientacji seksualnej, sprawności czy też innych cech chronionych na mocy kodeksów praw człowieka. Należy zapobiegać nierównościom wynikającym z braku dostępu do technologii lub z uprzedzeń. Ostatnią zasadą zaproponowaną przez WHO jest promowanie AI, które jest responsywne i zrównoważone. AI opiera się na danych zdrowotnych które powinny być gromadzone, przechowywane i przetwarzane z poszanowaniem prywatności pacjentów, zgodnie z najwyższymi standardami bezpieczeństwa i etyki. Ponadto systemy powinny być na bieżąco projektowane w taki sposób, aby minimalizować negatywne konsekwencje zarówno wobec pacjentów, jak i w odniesieniu do utraty miejsc pracy z powodu korzystania z systemów zautomatyzowanych, gdyż w sektorach, gdzie powtarzalne zadania mogą być łatwo zautomatyzowane może dojść do redukcji kadr¹⁶. Jak podkreślił dr Jeremy Farrar z WHO - „technologie generatywnej sztucznej inteligencji mają potencjał, aby opiekę zdrowotną czynić lepszą, ale tylko wtedy, gdy ci, którzy opracowują, regulują i wykorzystują te technologie, zidentyfikują i w pełni uwzględnią związane z nimi ryzyko”.

Aspekt etyczny w odniesieniu do sztucznej inteligencji powinien uwzględniać perspektywę etyki medycznej i bioetyki. Systemy sztucznej inteligencji będą na tyle etyczne na ile etyką będą się kierować ich twórcy, dlatego szczególną uwagę w wyznaczaniu etycznych ram powinno się skupiać właśnie na nich. Twórcy systemów

¹⁶ Wytyczne WHO: Etyka i zarządzanie sztuczną inteligencją w ochronie zdrowia, World Health Organization 2021, [chrome-extension://efaidnbmnnnibpcajpegglefindmkaj/https://iris.who.int/bitstream/handle/10665/341996/9789240029200-eng.pdf?sequence=1](https://iris.who.int/bitstream/handle/10665/341996/9789240029200-eng.pdf?sequence=1) [26.11.2024]; zob. też *WHO publikuje pierwszy globalny raport na temat sztucznej inteligencji (AI) w ochronie zdrowia i sześć wytycznych dotyczących jej projektowania i stosowania*; <https://www.who.int/news/item/28-06-2021-who-issues-first-global-report-on-ai-in-health-and-six-guiding-principles-for-its-design-and-use> [26.11.2024].

AI dla sektora zdrowia powinni brać pod uwagę uwarunkowania kształtujące etyczne ramy pracy lekarza, do których należą takie zasady jak dobro pacjenta, jego autonomia i sprawiedliwość społeczna. Istotne z punktu widzenia etyki jest przewidywanie konsekwencji i szybka reakcja, co może być osiągalne z pomocą nowych instytucjonalnych mechanizmów oceny AI, które będą charakteryzować się interdyscyplinarnością, refleksą, systematycznością i sprawczością. Ważnym zagadnieniem pozostaje problem przestrzegania tajemnicy zawodowej co przy systemach AI, które funkcjonują dzięki danym, jest poważnym wyzwaniem¹⁷.

Gromadzenie, przechowywanie i analiza dużych zbiorów danych medycznych mogą stwarzać ryzyko naruszenia prywatności pacjentów. Każda osoba ma prawo do ochrony swoich danych zdrowotnych przed dostępem osób nieupoważnionych. Dlatego tak istotne jest zagwarantowanie bezpieczeństwa informacji poprzez zastosowanie odpowiednich rozwiązań technologicznych oraz przestrzeganie aktualnych przepisów dotyczących ochrony danych osobowych. Konieczne jest również opracowanie jasnych procedur regulujących wykorzystanie sztucznej inteligencji w medycynie, które będą zgodne z obowiązującymi normami prawnymi, zasadami etyki i standardami bezpieczeństwa. Wprowadzenie takich procedur pozwoli placówkom medycznym korzystać z AI w sposób odpowiedzialny, bezpieczny i etyczny, przyczyniając się do podniesienia jakości opieki zdrowotnej. Ważnym elementem jest także organizowanie szkoleń dla personelu medycznego, które będą obejmować zarówno aspekty etyczne, jak i praktyczne związane z wykorzystaniem sztucznej inteligencji w codziennej pracy. Niezbędne jest również informowanie pacjentów o możliwościach, korzyściach i potencjalnych zagrożeniach wynikających z użycia AI, a także o ich prawach związanych z ochroną danych medycznych. Dodatkowo istotne jest wdrożenie systemów monitorowania działań związanych z zastosowaniem sztucznej inteligencji, aby śledzić decyzje podejmowane przez algorytmy, identyfikować ewentualne błędy i podejmować odpowiednie działania naprawcze. Nie mniej ważne jest wprowadzenie mechanizmów kontroli i oceny jakości działań, podejmowanych z wykorzystaniem sztucznej inteligencji, w tym danych wykorzystywanych do uczenia maszynowego, co zapewni ich rzetelność, kompletność i spójność.

Odpowiedzialność za błędy popełnione z wykorzystaniem sztucznej inteligencji

Odpowiedzialność za błędy wyrządzone przez sztuczną inteligencję to złożone zagadnienie, które obejmuje kwestie prawne, etyczne i techniczne. Już na wstępie należy zaznaczyć, że w przeciwieństwie do człowieka, AI nie posiada osobowości prawnej ani zdolności do ponoszenia odpowiedzialności, dlatego odpowiedzialność za jej działania spada na ludzi lub organizacje, które ją stworzyły, wdrożyły lub

¹⁷ R. Sroka, *Wyzwania etyczne AI w sektorze zdrowia*, [w:] R. Siewiorek, M. Lach-Kuśnierz (red.), *Iloraz Sztucznej Inteligencji tom 4. Potencjał Sztucznej Inteligencji w Sektorze Ochrony Zdrowia*, Fundacja AI Law Tech, Warszawa 2021, s. 76-79.

nadzorowały. Należy uznać, że podmiotowość prawna nie może być przyznana bytom, które nie mają wolnej woli lub nie są rozumne. Systemy AI nie mają swojego własnego interesu prawnego w osiągnięciu określonego celu, cel ten zaś przyświeca ich twórcom, programatorom czy podmiotom, dla których dobra zostały stworzone. Należy stwierdzić, że prawidłowe jest (przynajmniej na tym etapie rozwoju) nieprzypisywanie sztucznej inteligencji zdolności prawnej ani zdolności do czynności prawnych. Zgodnie z art. 8 ustawy z dnia 23 kwietnia 1964 r. - Kodeks cywilny „każdy człowiek od chwili urodzenia ma zdolność prawną”¹⁸. Nie ma żadnych podstaw, bazując na tym przepisie, aby przyznać ową zdolność sztucznej inteligencji. Wynika z tego, że nie należy również przypisywać osobowości prawnej AI, co można rozumieć jako zdolność osób prawnych do bycia podmiotem praw oraz obowiązków, a także dokonywania we własnym imieniu czynności, które powodują skutki prawne. Rozważając jednak kwestie przyznania sztucznej inteligencji tych atrybutów pojawia się wiele pytań. M.in. kiedy miałyby nastąpić moment „urodzenia” czy „śmierci” AI? W jakim rejestrze widniała by sztuczna inteligencja - jako osoba fizyczna a może prawna? W jaki sposób system miałby ponosić odpowiedzialność?

W wielu krajach trwają prace nad regulacjami prawnymi dotyczącymi odpowiedzialności za działania AI. Brak odpowiednich regulacji może prowadzić do niejasności w kwestii odpowiedzialności prawnej, co może skutkować długotrwałymi i kosztownymi sporami sądowymi. Brak standardów może utrudniać współpracę międzynarodową i integrację systemów AI z innymi krajami, co może ograniczać możliwości rozwoju i wdrażania nowych technologii.

Wprowadzenie zharmonizowanych wytycznych etycznych jest niezbędne do prawidłowego wdrożenia AI w globalnym zdrowiu. Parlament Europejski w Rezolucji z dnia 16 lutego 2017 r. zawierającej zalecenia dla Komisji w sprawie przepisów prawa cywilnego dotyczących robotyki wskazał, iż należy w odpowiednich przypadkach zaktualizować i uzupełnić obecne unijne ramy prawne o główne zasady etyczne, które będą odzwierciedlać złożoność robotyki i jej liczne implikacje społeczne, medyczne i bioetyczne. Parlament zwrócił uwagę na potrzebę zdefiniowania minimalnych wymogów zawodowych, jakie musi posiadać chirurg, aby móc posługiwać się robotami chirurgicznymi i obsługiwać je. Ponadto trafnie podkreśla że zasadnicze znaczenie ma przestrzeganie zasady kontrolowanej autonomii robotów, zgodnie z którą ostateczną decyzję zawsze podejmuje lekarz. W rezolucji tej wskazano na potrzebę utworzenia obowiązkowego systemu ubezpieczeń, który mógłby opierać się na zobowiązaniu producenta do wykupienia ubezpieczenia w związku z wytwarzanymi przez siebie automatycznymi robotami. Straty w przypadkach, których ubezpieczenie nie obejmuje mogło by być pokryte ze specjalnie do tego utworzonego funduszu. Parlament podkreśla jednak, że wszelkie decyzje polityczne dotyczące zasad odpowiedzialności cywilnej mających zastosowanie do robotów i sztucznej inteligencji powinny być podejmowane po konsultacjach z podmiotami realizującymi ogólnoeuropejski projekt badawczo-

¹⁸ T.j. Dz. U. z 2024 r. poz. 1061 z późn. zm.

rozwojowy w zakresie robotyki i neuronauki, dlatego że naukowcy i eksperci są w stanie ocenić wszystkie powiązane rodzaje ryzyka i konsekwencje. Przyjęte rozwiązania powinny zapewnić ten sam poziom efektywności, przejrzystości i spójności w całej Unii Europejskiej¹⁹. Warte wspomnienia są w tym miejscu są również trzy rezolucje Parlamentu Europejskiego z dnia 20 października 2020 r. odnoszące się do sztucznej inteligencji. Nawiązują one do istotnych dla tego rodzaju systemów płaszczyzn tj. etyki (Rezolucja zawierająca zalecenia dla Komisji w sprawie ram aspektów etycznych sztucznej inteligencji, robotyki i powiązanych z nimi technologii²⁰), odpowiedzialności cywilnej za systemy AI (Rezolucja z zaleceniami dla Komisji w sprawie systemu odpowiedzialności cywilnej za sztuczną inteligencją²¹) oraz aspektów praw własności intelektualnej w stosunku do AI (Rezolucja Parlamentu Europejskiego w sprawie praw własności intelektualnej w dziedzinie rozwoju technologii sztucznej inteligencji)²². Ważnym postulatem Parlamentu Europejskiego w stosunku do wszystkich rezolucji jest zastrzeżenie, że przy stosowaniu AI za prawidłowe należy uznać kierowanie się zasadą „ograniczonego zaufania”, co oznaczanie pozostawianie robotom pełnej decyzyjności ani zbytniej swobody.²³ Kolejną ważną inicjatywą jest projekt dyrektywy w sprawie dostosowania przepisów dotyczących odpowiedzialności cywilnej pozaumownej do sztucznej inteligencji (Dyrektywa Parlamentu Europejskiego i Rady w sprawie dostosowania przepisów dotyczących pozaumownej odpowiedzialności cywilnej do sztucznej inteligencji, AILD)²⁴. Ma ona na celu wspieranie wdrożenia sztucznej inteligencji w sposób godny zaufania, aby w pełni wykorzystać jej potencjał na rynku wewnętrznym. Istotnym elementem jest zapewnienie poszkodowanym przez AI ochrony podobnej do tej, jaką mają osoby poszkodowane przez tradycyjne produkty. AILD uzupełnia i aktualizuje unijne ramy odpowiedzialności cywilnej, wprowadzając po raz pierwszy przepisy, które dotyczą szkód spowodowanych przez systemy sztucznej inteligencji. Głównym założeniem dyrektywy jest to, że obowiązujące przepisy krajowe dotyczące odpowiedzialności, w szczególności opartej na zasadzie winy, nie są dopasowane do rozpatrywania roszczeń z tytułu odpowiedzialności za szkody spowodowane przez produkty i usługi oparte na sztucznej inteligencji.

Zgodnie z przepisami krajowymi poszkodowani muszą udowodnić niezgodne

¹⁹ T.j. Dz. U. UE C 252/239, s. 239–257, 2015/2103(INL), <https://eur-lex.europa.eu/legal-content/PL/TXT/?uri=CELEX%3A52017IP0051> [29.11.2024].

²⁰ T.j. Dz. Urz. UE C 404 z 6.10.2021 r., s. 63–106, 2020/2012 (INL), https://www.europarl.europa.eu/doceo/document/TA-9-2020-0275_PL.html#title1 [29.11.2024].

²¹ T.j. Dz. Urz. UE C 404 z 6.10.2021 r., s. 107–128, 142020/2014(INL), https://eur-lex.europa.eu/legal-content/PL/TXT/?uri=OJ:JOC_2021_404_R_0006 [dostęp: 29.11.2024].

²² T.j. Dz. Urz. UE C 404 z 6.10.2021 r., s. 129–135, 152020/2015 (INI), https://www.europarl.europa.eu/doceo/do-cument/TA-9-2020-0277_PL.html [29.11.2024].

²³ K. Bączyk-Rozwadowska, *Odpowiedzialność cywilna za szkody wyrządzone w związku z zastosowaniem sztucznej inteligencji w medycynie*, „Przegląd Prawa Medycznego” 2021, nr 3–4(8), s. 8–9.

²⁴ COM(2022) 496 final, 2022/0303(COD), <https://eur-lex.europa.eu/legal-content/PL/TXT/?uri=CELEX%3A52022PC0496> [29.11.2024].

z prawem działanie lub zaniechanie osoby, która spowodowała szkodę, co przy charakterystycznych cechach AI, takich jak złożoność, autonomia i brak przejrzystości (tzw. efekt czarnej skrzynki) może być trudne i kosztowne. Dodatkowo, dochodzenie odszkodowania w takich przypadkach może wymagać dużych nakładów finansowych i wydłużać proces sądowy, co może zniechęcić poszkodowanych do podejmowania działań prawnych²⁵. Parlament Europejski w swojej rezolucji z 3 maja 2022 r. w sprawie sztucznej inteligencji w epoce cyfrowej podtrzymał te obawy i potwierdził potrzebę harmonizacji przepisów w zakresie odpowiedzialności za szkody spowodowane przez AI na poziomie Unii²⁶. Harmonizacja ta powinna dotyczyć przede wszystkim zasad dotyczących ciężaru dowodu w postępowaniach odszkodowawczych, co ma zapewnić większą pewność prawną i równe warunki działania na rynku dla produktów i usług opartych na AI. Nie będzie jednak obejmowała wszystkich aspektów odpowiedzialności cywilnej, które nadal będą regulowane przez prawo krajowe. Należy również zapewnić spójność regulacji z aktem w sprawie AI. Istotnym elementem jest również dostęp do informacji dotyczących systemów sztucznej inteligencji, zwłaszcza tych wysokiego ryzyka, które mogą być kluczowe w procesie dochodzenia odszkodowania. Ponadto w przypadku systemów sztucznej inteligencji wysokiego ryzyka akt w sprawie AI przewiduje szczególne wymogi dotyczące dokumentacji, informacji i rejestracji danych, ale nie zapewnia osobie poszkodowanej prawa dostępu do tych informacji. Należy w związku z tym ustanowić zasady ujawniania istotnych dowodów przez podmioty, które nimi dysponują, do celów ustalenia odpowiedzialności. Dyrektywa w sprawie odpowiedzialności za sztuczną inteligencję niejako łagodzi ciężar dowodu spoczywający na ofiarach dzięki wprowadzeniu „domniemania istnienia związku przyczynowego”. Jeśli ofiary są w stanie wykazać, że ktoś ponosi winę za niedopełnienie pewnego obowiązku, istotnego z punktu widzenia szkody, a związek przyczynowy z funkcjonowaniem sztucznej inteligencji jest racjonalnie prawdopodobny, sąd może założyć, że to niedopełnienie tego obowiązku było przyczyną szkody. Z drugiej strony osoba odpowiedzialna może obalić to domniemanie z racji że jest ono wzruszalne (na przykład udowadniając, że szkodę spowodowała inna przyczyna)²⁷. Dyrektywa nie narusza krajowych zasad dotyczących ustalania winy, ale wprowadza mechanizmy, które mają ułatwić dochodzenie roszczeń w sprawach związanych ze sztuczną inteligencją. Warto także wspomnieć w tym miejscu o tym, że Komisja Europejska oprócz projektu AILD opublikowała również projekt nowej dyrektywy w sprawie odpowiedzialności za produkty wadliwe (dyrektywa Parlamentu Europejskiego i Rady (UE) 2024/2853 z dnia 23 października 2024 r. w sprawie odpowiedzialności

²⁵ P. Dolniak, T. Kuźma, A. Ludwiński, K. Wasik, *1.4. Dyrektywa w sprawie odpowiedzialności za sztuczną inteligencję (AILD)*, [w:] P. Dolniak, T. Kuźma, A. Ludwiński, K. Wasik, *Sztuczna inteligencja w wymiarze sprawiedliwości. Między prawem a algorytmami*, Warszawa 2024, s. 131-133.

²⁶ T.j. Dz. U. UE. C.2022.465.65, 2020/2266(INI), https://www.europarl.europa.eu/doceo/document/TA-9-2022-0140_PL.html [29.11.2024].

²⁷ R. Bujalski, *Odpowiedzialność za sztuczną inteligencję [projekt UE]*, LEX/el. 2022.

za produkty wadliwe i uchylenia dyrektywy Rady 85/374/EWG, PLD)²⁸. PLD ma na celu zapewnienie, że w przypadku wadliwych systemów AI, które powodują fizyczne uszkodzenia, szkody w mieniu lub utratę danych, będzie możliwość dochodzenia odszkodowania od dostawcy systemu AI lub od każdego producenta, który integruje system AI z innym produktem. Warto podkreślić, że przez niemal 40 lat dyrektywa 85/374/EWG (dyrektywa Rady z dnia 25 lipca 1985 r. w sprawie zbliżenia przepisów ustawowych, wykonawczych i administracyjnych Państw Członkowskich dotyczących odpowiedzialności za produkty wadliwe) dotycząca odpowiedzialności za produkty zapewniała obywatelom ochronę, umożliwiając dochodzenie odszkodowań za szkody spowodowane przez wadliwe produkty²⁹. Jednakże, ponieważ została uchwalona w 1985 roku, nie uwzględnia nowoczesnych technologii cyfrowych, takich jak inteligentne urządzenia czy systemy oparte na sztucznej inteligencji. Z tego względu konieczne stało się zaktualizowanie przepisów, aby dostosować je do realiów ery cyfrowej, obejmujących oprogramowanie, AI oraz usługi cyfrowe³⁰.

Istnieje kilka reżimów odpowiedzialności, które mogą być istotne przy stosowaniu AI w ochronie zdrowia. Odpowiedzialność cywilna za wykorzystanie systemów AI w medycynie najczęściej spoczywa na użytkowniku, czyli lekarzu lub operatorze urządzenia. Zgodnie z rezolucją Parlamentu Europejskiego z dnia 16 lutego 2017 roku, odpowiedzialność ta powinna opierać się na winie. W polskim prawie cywilnym, choć brak regulacji dotyczących systemów autonomicznych, zastosowanie ma ogólna zasada wynikająca z art. 415 k.c. stanowiąca, że osoba która wyrządziła szkodę z własnej winy, jest zobowiązana do jej naprawienia. W kontekście AI oznacza to, że lekarz ponosi odpowiedzialność, jeśli szkoda wynikała z jego niedbałości i pozostaje w adekwatnym związku przyczynowym z jego działaniem (art. 361 k.c.). Przypisanie winy wymaga wykazania, że lekarz nie zachował standardów staranności, jakie obowiązują w jego specjalizacji, oraz że mógł zapobiec szkodzie, gdyby wykorzystał dostępne metody diagnostyczne i terapeutyczne. Jednak odpowiedzialność nie może być przypisana, jeśli szkoda wynikała z autonomicznych, nieprzewidywalnych decyzji AI, na które lekarz nie miał wpływu. Odpowiedzialność może również spoczywać na producencie systemu AI. Zgodnie z art. 449¹ k.c., producent odpowiada za szkody wyrządzone przez produkt niebezpieczny na zasadzie ryzyka. Oznacza to, że nawet bez winy musi zapewnić bezpieczeństwo użytkowania swojego produktu. Problemem jest jednak brak jednoznacznej definicji „produktu” w prawie unijnym. AI jako niematerialne oprogramowanie, może nie być uznane za produkt, chyba że stanowi element sprzętu fizycznego. W sytuacji, gdy szkoda wynika z działania systemu składającego się z komponentów pochodzących od różnych producentów, trudne może być określenie, kto jest odpowiedzialny. Dodatkowo, ustalenie związku

²⁸ T.j. Dz. U. UE.L.2024.2853.

²⁹ T.j. Dz. U. UE. L. z 1985 r. Nr 210, str. 29 z późn. zm.

³⁰ P. Dolniak, T. Kuźma, A. Ludwiński, K. Wasik, *I.4. Dyrektywa w sprawie ...*, op. cit., s.127-128.

przyczynowego między wadą systemu AI a wyrządzoną szkodą komplikuje wspomniane już wcześniej zjawisko tzw. „czarnej skrzynki”, gdy mechanizm działania AI jest nieprzejrzysty. Choć odpowiedzialność karna za stosowanie AI w medycynie nie została jeszcze precyzyjnie uregulowana, obecne przepisy dotyczące np. nieumyślnego spowodowania śmierci czy przeprowadzenia zabiegu bez zgody pacjenta mogą znaleźć zastosowanie. Możliwa jest również odpowiedzialność zawodowa, jeśli użycie lub zaniechanie użycia AI narusza standardy etyczne lub przepisy regulujące wykonywanie zawodu lekarza³¹. Podział odpowiedzialności za sztuczną inteligencję na zasadzie winy i ryzyka jest rozwiązaniem pożądanym, ponieważ systemy AI różnią się stopniem ryzyka, które generują. W zależności od poziomu autonomiczności takich systemów, odpowiedzialność operatorów lub producentów powinna być zróżnicowana. Jednak oparcie odpowiedzialności wyłącznie na autonomiczności AI ma swoje wady. Pierwszym wyzwaniem jest trudność w precyzyjnym określeniu, od jakiego momentu system AI staje się na tyle autonomiczny, by przypisać mu odpowiedzialność na zasadzie ryzyka. Wymagałoby to wprowadzenia pewnego stopnia uznaniowości. Kolejnym problemem są kwestie dowodowe - trudno jednoznacznie wykazać, że dany system osiągnął wysoki poziom autonomii. Rozwiązaniem mogłoby być domniemanie prawne, że każda AI jest autonomiczna, a ciężar obalenia tego domniemania spoczywałby na operatorze. Dodatkowo, istnieje ryzyko, że system pierwotnie nieautonomiczny, w miarę działania i „uczenia się”, stanie się autonomiczny bez ingerencji operatora. W związku z tym należałoby oceniać stopień autonomiczności AI w momencie powstania szkody, a nie podczas jej projektowania czy wprowadzania na rynek. W kontekście dochodzenia odszkodowań za szkody wyrządzone przez AI, kluczowe znaczenie ma właśnie stopień autonomiczności systemu, który powinien determinować odpowiedni reżim odpowiedzialności. Należy jednak pamiętać, że ostateczna odpowiedzialność za działanie AI zawsze spoczywa na człowieku, ponieważ to ludzkie decyzje i działania leżą u podstaw funkcjonowania nawet najbardziej zaawansowanych systemów sztucznej inteligencji³².

Podsumowanie

Sztuczna inteligencja niesie ogromny potencjał dla medycyny, pozwalając na bardziej precyzyjną diagnostykę, efektywne leczenie oraz lepsze zarządzanie opieką zdrowotną. Jednak jej wdrożenie wymaga rozważenia istotnych wyzwań prawnych, takich jak ochrona danych, odpowiedzialność za decyzje algorytmów czy regulacje dotyczące wyrobów medycznych. Dodatkowo pojawiają się kwestie etyczne, takie jak prywatność pacjenta, równość dostępu, autonomia w podejmowaniu decyzji

³¹ P. Kaźmierczyk (red.), M. Kupis, M. Maj, *Biała Księga AI w Praktyce Klinicznej. Stosowanie sztucznej inteligencji przy udzielaniu świadczeń zdrowotnych*, wZdrowiu, Warszawa 2022, s. 64-73.

³² P. Staszczuk, *Czy unijna regulacja odpowiedzialności cywilnej za sztuczną inteligencję jest potrzebna?*, „Europejski Przegląd Sądowy” 2022, nr 6, s. 24-30.

czy ryzyko uprzedzeń. Kluczowe dla skutecznego i bezpiecznego wdrożenia SI w medycynie będzie stworzenie odpowiednich ram prawnych i standardów etycznych, które zapewnią ochronę pacjentów i umożliwią pełne wykorzystanie możliwości, jakie oferuje ta technologia. Sztuczna Inteligencja, tak jak każda technologia, nie jest nieomylna i powinna być nadzorowana przez człowieka. Jest systemem komputerowym opierającym się o dane statystyczne. Sztuczna Inteligencja, jak każde inne narzędzie, sama w sobie nie jest ani dobra ani zła. Dopiero jej użycie świadczy o motywacji użytkownika. Pozostawiona sama sobie nie jest w stanie nikomu wyrządzić znaczącej szkody. To sposób jej zastosowania może pociągać za sobą różne reperkusje³³. Zdrowie i życie ludzkie są najwyższą wartością nie tylko indywidualną, ale i społeczną, czego konsekwencją jest nakaz stosowania najlepszych metod leczenia i najlepszych środków medycznych i technicznych, które temu służą. Należy pamiętać, że to lekarze muszą podejmować ostateczne decyzje diagnostyczne i terapeutyczne, traktując AI jako narzędzie wspierające, a nie zastępujące ich wiedzę. W kontekście regulacji prawnych niezbędne jest wprowadzenie jasnych przepisów dotyczących odpowiedzialności za błędy, które uwzględnią specyfikę autonomicznych systemów AI. Konieczne jest również opracowanie międzynarodowych standardów walidacji, testowania i monitorowania systemów AI, aby zapewnić ich bezpieczeństwo i skuteczność w różnych środowiskach medycznych. Równie ważne jest stworzenie mechanizmów ochrony danych osobowych oraz zapewnienie audytowalności i wyjaśnialności decyzji podejmowanych przez algorytmy. Długoterminowym celem powinno być również zwiększenie dostępności zaawansowanych technologii medycznych opartych na AI, tak aby korzyści płynące z ich stosowania były dostępne dla wszystkich pacjentów, niezależnie od ich statusu społecznego czy miejsca zamieszkania. Rozwój sztucznej inteligencji w medycynie oferuje ogromne możliwości, ale jednocześnie wymaga precyzyjnych ram prawnych i etycznych, które zagwarantują bezpieczeństwo pacjentów, ochronę ich danych oraz transparentność działania systemów. Tylko takie podejście pozwoli w pełni wykorzystać potencjał SI w służbie zdrowia, jednocześnie minimalizując związane z nią ryzyko.

Piśmiennictwo

- Bączyk-Rozwadowska K., *Odpowiedzialność cywilna za szkody wyrządzone w związku z zastosowaniem sztucznej inteligencji w medycynie*, „Przegląd Prawa Medycznego” 2021, nr 3-4(8).
- Bujalski R., *Odpowiedzialność za sztuczną inteligencję [projekt UE]*, LEX/el. 2022.
- Dembrower K. i in., *Sztuczna inteligencja w wykrywaniu raka piersi w*

³³ K. Karski, *Wybrane zastosowania sztucznej inteligencji w medycynie*, „Management and Quality - Zarządzanie I Jakość” 2022, tom 4 nr 4, s. 209-212.

mammografii przesiewowej w Szwecji: prospektywne, populacyjne, badanie typu non-inferiority z udziałem czytających w parach, „The Lancet Digital Health” tom 5, numer 10.

- Dolniak P., Kuźma T., Ludwiński A., Wasik K., 1.4. *Dyrektywa w sprawie odpowiedzialności za sztuczną inteligencję (AILD)*, [w:] Dolniak P., Kuźma T., Ludwiński A., Wasik K., *Sztuczna inteligencja w wymiarze sprawiedliwości. Między prawem a algorytmami*, Warszawa 2024.
- Guo H., Tian Y., Shi J., Yang P., Yang J., Ding P., Zhao X., Zhang Z., Zhao Q., *Pioneering case: Robot-assisted remote radical distal gastrectomy for gastric cancer based on 5G communication technology*, „Intelligent Surgery” 2024, tom 7.
- Karski K., *Wybrane zastosowania sztucznej inteligencji w medycynie*, „Management and Quality – Zarządzanie I Jakość” 2022, tom 4, nr 4.
- Kaźmierczyk P. (red.), Kupis M., Maj M., *Biała Księga AI w Praktyce Klinicznej. Stosowanie sztucznej inteligencji przy udzielaniu świadczeń zdrowotnych*, w *Zdrowiu*, Warszawa 2022.
- Sroka R., *Wyzwania etyczne AI w sektorze zdrowia*, [w:] Siewiorek R., Lach-Kuśnierz M., (red.) *Iloraz Sztucznej Inteligencji* tom 4. *Potencjał Sztucznej Inteligencji w Sektorze Ochrony Zdrowia*, Fundacja AI Law Tech, Warszawa 2021.
- Staszczuk P., *Czy unijna regulacja odpowiedzialności cywilnej za sztuczną inteligencję jest potrzebna?*, „Europejski Przegląd Sądowy” 2022, nr 6.
- Ting D., Pasquale L., Peng L., Campbell J., Lee A., Raman R., Tan G., Schmetterer L., Keane P., Wong T. Y., *Artificial intelligence and deep learning*, „British Journal of Ophthalmology” 2018, 103(2)
- Woźniacka A., Patrzyk S., *Sztuczna inteligencja w medycynie*, Wydawnictwo Uniwersytetu Medycznego w Łodzi, Łódź 2022.
- Zalewski T., *Prawo sztucznej inteligencji*, CH Beck, Warszawa 2020.

ARTIFICIAL INTELLIGENCE IN MEDICINE: APPLICATION POSSIBILITIES AND ETHICAL AND LEGAL CHALLENGES

Summary: The primary objective of this publication is to present the growing role of artificial intelligence in medicine, analyzing both its potential and challenges related to regulatory and ethical aspects. The article discusses key areas in which artificial intelligence can improve diagnostic and therapeutic processes. The principles according to which the aforementioned technology should be used in health-related areas are indicated. Attention has also been paid to the problem of liability for mistakes made in connection with the use of artificial intelligence algorithms. Both ethical challenges and legal issues related to the development of this technology in medical fields are presented.

Keywords: artificial intelligence, medicine, treatment, risk, algorithms.