

Medycyna w cieniu algorytmów: Wyzwania wiarygodnej AI



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł podejmuje krytyczną analizę wdrażania systemów sztucznej inteligencji w sektorze ochrony zdrowia, koncentrując się na paradygmacie Trustworthy AI. Autor wskazuje, że techniczna sprawność algorytmów to za mało – niezbędna jest architektura zaufania oparta na przejrzystości, wyjaśnialności i rygorystycznym zarządzaniu ryzykiem w całym cyklu życia produktu (Total Product Life Cycle). Tekst przybliży standardy międzynarodowe, takie jak wytyczne WHO, NIST oraz regulacje FDA i UE, które definiują oprogramowanie medyczne jako systemy wysokiego ryzyka. Kluczowym elementem jest przejście od „czarnych skrzynek” do modeli, które pozwalają lekarzom na świadome podejmowanie decyzji terapeutycznych, minimalizując ryzyko biasu i błędów algorytmicznych w realnym środowisku klinicznym.

This article critically examines the implementation of artificial intelligence systems in the healthcare sector, focusing on the Trustworthy AI paradigm. The author argues that the technical efficiency of algorithms is not enough – a trust architecture based on transparency, explainability, and rigorous risk management throughout the product life cycle (TPLC) is essential. The article explores international standards, such as WHO and NIST guidelines, as well as FDA and EU regulations, which define medical software as high-risk systems. A key element is the shift from "black boxes" to models that enable physicians to make informed therapeutic decisions, minimizing the risk of bias and algorithmic errors in real-world clinical environments.

Współczesna medycyna stoi u progu cywilizacyjnej zmiany, w której sztuczna inteligencja przestaje być jedynie narzędziem diagnostycznym, stając się współtwórcą infrastruktury decyzji terapeutycznych. Artykuł analizuje paradygmat wiarygodnej AI (Trustworthy AI) jako fundamentalny ład instytucjonalny, niezbędny do ochrony zdrowia i podmiotowości pacjenta. Autorzy stawiają tezę, że techniczna sprawność

algorytmów – choć imponująca – jest niewystarczająca bez rygorystycznej wyjaśnialności, przejrzystości danych oraz jasnej odpowiedzialności prawnej. W obliczu dynamicznego rozwoju neurotechnologii, takich jak głęboka stymulacja mózgu czy systemy BCI, medycyna musi odejść od technokratycznego optymizmu na rzecz odpowiedzialnego nadzoru. Tekst demaskuje niebezpieczeństwo „metafizyki venture capital”, gdzie złożoność kodu bywa mylona z rzetelnością naukową. Głównym wyzwaniem staje się zatem zbudowanie systemu, w którym innowacja nie kolonizuje intymności i nie pogłębia systemowych nierówności. W obliczu unijnych regulacji, takich jak AI Act, artykuł wskazuje na konieczność ochrony integralności kognitywnej jednostki, argumentując, że zaufanie do technologii musi być twardym efektem zarządzania ryzykiem, a nie dziecinną fascynacją cyfrowym postępem.

SŁOWA KLUCZOWE

Trustworthy AI

wyjaśnialność AI

Total Product Life Cycle

bias w danych

transparentność modeli

medycyna nawigowana przez modele

modele wielomodalne

nadzór człowieka

odpowiedzialność deliktowa

architektura zaufania

algorytmiczna diagnostyka obrazowa

zarządzanie ryzykiem

standardy etyczne AI

integralność kognitywna

Wiarygodna AI: Od technicznej sprawności do ładu zaufania

Istnieją epoki, które rozpoznają własną wielkość zbyt wcześnie, i takie, które własne niebezpieczeństwa dostrzegają zdecydowanie zbyt późno. Nasza, z właściwą sobie pychą, zdaje się mieć ambicję należeć do obu tych kategorii naraz. O sztucznej inteligencji mówi się bowiem tonem niemal mesjańskim, lecz wdraża się ją coraz częściej tam, gdzie pomyłka nie jest jedynie kompromitacją inżyniera, lecz cudzym bólem, kalectwem czy utratą autonomii. Dlatego właśnie pojęcie Trustworthy AI, czyli koncepcja wiarygodnej sztucznej inteligencji opartej na jakości danych, przejrzystości modelu i jasnej odpowiedzialności, nie jest tylko ozdobnym ornamentem do raportów compliance. To raczej nazwa nowego, fundamentalnego sporu o to, czy medycyna cyfrowa stanie się trwałym ustrojem zaufania, czy jedynie przemysłem zautomatyzowanego ryzyka.

Współczesny paradygmat naukowy przestał sprowadzać się do prostego pytania o to, czy algorytm potrafi coś przewidzieć. Dziś pytamy raczej, czy czyni to w sposób dający się obronić technicznie, etycznie i cywilizacyjnie, co widać w najnowszych wytycznych WHO dotyczących modeli wielomodalnych. Organizacja ta podkreśla, że generatywna AI w ochronie zdrowia wymaga odrębnej refleksji nad nadzorem, by nie stała się nowoczesną formą bezpańskości, w której nikt nie bierze odpowiedzialności za błąd. NIST z kolei ujmuje sprawę jeszcze trzeźwiej, traktując zarządzanie ryzykiem jako warunek wbudowania wiarygodności w cały cykl życia systemu, a nie jako kosmetykę nakładaną po fakcie. Unia Europejska, działając jako mocarstwo normatywne budujące przewagę przez rygorystyczne standardy etyczne, wpisała tę intuicję do prawa, uznając oprogramowanie medyczne za kategorię wysokiego ryzyka.

To już nie jest dyskusja salonowa, lecz zmiana samej konstytucji praktyki technicznej, gdzie wiarygodna AI staje się próbą ucywilizowania maszynowej mocy predykcyjnej przez włączenie jej do ładu instytucjonalnego. Nie wystarczy bowiem, że system działa często; musi on działać niezawodnie w każdych warunkach. Nie wystarczy, że algorytm osiąga wysoką skuteczność na sterylnym zbiorze walidacyjnym, skoro musi zachowywać bezpieczeństwo w chaotycznym, realnym środowisku klinicznym. Nie wystarczy również, że model nie deklaruje uprzedzeń, jeśli w praktyce reprodukuje je w skali, która z czasem przeobraża się w systemową niesprawiedliwość. Wreszcie, nie wystarczy, że technologia zachwyca inwestorów swoją sprawnością, gdyż musi ona pozostawać kompatybilna z prawami pacjenta, wymogami dowodowymi medycyny oraz logiką odpowiedzialności deliktowej.

W tym ujęciu medycyna nawigowana przez modele, czyli system, w którym AI nie tylko klasyfikuje dane, ale współtworzy infrastrukturę decyzji terapeutycznych, musi uwzględniać fakt, że człowiek nie jest tylko serią cech wejściowych. Jesteśmy istotami ucieleśnionymi i historycznymi, a nie jedynie zbiorem danych w Total Product Life Cycle, czyli podejściu regulacyjnym obejmującym zarządzanie ryzykiem od projektu po rynek. Dlatego coraz częściej zamiast o samym kodzie mówi się o architekturze zaufania, która musi otaczać każdą innowację, od prostych diagnoz po aDBS, czyli spersonalizowaną stymulację mózgu dostosowującą się do pacjenta w czasie rzeczywistym. Zaufanie w tym kontekście przestaje być ulotną emocją czy przejawem dziecinnej fascynacji skutecznością liczbową. Staje się ono twardym, instytucjonalnym fundamentem, bez którego technologia w medycynie pozostaje jedynie ryzykownym eksperymentem na żywym organizmie społecznym.

Wyjaśnialność AI: Od technicznej metafizyki do praktyki

Wiarygodność algorytmu medycznego nie jest darem niebios ani efektem ubocznym sprawnego marketingu, lecz wypadkową rygorystycznej higieny danych, przejrzystości architektury oraz odporności na nieuchronne zmiany kontekstu klinicznego. Amerykańska FDA, daleka od poetyckich uniesień, sprowadza ten problem do konkretnej kategorii transparentności urządzeń opartych na uczeniu maszynowym, czyli stopnia komunikowania informacji o ich przeznaczeniu, rozwoju i wydajności. Według regulatora jasność co do logiki działania systemu stanowi fundament bezpieczeństwa, pozwalając na wczesne wykrywanie błędów, spadków formy algorytmu czy ukrytego biasu. W medycznej codzienności każda „czarna skrzynka” prędzej czy później przestaje być zgrabną metaforą, a staje się bardzo konkretnym problemem dowodowym na sali sądowej. Musimy zatem przyjąć, że wyjaśnialność nie jest jedynie luksusem poznawczym dla akademików o skłonnościach do metafizyki, lecz twardym narzędziem porządkowania odpowiedzialności.

W istocie wyjaśnialność (explainability) służy przede wszystkim jako precyzyjny mechanizm nawigacji w trójkącie pacjent-lekarz-maszyna, gdzie stawką jest zdrowie i życie. Jeżeli system rekomenduje inwazyjną interwencję lub sugeruje intensywność stymulacji neurologicznej, klinicysta musi rozumieć mechanizm tej sugestii na tyle, by odróżnić rzetelną wskazówkę od zwykłego błędu. Chodzi o przejście od medycyny opartej na intuicji kodu do epistemologii praktycznej, gdzie lekarz potrafi oddzielić korelację od przyczyny, a sygnał diagnostyczny od technicznego artefaktu. NIST wpisuje te starania w szerszy ład Trustworthy AI, czyli koncepcję wiarygodnej sztucznej inteligencji opartej na jakości danych i pełnej przejrzystości modelu. Z kolei Komisja Europejska traktuje nadzór człowieka nad systemami wysokiego ryzyka jako cywilizacyjny standard, a nie opcjonalny dodatek do oprogramowania.

Obserwujemy dziś zbawienne odejście od dziecinnej fascynacji skutecznością liczbową na rzecz dojrzałego pytania o uzasadnialność każdego podjętego przez algorytm kroku. Wyjaśnialność nie obiecuje nam przecież cudu pełnego odsłonięcia każdej matematycznej warstwy modelu głębokiego, co byłoby zresztą poznawczo bezużyteczne. Zamiast tego oferuje wgląd w to, które cechy wejściowe przeważały o wyniku i na jakich populacjach system zaczyna tracić swoją pierwotną skuteczność. Pozwala to monitorować stabilność narzędzia po zmianie środowiska, co w medycynie nawigowanej przez modele, czyli systemie współtworzącym infrastrukturę decyzji terapeutycznych, jest kwestią elementarnego bezpieczeństwa. To właśnie ta pragmatyka wymusza na przemyśle technologicznym budowę całego instrumentarium nadzoru, które ma okiełznać spontaniczną innowacyjność.

Giganci tacy jak IBM czy AWS od lat rozwijają środowiska mające zapewniać wgląd w trzewia algorytmów, wykrywać dryf danych oraz wspierać atrybucję cech. IBM utrzymuje Watson OpenScale jako platformę do walki z biasem, podczas gdy Amazon stawia na SageMaker Clarify i matematyczne wartości Shapleya do wyjaśniania predykcji. Microsoft z kolei osadza swoje działania w ramach Responsible AI Standard, wiążąc technologię z zasadami sprawiedliwości, niezawodności oraz szeroko pojętej prywatności i bezpieczeństwa. Na tym tle niezwykle intrygująco wygląda niedawna wolta Google, które zdecydowało się wygasić swój sztandarowy projekt Vertex Explainable AI do marca 2027 roku. Ta decyzja nie jest jedynie nudnym detalem technicznym, lecz sygnałem głębokiej reorganizacji całego pola bitwy o zaufanie użytkownika.

Wycofywanie sprawdzonych narzędzi pokazuje, że zaufanie do sztucznej inteligencji nie rodzi się w świecie stabilnych dogmatów, lecz w tyglu instytucjonalnego eksperymentu pod presją regulacyjną. Jeszcze kilka lat temu część środowiska technologicznego zachowywała się tak, jakby najlepszym argumentem za bezpieczeństwem systemu była imponująca liczba slajdów na konferencji. Panowało przekonanie, że im więcej gradientów i skomplikowanych wykresów ROC, tym mniej pytań o etykę i realną skuteczność kliniczną. Była to oczywiście metafizyka venture capital przebrana za metodologię nauki, w której złożoność algorytmu miała automatycznie gwarantować jego słuszność. Dzisiejszy paradygmat jest już znacznie mniej naiwny, odrzucając kapitulację rozumu ubraną w język postępu na rzecz twardej, weryfikowalnej transparentności.

Współczesne podejście, takie jak Total Product Life Cycle, czyli strategia regulacyjna obejmująca zarządzanie ryzykiem na każdym etapie życia produktu, od projektu po rynek, zmienia reguły gry. Nie wystarczy, że algorytm działał w laboratorium; musi on zachować swoją integralność w brudnej rzeczywistości szpitalnego oddziału. Unia Europejska, działając jako mocarstwo normatywne, czyli podmiot budujący przewagę poprzez rygorystyczne standardy etyczne i prawne, stawia na ochronę pacjenta ponad czystą moc obliczeniową. W tym nowym rozdaniu wyjaśnialność staje się pomostem między abstrakcyjnym kodem a konkretnym ciałem pacjenta poddawany procedurom takim jak aDBS. Ta spersonalizowana, głęboka stymulacja mózgu wymaga od nas pełnego zrozumienia, dlaczego impuls elektryczny został wysłany właśnie w tej milisekundzie, a nie innej.

Wiarygodność AI: Między techniczną trafnością a odpowiedzialnością

Globalne instytucje, takie jak WHO, NIST czy FDA, porzucają powoli technokratyczny optymizm na rzecz twardego realizmu prawniczego. Unia Europejska, budując swój wizerunek jako mocarstwo normatywne, czyli podmiot kreujący globalną przewagę poprzez rygorystyczne

standardy etyczne, stawia sprawę jasno. Wdrożenie algorytmów w medycynie wymaga nie tylko potężnej mocy obliczeniowej, lecz przede wszystkim żelaznych procedur dowodowych i mechanizmów nadzoru. Komisja Europejska kładzie nacisk na ramy prawne umożliwiające skuteczne dochodzenie odszkodowań od producentów w przypadku wystąpienia błędu. Ten pozornie nudny, juretyczny detal odsłania fundamentalną prawdę o naturze współczesnej technologii. Wiarygodność nie jest bowiem cechą kodu, lecz relacją między procesem projektowania a możliwością przypisania realnych konsekwencji.

Maszyna, której nikt nie może pociągnąć do odpowiedzialności, w istocie nie zasługuje na nasze zaufanie. Pozostaje ona jedynie nowoczesną formą bezpańskości, gdzie błąd rozplywa się w cyfrowym niebycie. Na tym tle refleksja Fundacji Dobre Państwo jawi się jako rzadki przykład intelektualnej dojrzałości. Zamiast ulegać dziecięcej egzaltacji wobec technologicznych nowinek, autorzy słusznie wskazują na strukturalne asymetrie, które algorytmy mogą niebezpiecznie utrwalać. Szczególnie trafna okazuje się analiza luki w danych, gdzie historycznie męskocentryczne zbiory prowadzą do systemowego zafałszowania diagnoz. Taka sytuacja to nie tylko błąd techniczny, ale przede wszystkim reprodukcja dyskryminacji wobec kobiet.

To ujęcie jest nie tylko etycznie nośne, ale również metodologicznie bezbłędne w swojej surowości. Gdy Komisja Europejska podkreśla wagę wysokiej jakości danych, promuje ona Trustworthy AI, czyli koncepcję wiarygodnej sztucznej inteligencji opartej na przejrzystości i odpowiedzialności. FDA łączy transparentność z ideą sprawiedliwości zdrowotnej, co potwierdza intuicję, że nierówność w danych to forma redystrybucji ryzyka. Tam, gdzie zbiory treningowe nie reprezentują populacji, model nie tylko myli się częściej w sensie statystycznym. On świadomie przerzuca koszty tych pomyłek na grupy, które już wcześniej były słabiej widzialne. W ekonomii nazwalibyśmy to asymetrią informacyjną, natomiast w prawie stanowi to gotowy materiał pod spór o dyskryminację.

Właśnie w tym punkcie zaczyna się prawdziwa polityka medycyny algorytmicznej, rozumiana jako walka o kształt naszej infrastruktury poznawczej. Musimy pytać, kto dostarcza dane i kto ustala progi czułości dla systemów ratujących życie. Kto definiuje akceptowalny poziom błędu i kto ostatecznie ponosi koszt fałszywie ujemnego wyniku? Kto finansuje walidację na populacjach rzadziej badanych, a kto zyskuje na samej automatyzacji procesów klinicznych? To nie są pytania techniczne, lecz fundamentalne kwestie dotyczące podziału władzy w systemie ochrony zdrowia. Odpowiedzi na nie zadecydują, czy technologia będzie nas wyzwalać, czy jedynie pogłębiać istniejące od dekad nierówności.

Fundacja Dobre Państwo trafnie wiąże cyfrowe technologie z problemem wolności pozytywnej oraz integralności kognitywnej. W tekstach inspirowanych myślą Snydera ochrona przed manipulacją poznawczą zostaje wpisana w samą architekturę nowoczesnej wolności. Ma to kolosalne znaczenie

dla medycyny nawigowanej przez modele, czyli modelu opieki, w którym AI współtworzy infrastrukturę decyzji terapeutycznych. System może być technicznie skuteczny, a zarazem politycznie regresywny, jeśli zwiększa zależność pacjenta od nieprzejrzystych monopolii. Zaufanie nie powstaje przecież wtedy, gdy maszyna przemawia do nas kojącym, syntetycznym głosem. Powstaje ono dopiero wtedy, gdy instytucja gwarantuje obywatelowi realną zdolność rozumienia i sprzeciwu.

Gdy przejdziemy do diagnostyki obrazowej, zrozumiemy, dlaczego AI odniosła tam tak spektakularny sukces rynkowy. Obraz medyczny jest dla modelu atrakcyjny, ponieważ dostarcza ustrukturyzowany materiał i pozwala na wykrywanie wzorców niewidocznych dla ludzkiego oka. Dlatego radiologia i patomorfologia stały się pierwszymi poligonami masowej adopcji tych rozwiązań w praktyce klinicznej. FDA prowadzi oficjalną listę urządzeń, aby zwiększać transparentność rozwiązań ML, czyli stopień komunikowania informacji o ich przeznaczeniu i logice działania. Większość dopuszczonych rozwiązań należy do obszaru radiologii, co wynika bezpośrednio z dojrzałości ścieżek komercjalizacji. Niemniej jednak regulatorzy słusznie zauważają, że bezpieczeństwo nie kończy się na samym momencie dopuszczenia do obrotu.

W 2025 roku FDA opublikowała wytyczne kładące nacisk na Total Product Life Cycle, czyli podejście regulacyjne obejmujące zarządzanie ryzykiem na każdym etapie życia produktu. Oznacza to ostateczne odejście od postrzegania urządzenia medycznego jako bytu statycznego i niezmiennego. Model jest procesem, a nie tylko gotowym produktem zamkniętym w pudełku z instrukcją. Proces pozbawiony ciągłego monitoringu staje się bardzo szybko ruiną z opóźnionym zapłonem dla całego systemu. W tym kontekście wyjaśnialność w obrazowaniu medycznym przestaje być jedynie dydaktycznym dodatkiem dla studentów. Staje się ona kluczowym instrumentem obrony przed dwiema groźnymi iluzjami poznawczymi, które trawią współczesną medycynę.

Pierwsza iluzja głosi, że skoro model przewiduje trafnie, to zapewne rozumie on istotę problemu klinicznego. Druga iluzja sugeruje, że jeśli radiolog nie widzi błędu gołym okiem, to błąd ten po prostu nie istnieje. Współczesna literatura pokazuje jednak, że systemy wyjaśnialne mają sens tylko wtedy, gdy pozwalają odróżnić sygnał od artefaktów. W praktyce musimy wiedzieć, czy system kierował się patologią, czy może specyfiką aparatu lub znacznikiem instytucji. To jest właśnie miejsce, w którym filozofia nauki spotyka się z brutalną praktyką kliniczną. Model bez wyjaśnialności to często metafizyka venture capital przebrana za metodologię nauki. Trafność statystyczna nie jest bowiem tym samym, co rzetelność epistemiczna.

Algorytmiczna nawigacja: między wsparciem a kolonizacją intymności

Radiologia i neurologia przestały być polami czystej obserwacji, stając się poligonem, na którym sztuczna inteligencja buduje nową infrastrukturę ludzkiego działania. W planowaniu głębokiej stymulacji mózgu algorytmy nie tylko wspierają fuzję obrazów CT i MRI, ale wręcz rekonstruuja architekturę istoty białej. Mamy tu do czynienia z medycyną nawigowaną przez modele, czyli systemem, w którym AI współtworzy ramy decyzji terapeutycznej, zamiast tylko opisywać statyczne zdjęcia. W praktyce oznacza to, że chirurg staje się operatorem interfejsu, którego precyzja zależy od jakości cyfrowego rozpoznania topografii neuronów. To niewątpliwie cywilizacyjny przełom, niemniej rzuca on rękawicę tradycyjnemu pojmowaniu odpowiedzialności zawodowej lekarza.

Sytuacja, w której autorytet naukowy firmuje swoim nazwiskiem rekomendację systemu zrozumiałego jedynie dla garstki inżynierów, ociera się o technologiczną groteskę. Zamiast partnerskiej współpracy człowieka z maszyną obserwujemy arystokratyczną formę podporządkowania się nowej scholastyce kodu, gdzie algorytm staje się ostateczną wyrocznią. Im bardziej neurochirurg ufa warstwie obliczeniowej, tym mocniej państwo musi weryfikować standardy jej walidacji oraz ciągłego monitorowania. Jest to swoista kapitulacja rozumu ubrana w język postępu, jeśli nie zapewnimy pełnej transparentności urządzeń ML. Przejrzystość ta, czyli jasne komunikowanie logiki działania modelu, stanowi wszak jedyny bezpiecznik chroniący nas przed dyktaturą czarnej skrzynki.

W psychiatrii stawka rośnie, ponieważ przedmiotem interwencji nie jest już tylko trajektoria elektrody, lecz sama tkanka ludzkiego nastroju i języka. Modele oparte na cyfrowym fenotypowaniu, czyli zbieraniu danych z urządzeń ubieralnych w celu mapowania zachowań, obiecują nam dziś wczesne wykrywanie kryzysów psychicznych. Systematyczne przeglądy literatury z lat 2025 i 2026 wskazują, że rozwój ten ogniskuje się wokół diagnozy, predykcji ryzyka oraz cyfrowych interwencji. Niemniej ta sama literatura akcentuje bariery interpretowalności oraz ryzyko systemowego uprzedzenia, które może zniekształcać obraz cierpienia pacjenta. Wszak algorytm nie posiada empatii, a jedynie statystyczną korelację między słowem a prawdopodobnym stanem afektywnym. To nie jest już medycyna, lecz metafizyka venture capital przebrana za metodologię nauki.

Nie jest dziełem przypadku, że Fundacja Dobre Państwo w swoich postulatach dotyczących wolności pozytywnej tak silnie akcentuje ochronę zdrowia psychicznego młodzieży. W dobie monopolu danych granica między wsparciem terapeutycznym a kolonizacją intymności staje się niebezpiecznie cienka i łatwa do przeoczenia. AI w psychiatrii może bowiem niepostrzeżenie przejść

od roli asystenta do roli zarządcy zachowań, profilując użytkownika pod kątem jego podatności poznawczej. Musimy zatem zachować zimną krew, aby odróżnić autentyczną opiekę od subtelnej formy manipulacji ukrytej pod płaszczem optymalizacji dobrostanu. W tym kontekście mocarstwo normatywne, czyli strategia Unii Europejskiej budująca przewagę poprzez rygorystyczne standardy etyczne, wydaje się jedyną sensowną drogą.

Chatboty oferujące elementy terapii poznawczo-behawioralnej mogą być użyteczne jako narzędzia niskiego progu dostępu, zwłaszcza w systemach chronicznie niedofinansowanych. Trzeba jednak wyraźnie zaznaczyć, że powszechna dostępność cyfrowego asystenta nie jest tożsama z równoważnością terapeutyczną profesjonalnej relacji z człowiekiem. WHO przypomina, że generatywne modele wielomodalne wymagają rygorystycznego podejścia do ryzyka, aby nie stały się źródłem jatrogennych szkód. Z kolei FDA promuje podejście Total Product Life Cycle, czyli zarządzanie bezpieczeństwem produktu od fazy projektu aż po jego rynkową obecność. Tylko taka kontrola pozwala ufać, że Trustworthy AI, oparta na przejrzystości i odpowiedzialności, faktycznie służy pacjentowi.

Od wsparcia decyzji do rekonstrukcji układu nerwowego

Musimy zacząć od nakreślenia bardzo wyraźnej linii demarkacyjnej, ponieważ w obecnym dyskursie wsparcie technologiczne zbyt często mylone jest z leczeniem tylko dlatego, że posiada interfejs imitujący ludzką rozmowę. Człowiek znajdujący się w egzystencjalnym kryzysie nie potrzebuje przecież kolejnego technologicznego fetyszu responsywności, lecz domaga się bezpiecznej i stabilnej architektury opieki. Tam, gdzie cyfrowe protezy stają się tanim substytutem niewydolnych instytucji, zamiast ich mądrym rozszerzeniem, kończy się medycyna, a zaczyna polityka cynicznej iluzji. Interfejs człowiek-maszyna w obszarze zdrowia jest właśnie tym newralgicznym punktem, w którym cała filozofia Trustworthy AI, czyli koncepcja wiarygodnej sztucznej inteligencji opartej na jakości danych i przejrzystości modelu, staje się boleśnie namacalna. W klasycznym ujęciu inżynierskim systemy te projektowano z myślą o intuicyjności, efektywności oraz ergonomii, lecz w dobie algorytmizacji to już zdecydowanie za mało.

Współczesny interfejs musi uwzględniać fakt, że jego użytkownik jest niemal zawsze obciążony poznawczo, cierpi, działa pod ogromną presją czasu lub bywa po prostu zdezorientowany i zmęczony wiekiem. W praktyce klinicznej oznacza to, że medyczny HMI przestał być wyłącznie domeną użyteczności, stając się fundamentalną kwestią bezpieczeństwa oraz ludzkiej godności. Amerykańska FDA słusznie łączy transparentność urządzeń ML, czyli stopień komunikowania

informacji o logice działania systemów uczenia maszynowego, z paradygmatem projektowania zorientowanego na człowieka. Takie podejście powinno obejmować całość doświadczenia pacjenta, walidować przejrzystość komunikatów i dostarczać wszelkich niezbędnych informacji o stanie urządzenia bez zbędnego żargonu. Ma to szczególną wagę w momentach, gdy urządzenia medyczne zyskują warstwę konwersacyjną, głosową lub predykcyjną, co naturalnie zmienia dynamikę zaufania. Absurd naszej współczesności polega bowiem na tym, że człowiek potrafi ufać syntetycznemu głosowi bardziej niż papierowej instrukcji obsługi, byle tylko ten głos brzmiał odpowiednio kojąco.

Dlatego właśnie projektowanie interfejsów w medycynie musi być traktowane zarazem jako precyzyjna inżynieria i swoista pedagogika ograniczeń, która uczy użytkownika realnych możliwości maszyny. Z tego samego powodu integracja AI z Internetem Rzeczy oraz technologiami ubieralnymi nie może być opisywana wyłącznie w kategoriach rynkowej wygody czy mitycznej skalowalności. Owszem, połączenie czujników, domowych urządzeń medycznych i zdalnej analityki otwiera fascynującą drogę do wcześniejszego wykrywania pogorszenia stanu zdrowia oraz redukcji obciążenia szpitali. Komisja Europejska podkreśla przy tym, że rozwój ten wymaga dostępu do zróżnicowanych danych, a EHDS, czyli europejska przestrzeń ułatwiająca wtórne wykorzystanie elektronicznych danych zdrowotnych, ma stać się paliwem dla innowacji. Niemniej jednak, równoległe do tych obietnic rośnie stawka sporu o interoperacyjność, prywatność oraz granice wtórnego użycia informacji o najbardziej intymnych sferach naszego życia.

Im więcej system wie o pacjencie, tym głośniej i częściej musimy pytać o to, kto właściwie posiada wiedzę o samym systemie i jego wewnętrznych procesach. Musimy wiedzieć, kto nadzoruje jego autonomiczne decyzje, kto audytuje błędy algorytmiczne, kto ma dostęp do logów i kto ostatecznie odpowiada za ewentualny wyciek danych. W tym sensie IoT w medycynie nie jest jedynie niewinnym dodatkiem do sztucznej inteligencji, lecz stanowi realne rozszerzenie jej jurysdykcji nad codziennością pacjenta. Wiarygodna AI w medycynie nie jest więc problemem jednej, konkretnej technologii, lecz wyzwaniem dotyczącym całego ładu cywilizacyjnego, w którym obecnie funkcjonujemy. Wymaga ona równoczesnego myślenia kategoriami ekonomisty, prawnika oraz medyka, ponieważ stawką jest tutaj alokacja ryzyka, ochrona praw podstawowych oraz trafność kliniczna.

Fundacja Dobre Państwo, wskazując na lukę w danych oraz na konieczność ochrony integralności kognitywnej, trafia w sam środek tego wielowarstwowego sporu o przyszłość podmiotowości. Nie chodzi przecież o to, by spowalniać rozwój technologii przez irracjonalny lęk, lecz o to, by nie budować kolejnej cywilizacji wydajności, która każdą nieprzejrzystość nazywa innowacją. Często mamy tu do czynienia z metafizyką venture capital przebraną za metodologię nauki, gdzie każdy

koszt społeczny przerzuca się na barki najsłabszych uczestników systemu. To jednak dopiero początek drogi, ponieważ prawdziwe wyzwania zaczynają się tam, gdzie AI łączy się z interfejsami mózg-komputer oraz robotyką rehabilitacyjną. Kiedy zejdziemy głębiej, w obszar nanotechnologii i systemów typu closed-loop, okaże się, że stawką przestaje być tylko wspomaganie decyzji, a staje się nią rekonstrukcja kanału komunikacji.

Pierwsza warstwa sporu o wiarygodną sztuczną inteligencję dotyczyła tego, czy wolno jej ufać w diagnozie, natomiast druga dotyka już kwestii znacznie bardziej radykalnej i intymnej. Nie chodzi wyłącznie o to, czy maszyna dobrze rozpoznaje symptomy, lecz o to, czy może stać się integralnym przedłużeniem ludzkiego układu nerwowego. Tu kończy się komfortowy obszar administracji medycznej, a zaczyna właściwe laboratorium nowej antropotechniki, w którym maszyna filtruje sygnały, których człowiek nie potrafi już samodzielnie przetworzyć. Integracja AI oraz BCI nie jest tylko kolejnym rozdziałem w podręczniku inżynierii, lecz próbą przebudowy relacji między mózgiem, ciałem a instytucją leczniczą. Nowoczesne badania sugerują, że rozwój ten zmierza w stronę zamkniętych pętli sterowania, choć największą przeszkodą pozostaje wciąż długoterminowa stabilność sygnału i jego bezpieczna interpretacja.

To jest właśnie miejsce, w którym technologiczna obietnica i etyczna odpowiedzialność spotykają się w swojej najostrzejszej, niemal bezwzględnej formie. W języku najbardziej oszczędnym BCI jest próbą zamiany sygnału neuronalnego na rozkaz dla świata zewnętrznego lub odwrotnie, co dawniej brzmiało jak czysta neurofantastyka. Dziś jest to coraz częściej dojrzała praktyka kliniczna, wspierająca komunikację u osób z porażeniami oraz obsługę zaawansowanych neuroprotez ruchowych. Największa zmiana polega jednak na tym, że sztuczna inteligencja przestała być w tych systemach prostym klasyfikatorem danych wejściowych. Staje się ona adaptacyjnym negocjatorem między niestabilną, ludzką fizjologią a sztywnymi wymaganiami technicznymi urządzenia, z którym pacjent musi współistnieć.

Zamiast szukać jednej, uniwersalnej reguły decyzyjnej, buduje się dziś systemy uczące się konkretnego pacjenta, jego unikalnego rytmu neuronalnego oraz dobowych zmian zmęczenia. To prowadzi nas bezpośrednio do głębokiej stymulacji mózgu, która jest znakomitą przykładem przejścia medycyny od logiki prostego urządzenia do logiki systemu cyberfizycznego. Klasyczny DBS działał w gruncie rzeczy jak dość toporny mechanizm podający zaprogramowaną stymulację wedle sztywnych, ustalonych odgórnie parametrów. Był on skuteczny w chorobie Parkinsona, ale obciążony ograniczeniami wynikającymi z ciągłego i całkowicie niekontekstowego pobudzania struktur mózgowych. Nowsze podejścia opisują przejście do aDBS, czyli spersonalizowanej stymulacji dostosowującej się do patologicznych oscylacji w czasie rzeczywistym, co stanowi rewolucyjny skok jakościowy.

Jest to przejście od medycyny uśrednionej do medycyny sytuacyjnej, gdzie nie stymuluje się już statystycznego pacjenta, lecz konkretnego człowieka w określonym momencie jego aktywności. W badaniach z 2025 roku adaptacyjny DBS opisywany jest jako jeden z najbardziej obiecujących kierunków, choć autorzy słusznie studzą entuzjazm, wskazując na problemy z trwałością efektu. Warto tutaj zachować intelektualną dyscyplinę, ponieważ w tym miejscu wyjątkowo łatwo jest popaść w technologiczną egzaltację i zapomnieć o ograniczeniach materii. Adaptacyjny DBS nie jest przecież magicznym urządzeniem, które "czyta w myślach" w sensie popkulturowym, lecz systemem sterowania opartym na biologicznym sprzężeniu zwrotnym. W praktyce chodzi o to, by wyodrębnić sygnał uznany za istotny klinicznie i w odpowiedzi na niego precyzyjnie modulować parametry stymulacji, co jest raczej triumfem inżynierii niż kapitulacją rozumu ubraną w język postępu.

Wiarygodność AI w neurotechnologii: Pętla sprzężenia zwrotnego

W chorobie Parkinsona głównym kandydatem na biomarker są patologiczne oscylacje w paśmie beta, stanowiące swoisty rytmiczny odcisk palca tego schorzenia. Doniesienia kliniczne oraz badania długoterminowe sugerują, że spersonalizowany aDBS, czyli adaptacyjna głęboka stymulacja mózgu dostosowująca się do potrzeb pacjenta w czasie rzeczywistym, może znacząco poprawić kontrolę objawów. System ten obiecuje zmniejszenie ekspozycji na zbędne impulsy elektryczne oraz ograniczenie dotkliwych działań niepożądanych, które często towarzyszą tradycyjnej terapii. Niemniej jednak urządzenia te nieustannie zmagają się z tak zwanym artefaktem stymulacyjnym. Sygnał, który pragniemy monitorować, bywa skutecznie zasłaniany przez aktywność samej interwencji. Jest to analogiczne do próby usłyszenia szeptu pacjenta przy jednoczesnym działaniu głośnego źródła dźwięku.

Właśnie dlatego współczesny paradygmat zamkniętej pętli nie jest wyłącznie opowieścią o coraz inteligentniejszych algorytmach, lecz przede wszystkim o mozolnym doskonaleniu technik tłumienia artefaktów i doboru biomarkerów. Systemy typu closed loop stanowią jeden z najbardziej wymownych dowodów na to, że Trustworthy AI, czyli koncepcja wiarygodnej sztucznej inteligencji opartej na jakości danych i przejrzystości, nie może oznaczać jedynie estetycznego interfejsu. Tutaj wiarygodność musi obejmować samą fizjologię sprzężenia zwrotnego, stając się fundamentem dopuszczalności technologii w sferze ludzkiego mózgu. System ma reagować szybko, trafnie i proporcjonalnie do zaistniałej potrzeby neuronalnej. Ma on nie tylko minimalizować błąd predykcyjny, lecz także unikać pułapki nadmiernej stymulacji, ograniczać koszty energetyczne i

pozostawać pod ścisłą kontrolą klinicysty. Wszak w medycynie precyzja bez nadzoru to jedynie kosztowna iluzja bezpieczeństwa.

W tym miejscu ekonomista dostrzega klasyczny problem optymalizacji przy twardych ograniczeniach, podczas gdy prawnik analizuje widmo odpowiedzialności za automatyczną decyzję terapeutyczną. Medyk natomiast rozpoznaje dylemat, czy większa precyzja techniczna rzeczywiście przekłada się na realną korzyść dla cierpiącego człowieka. Nauka ostatnich lat odpowiada na te pytania z dużą ostrożnością, świadoma, że potencjał technologii jest ogromny, ale jej translacja kliniczna bywa bolesna i nierówna. Różnica między sukcesem w badaniu a codzienną praktyką bywa znacząca, przypominając nam o ryzyku, jakim jest kapitulacja rozumu ubrana w język postępu. Często bowiem zapominamy, że zaawansowany model to wciąż tylko uproszczenie rzeczywistości, a nie sama rzeczywistość. Nie wystarczy, że system działa w laboratorium; musi on przetrwać zderzenie z nieprzewidywalnością ludzkiego życia.

W tej skomplikowanej architekturze ogromne znaczenie zyskuje SPES, czyli stymulacja pojedynczym impulsem elektrycznym, oraz analiza potencjałów korowo-korowych (CCEP). Dla laika terminy te brzmią zapewne jak ezoteryczny detal neurofizjologii, jednak dla współczesnej neuronauki klinicznej jest to jedna z najważniejszych metod przejścia od prostych korelacji do funkcjonalnej mapy przyczynowej mózgu. SPES pozwala wywołać kontrolowaną odpowiedź w określonych sieciach neuronalnych i obserwować, jak aktywność rozchodzi się po całym układzie. W przeciwieństwie do wielu metod czysto obserwacyjnych, technika ta nie poprzestaje na rejestracji współwystępowania sygnałów, lecz daje wgląd w efektywną łączność sieci. Pozwala to badaczom nie tylko obserwować mózg, ale wręcz zadawać mu konkretne pytania i otrzymywać czytelne odpowiedzi.

Nowsze opracowania podkreślają, że CCEP są coraz szerzej wykorzystywane do badania sieci funkcjonalnych, planowania chirurgii padaczki oraz śródoperacyjnego monitorowania pacjenta. Służą one rozwijaniu zautomatyzowanych metod analizy sygnału, co ma kluczowe znaczenie dla medycyny nawigowanej przez modele, czyli podejścia, w którym AI współtworzy infrastrukturę decyzji terapeutycznych. Jednocześnie przeglądy metodologiczne pokazują ogromną różnorodność podejść do detekcji i interpretacji tych niezwykle czułych sygnałów. Ta wielość standardów sugeruje, że technologia wciąż wyprzedza nasze ramy pojęciowe. Wiarygodność systemu wymaga zatem nie tylko doskonałego kodu, ale i konsensusu co do tego, co właściwie mierzymy. Bez tego nawet najdoskonalsza pętla sprzężenia zwrotnego pozostanie jedynie wyrafinowanym instrumentem grającym do pustej sali.

Między obietnicą a kliniczną prozą: granice medycznej AI

To zjawisko odsłania klasyczny, niemal podręcznikowy kryzys wieku dojrzewania nowoczesnej technologii, w którym tempo wzrostu pola badawczego drastycznie wyprzedza procesy standaryzacji i rzetelnej weryfikacji. W tym gąszczu danych sztuczna inteligencja staje się narzędziem nieodzownym, lecz biada tym, którzy zaczną ją fetyszyzować jako ostateczną wyrocznię posiadającą klucze do prawdy absolutnej. Wielowymiarowe zbiory danych iEEG oraz potencjałów wywołanych są strukturami zbyt gęstymi, by mogły poddać się wyłącznie tradycyjnej, manualnej interpretacji badacza, która w obliczu takiej skali staje się anachronizmem. Algorytmy świetnie radzą sobie z czyszczeniem sygnału z artefaktów, precyzyjnym wyłapywaniem komponentów N1 i N2 czy modelowaniem hierarchii sieciowych w czasie rzeczywistym. Niemniej jednak wyciąganie z tej technicznej sprawności wniosku, że AI wkrótce „rozwiąże zagadkę mózgu”, jest przejawem uroczej, choć niebezpiecznej naiwności, którą można by nazwać ślepym zauroczeniem cyfrową sprawnością.

Algorytm rozwiązuje bowiem konkretny problem obliczeniowy, polegający na ekstrakcji ukrytych struktur z szumu informacyjnego, a nie problem ontologiczny czy kliniczny. Sam sens medyczny tych struktur musi wciąż przechodzić przez gęsty filtr wiedzy neurochirurgicznej, fizjologicznej i epileptologicznej, który nie daje się łatwo zamknąć w eleganckim kodzie. Mówiąc dosadniej: maszyna może wskazać statystycznie istotny wzorec, ale nie posiada kompetencji do wydania sądu o jego realnym znaczeniu dla dobrostanu konkretnego człowieka. Kto pomyli te dwa porządki, ten nieuchronnie zamieni rzetelną analizę sygnału w nową, cyfrową scholastykę automatyzmu, w której liczy się wynik, a nie jego zrozumienie. W medycynie ostateczny werdykt wciąż należy do człowieka, a każda próba obejścia tej zasady kończy się intelektualną rejteradą ubraną w szaty postępu.

Jeszcze wyraźniej ten rozdźwięk widać w obszarze rehabilitacji robotycznej, gdzie wielkie obietnice inżynierskie zderzają się z najbardziej upartym faktem klinicznym – żywym, nieprzewidywalnym pacjentem. Egzoszkielety, systemy typu end-effector czy roboty do treningu chodu tworzą dziś imponujący wachlarz narzędzi, który na targach medtech wygląda jak spełnienie marzeń o ostatecznej cyborgizacji medycyny. Nowsze metaanalizy sugerują wprawdzie, że robotycznie wspomagany trening chodu po udarze może poprawiać funkcje motoryczne i równowagę, zwłaszcza w duecie z terapią konwencjonalną. Jednak literatura przedmiotu daleka jest od euforycznej jednolitości, co powinno studzić zapał entuzjastów wierzących, że stal i krzem bezrefleksyjnie zastąpią dotyk terapeuty.

Część przeglądów wskazuje na korzyści istotne, inne zaś brutalnie punktuja brak przekroczenia progu klinicznie znaczącej poprawy w kluczowych wynikach funkcjonalnych pacjenta. W rehabilitacji kończyny górnej sytuacja wygląda podobnie: mamy mnóstwo obiecujących sygnałów, ale i wyraźne ostrzeżenia przed przedwczesnym triumfalizmem, który bywa zgubny. To nie jest jeszcze etap, na którym nauka może wydać jeden, rozstrzygający werdykt, lecz raczej pole mozolnej, codziennej optymalizacji parametrów, takich jak intensywność czy dobór grupy docelowej. Technologie rehabilitacyjne często błyszczą w kontrolowanych warunkach laboratoryjnych, by ostatecznie przegrać z szarą prozą adherence, czyli stopnia, w jakim pacjent faktycznie angażuje się w proces i przestrzega zaleceń.

Pacjent nie jest przecież bezwolnym przedłużeniem protokołu badawczego, lecz istotą, która się męczy, nudzi, boi, a czasem po prostu rezygnuje z wysiłku pod wpływem osamotnienia. Właśnie dlatego sztuczna inteligencja w rehabilitacji nie powinna być postrzegana wyłącznie jako zaawansowany sterownik mechanicznego ruchu kończyny czy siłownik egzoszkieletu. Jej realna siła drzemie w zdolności do wykrywania spadku zaangażowania, personalizacji poziomu trudności oraz inteligentnej grywalizacji procesu zdrowienia, która przywraca pacjentowi podmiotowość. Kiedy robotyka łączy się z interfejsami mózg-komputer, system zaczyna reagować nie tylko na fizyczny ruch, ale na samą intencję i wzorce aktywacji neuronalnej. Rehabilitacja przestaje być wtedy serią nudnych wizyt, a staje się środowiskiem ciągłego uczenia się, o ile nie zapomnimy, że celem jest człowiek, a nie sprawność urządzenia.

W tym kontekście nanotechnologia przestaje być egzotycznym dodatkiem do głównej narracji, stając się próbą rozwiązania brutalnego problemu jakości kontaktu elektrody z tkanką mózgową. To nie algorytm jako pierwszy styka się z tajemnicą świadomości, lecz konkretny materiał, który bywa wobec biologii wyjątkowo zdradliwy i agresywny. Odpowiedź zapalna, obrastanie glejem czy wzrost impedancji sprawiają, że nawet najbardziej genialny model matematyczny nie uratuje wadliwego, niestabilnego interfejsu, który traci sygnał po kilku miesiącach. Rozwój elastycznych elektrod grafenowych i hybrydowych nanostruktur to dziś kluczowy kierunek walki o trwałość systemów implantowalnych, bez których neuromodulacja pozostanie jedynie drogim eksperymentem.

Literatura z pogranicza AI i nanomedycyny coraz częściej podkreśla znaczenie uczenia maszynowego w projektowaniu inteligentnych systemów dostarczania leków i optymalizacji struktur materiałowych. Niemniej jednak w tym obszarze sceptycyzm jest wręcz obowiązkiem intelektualnym, chroniącym nas przed myleniem laboratorium z gotową, zatwierdzoną ścieżką kliniczną. Gdy słyszymy o kwantowych znacznikach i symbiozie AI z nanomateriałami, łatwo ulec złudzeniu, że przyszłość już nadeszła i jest dostępna na wyciągnięcie ręki. To byłby błąd, bowiem

większość tych rozwiązań ma obecnie status obiecujących, ale wciąż nieustabilizowanych i ryzykownych prototypów, które muszą przejść przez czyszczenie badań toksykologicznych.

Najuczciwszy opis obecnego stanu rzeczy brzmi następująco: nanotechnologia może poprawić jakość sygnału, a AI może ten sygnał lepiej wykorzystać w procesie terapeutycznym. Jednak droga od tego „może” do powszechnego standardu opieki jest długa, wybrukowana kwestiami produkowalności, kosztów i rygorystycznych zgód regulacyjnych. Każdy, kto te etapy lekceważy, uprawia w istocie metafizykę venture capital przebraną za metodologię nauki, oferując kolorowe wizje zamiast twardych dowodów klinicznych. Prawdziwy postęp wymaga cierpliwości i pokory wobec biologicznej materii, która nie zawsze chce współpracować z naszymi cyfrowymi ambicjami i marzeniami o technologicznej nieśmiertelności.

Szczególne miejsce w tej układance zajmuje integracja BCI z Internetem Rzeczy, co pozwala na rozszerzenie medycyny poza sterylne mury gabinetu lekarskiego do naturalnego środowiska pacjenta. Gdy sygnały neuronalne i czujniki noszone zaczynają pracować w jednej sieci, choroba przestaje być incydem obserwowanym pod mikroskopem, a staje się stale mierzoną dynamiką życia. To przejście do medycyny nawigowanej przez modele, czyli modelu opieki, w którym AI nie tylko klasyfikuje dane, ale współtworzy infrastrukturę decyzji terapeutycznych, zmienia fundamentalnie rolę lekarza. Rehabilitacja staje się procesem ciągłym, zintegrowanym z codziennością, co daje wspaniałą obietnicę skuteczności, ale i rodzi pytania o granice naszej prywatności.

Unia Europejska, działając jako mocarstwo normatywne, czyli podmiot budujący przewagę poprzez rygorystyczne standardy etyczne i prawne, wiąże rozwój AI z budową Europejskiej Przestrzeni Danych dotyczących Zdrowia. Ma to gwarantować Trustworthy AI, czyli koncepcję wiarygodnej sztucznej inteligencji opartej na jakości danych, przejrzystości modelu i jasnej odpowiedzialności prawnej za błędy algorytmu. W perspektywie neurorehabilitacji oznacza to lepszą personalizację, ale i nowe pole konfliktu o prywatność neuronalną oraz zakres automatycznej ingerencji w ludzkie ciało. Transparentność urządzeń ML, czyli stopień komunikowania informacji o przeznaczeniu i logice działania systemów uczenia maszynowego, staje się tu fundamentem bezpieczeństwa. Musimy uważać, by fascynacja skutecznością liczbową nie sprawiła, że zamiast kapłana z kadzidłem będziemy mieli dewelopera z notebookiem decydującego o naszej integralności.

AI w medycynie: od technicznego gadżetu do ustroju zaufania

Kiedy współczesny system medyczny ewoluuje w stronę stale uczącej się sieci, musimy zacząć stawiać pytania wykraczające poza prostą inżynierię wiedzy. Nie chodzi już tylko o to, co taka sieć potrafi wywnioskować z oceanu danych, lecz przede wszystkim o to, czego wiedzieć jej nie wolno i jakich granic nie może przekroczyć, nawet jeśli dysponuje technicznym potencjałem do ich sforsowania. W tym kontekście najsilniejsza teza o konieczności nadzoru pozostaje nienaruszona, szczególnie gdy wchodzimy na najbardziej zaawansowany poziom neurotechnologii. Trustworthy AI, czyli koncepcja wiarygodnej sztucznej inteligencji opartej na jakości danych, przejrzystości modelu i jasnej odpowiedzialności, nie jest jedynie miękką, PR-ową etykietą dla technologii, którą chcielibyśmy po prostu polubić. Stanowi ona twardy, niemal biologiczny warunek cywilizacyjnej dopuszczalności algorytmów w naszym życiu, będąc swoistym paszportem bezpieczeństwa. Im bliżej ludzkiego mózgu operuje system, tym mniej miejsca zostaje na jakąkolwiek epistemiczną bylejakość czy metodologiczne pójście na skróty.

Wyjaśnialność w neurotechnologii nie jest przecież hołdem składanym akademickiej próżności, lecz fundamentalnym narzędziem kontroli każdej klinicznej interwencji. Bezpieczeństwo nie może być postrzegane jako opcjonalny dodatek do skuteczności, ponieważ w medycynie to właśnie ono jest nadrzędnym warunkiem jej zaistnienia. Fairness, rozumiane jako sprawiedliwość algorytmiczna, nie oznacza wyłącznie równego rozkładu błędów statystycznych w tabeli Excela, ale realną równość dostępu do walidacji i reprezentatywnych danych. Obejmuje ono także prawo do kompensacji szkód oraz wymóg, by każda innowacja przynosiła pacjentowi wymierną korzyść kliniczną, a nie tylko zysk producentowi. Prywatność zaś przestaje być w tym ujęciu wyłącznie kwestią ochrony cyfrowej dokumentacji medycznej ukrytej na serwerze. W dobie systemów zamkniętej pętli zaczyna ona dotyczyć samej architektury sygnału neuronalnego, który poprzedza każdy nasz ruch, mowę, afekt, a wreszcie każdą autonomiczną decyzję.

Warto w tym miejscu sformułować diagnozę być może niewygodną dla technoentuzjastów, ale absolutnie konieczną z punktu widzenia etyki lekarskiej. Współczesna medycyna nie potrzebuje kolejnej fali maszynowego prestiżu, budowanego na dziecinnej fascynacji skutecznością liczbową i efektownymi wykresami. Potrzebuje ona raczej solidnego ustroju zaufania, w którym algorytm, elektroda, klinicysta, pacjent oraz regulator tworzą spójny porządek odpowiedzialnej współpracy. Bez takiego fundamentu nawet najbardziej olśniewający postęp technologiczny pozostanie jedynie elegancko opakowaną formą hazardu, w którym stawką jest ludzkie zdrowie. A hazard w obszarze neurologii i psychiatrii nie ma w sobie nic z romantyzmu, gdyż jest po prostu niewybaczalnie

kosztowny dla cudzych synaps, rozbitych rodzin i niepewnej przyszłości pacjentów. To nie jest gra o sumie zerowej, lecz realne ryzyko trwałej degradacji ludzkiej podmiotowości.

Dochodzimy zatem do punktu, w którym spór o obecność AI w medycynie przestaje być techniczną rozmową o narzędziu, a staje się debatą o ustroju odpowiedzialności. Im głębiej systemy diagnostyczne, neuromodulacyjne i rehabilitacyjne wchodzi w samo centrum decyzji klinicznej, tym mniej wolno nam myśleć o nich w kategoriach gadżetu. Musimy zacząć postrzegać je jako krytyczną infrastrukturę publicznego zaufania, analogiczną do sieci energetycznych czy systemów bankowych. Współczesny paradygmat naukowy i regulacyjny jest w tej kwestii wyjątkowo jednoznaczny i nie pozostawia miejsca na interpretacyjną dowolność. Unijny AI Act, który wszedł w życie 1 sierpnia 2024 roku, potraktował oprogramowanie medyczne oparte na algorytmach jako obszar wysokiego ryzyka. Wymusza to na producentach rygorystyczne zarządzanie ryzykiem, dbałość o najwyższą jakość danych treningowych oraz zapewnienie stałego nadzoru człowieka nad maszyną.

Jednocześnie na poziomie unijnym trwa proces doprecyzowywania relacji między poszczególnymi aktami sektorowymi, co pokazuje, że prawo przestało próbować doganiać technologię w desperackim sprincie. Zamiast tego ustawodawca buduje tor przeszkód z rozsądnie rozstawionymi barierami, które mają odsiać rozwiązania niedojrzałe od tych naprawdę bezpiecznych. To bardzo dobra wiadomość dla pacjentów, ponieważ medycyna nie potrzebuje deregulowanego zachwyty, lecz wysokiego progu dowodowego. Równie istotne jest to, że globalne instytucje nie opisują dziś zaufania do AI jako ulotnej cechy psychologicznej, ale jako wymierny efekt systematycznego zarządzania ryzykiem. Amerykański NIST ujmuje AI RMF jako ramę służącą lepszemu panowaniu nad zagrożeniami dla ludzi i całych społeczeństw, łącząc trustworthiness bezpośrednio z kulturą odpowiedzialności.

Światowa Organizacja Zdrowia w swoich wytycznych dla dużych modeli wielomodalnych również podkreśla potrzebę rozbudowanego nadzoru i instytucjonalnej ostrożności. W praktyce oznacza to definitywny koniec infantylnej epoki, w której model uznawano za godny zaufania tylko dlatego, że dobrze prezentował się na konferencji. Dziś wiarygodność musi być rygorystycznie wykazana w całym cyklu życia systemu, od pierwszego zestawu danych, przez walidację kliniczną, aż po monitoring po wdrożeniu. W tym sensie źródłowa teza o konieczności budowania zaufania zasługuje nie tylko na obronę, ale wręcz na radykalne wyostrenie. Trustworthy AI nie jest pojęciem miękkim, lecz ustrojowym, ponieważ oznacza porządek, w którym nie da się oddzielić niezawodności od prawa, a bezpieczeństwa od etyki.

Jest to szczególnie kluczowe w obszarze neurologii i technologii BCI, czyli interfejsów mózg-komputer, gdzie system nie tylko analizuje dane, ale wchodzi w obszar samej cielesności i intencji.

Tu nie wystarczy już pytać, czy dany model statystycznie działa, bo musimy wiedzieć, kto go audytuje i na jakiej populacji był on trenowany. Musimy rozumieć, jak algorytm reaguje na zmianę kontekstu klinicznego i kto posiada uprawnienia, by w razie awarii natychmiast zatrzymać jego działanie. Na tym tle recepcja problemu w środowiskach analitycznych okazuje się zaskakująco przenikliwa, gdyż nie zatrzymuje się na powierzchownej opozycji między entuzjazmem a lękiem. Analiza luki w danych pokazuje dobitnie, że brak reprezentatywnych informacji nie jest neutralnym niedostatkiem metodologicznym, lecz nowoczesną formą bezpieczeństwa i systemowego wykluczenia.

W medycynie takie wykluczenie prowadzi do realnych kosztów zdrowotnych, które uderzają w grupy najsłabiej reprezentowane w bazach danych. Z kolei koncepcja architektury wolności pozytywnej i pojęcie habeas mentem, oznaczające prawo do integralności własnego umysłu, przesuwają dyskusję na poziom głębszy. Przypomina nam ono, że technologie cyfrowe nie działają w próżni, lecz wchodzi w obszar autonomii jednostki i jej zdolności do oporu wobec manipulacji. Połączenie tych dwóch wątków daje niezwykle mocną diagnozę współczesności. Algorytm może być bowiem formalnie nowoczesny, a zarazem ustrojowo regresywny, jeżeli bazuje na zawężonym obrazie człowieka i osłabia jego podmiotowość. To jest punkt, w którym antropologia, ekonomia i prawo zaczynają wreszcie mówić jednym, spójnym głosem.

Ekonomista ująłby tę kwestię w sposób brutalnie prosty: każdy system AI dokonuje redystrybucji kosztów błędów, a kluczowe pytanie brzmi, na kogo te koszty zostaną ostatecznie przerzucone. Jeśli model diagnostyczny jest trenowany na danych, które niedoreprezentują kobiet, osób starszych czy pacjentów z chorobami rzadkimi, to jego wydajność jest subsydiowana przez tych, którzy ponoszą wyższe ryzyko pomyłki. Z ekonomicznego punktu widzenia nie mamy wtedy do czynienia z innowacją neutralną, lecz z mechanizmem przesuwania ryzyka na grupy najsłabsze. Z perspektywy prawnej zbliżamy się wówczas do niebezpiecznego gruntu odpowiedzialności za wadliwy produkt lub dyskryminację pośrednią. Kulturowo zaś powraca stara opowieść o tym, że to, co ogłasza się jako uniwersalne, jest często jedynie lokalnym doświadczeniem silniejszych. Dobre państwo słusznie demaskuje ten mechanizm jako mit neutralności danych, a globalne instytucje, choć używają języka bardziej urzędowego, potwierdzają tę samą, gorzką prawdę.

Konstytucja zaufania: Jak ocalić podmiotowość w erze algorytmów

Dane niskiej jakości to nie tylko problem statystyczny, lecz przede wszystkim paląca kwestia sprawiedliwości społecznej i medycznej. Z tego powodu dzisiejsze prawo regulujące sztuczną inteligencję w medycynie nie powinno być postrzegane jako biurokratyczny wróg innowacji, lecz

jako niezbędny warunek jej cywilizacyjnej dojrzałości. Amerykańska FDA w swoich najnowszych projektach wytycznych kładzie silny nacisk na Total Product Life Cycle, czyli podejście obejmujące zarządzanie ryzykiem na każdym etapie życia produktu, od projektu po rynek. Osobne regulacje dotyczą Predetermined Change Control Plans, czyli z góry określonych planów zmian, które pozwalają systemom na iteracyjne ulepszanie bez utraty gwarancji bezpieczeństwa. To sygnał o ogromnym znaczeniu filozoficznym, sugerujący, że regulator wreszcie przestał traktować algorytm jak statyczny przedmiot. Uznano, że model uczący się jest bytem dynamicznym, wymagającym równie dynamicznej formy odpowiedzialności, co wpisuje się w strategię mocarstwa normatywnego, budującego przewagę poprzez rygorystyczne standardy etyczne. Tu nie da się już dłużej udawać, że jednorazowa certyfikacja urządzenia rozwiązuje problem bezpieczeństwa raz na zawsze.

W obszarach takich jak neurologia czy psychiatria, konsekwencje tego paradygmatycznego przejścia stają się wyjątkowo ostre i namacalne dla każdego praktyka. Systemy BCI, czyli interfejsy mózg-komputer, oraz adaptacyjne DBS, czyli głęboka stymulacja mózgu dostosowująca się do patologicznych oscylacji w czasie rzeczywistym, przestają być tylko narzędziami pomocniczymi. One zmieniają samą strukturę działania klinicznego, wprowadzając medycynę nawigowaną przez modele, w której AI współtworzy infrastrukturę decyzji terapeutycznych, zamiast tylko klasyfikować dane. Klinicysta rzadko kiedy ogląda już surowy wynik, częściej natomiast współpracuje z maszynowo wygenerowanym obrazem sytuacji pacjenta, co zmienia jego rolę na zawsze. To nieuchronnie przesuwaa ciężar epistemiczny z indywidualnego mistrzostwa lekarza na złożony układ człowiek-model-instytucja, wymagając nowej formy profesjonalizmu. W takim świecie wyjaśnialność algorytmu przestaje być kwestią komfortu poznawczego, a staje się warunkiem utrzymania rozproszonej odpowiedzialności w ryzach.

Jeżeli lekarz ma pozostać lekarzem, a nie tylko notariuszem algorytmu, musi posiadać realną możliwość zrozumienia logiki interwencji oraz świadomego przejęcia decyzji. W tym miejscu należy stanowczo zaatakować konformizm części środowiska technologicznego, który karmi nas wyjątkowo szkodliwymi mitami o nieuchronności postępu. Najbardziej niebezpieczne przekonanie głosi, że skoro system jest bardziej złożony od ludzkiego umysłu, to człowiek powinien pokornie pogodzić się z jego nieprzejrzyistością. To nie jest żadna dojrzała teoria wiedzy, lecz zwykła kapitulacja rozumu ubrana w język postępu, na którą nie możemy się zgodzić. W medycynie taki gest jest intelektualnie niedopuszczalny, gdyż stawką jest ludzkie życie, zdrowie oraz nasza fundamentalna integralność kognitywna. Można zaakceptować częściową nieprzezroczystość matematyczną modelu, ale nigdy nie wolno zaakceptować całkowitej nieodpowiedzialności procesu decyzyjnego, który bezpośrednio wpływa na pacjenta.

Można pracować z narzędziem, którego wewnętrznej mechaniki nikt intuicyjnie nie ogarnia, pod warunkiem istnienia procedur walidacji, audytu i monitoringu driftu. Inaczej wracamy do stanu przednaukowego, tyle że zamiast kapłana z kadzidłem mamy dewelopera z notebookiem, co zakrawa na ponury żart. Humor absurdu pisze się tu sam, gdy współczesna wyrocznia siedzi na chmurze obliczeniowej i rzadko wyjaśnia, skąd właściwie czerpie swą rzekomą wiedzę. Wyrocznia z Delf siedziała przynajmniej na trójnogu, podczas gdy dzisiejsza wymaga potężnych serwerowni i gigantycznych nakładów energii, nie oferując w zamian większej jasności. Współczesny paradygmat naukowy zaczyna to na szczęście rozumieć, co widać w tekstach dotyczących neuroradiologii, analizy SPES czy robotyki rehabilitacyjnej. Dominują w nich trzy motywy: entuzjazm wobec personalizacji, nacisk na translację kliniczną oraz żądanie rygorystycznej standaryzacji i walidacji.

Translacja kliniczna to brutalne pytanie o to, czy dane rozwiązanie faktycznie działa poza sterylnym środowiskiem laboratorium i czy przynosi realną korzyść. Nauka dojrzała nie obiecuje już rewolucji na każdym kroku, lecz umie odróżnić teoretyczny potencjał od realnego standardu opieki medycznej. Najbardziej wiarygodne publikacje to te, które z chłodną precyzją wskazują, gdzie kończy się twardy dowód, a zaczyna jedynie pobożna nadzieja. Z perspektywy biznesowej AI w medycynie przetrwa tylko wtedy, gdy wpisze się w realną ekonomię instytucji leczniczej i jej codzienną rutynę. Oznacza to, że technologia musi ograniczać koszty błędu i poprawiać przepływ pracy, zamiast tworzyć nową warstwę kosztownego chaosu organizacyjnego. Wiele rozwiązań upada nie dlatego, że są technicznie złe, lecz dlatego, że ignorują rzeczywisty rytm pracy klinicznej i wymogi interoperacyjności.

Interfejs człowiek-maszyna w medycynie nie może być po prostu efektywny, lecz musi być przede wszystkim poznawczo uczciwy wobec użytkownika. Ma on wspierać decyzję lekarza, a nie mnożyć informacyjny szum, objaśniając ograniczenia zamiast jedynie eksponować pewność predykcji algorytmu. W psychiatrii sprawa jest jeszcze bardziej drażliwa, gdyż systemy AI wchodzą w obszar, gdzie wsparcie może łatwo stać się symulacją relacji. To właśnie tutaj pojawia się ryzyko "fałszywego spomiedzy", czyli iluzji bliskości oferowanej przez chatboty bez żadnej wzajemności czy ludzkiego kosztu obecności. Taka iluzja może chwilowo pomagać, ale długofalowo grozi osłabieniem więzi społecznych i zastąpieniem instytucji opieki tanim ersatzem relacyjności. Musimy bronić tezy, że sztuczna inteligencja ma wspierać człowieka, a nie kolonizować jego samotność pod pozorem nowoczesnej troski.

Integracja AI z diagnostyką, obrazowaniem czy nanotechnologią nie jest tylko kolejnym etapem rozwoju aparatury, lecz próbą zbudowania zupełnie nowego reżimu poznawczego. W świecie, w którym granica między analizą a interwencją staje się coraz cieńsza, musimy zacząć mówić o konstytucji zaufania i prawie do technologii. Trustworthy AI, czyli koncepcja wiarygodnej sztucznej

inteligencji opartej na przejrzystości, staje się twardym warunkiem dopuszczalności tych systemów w naszym życiu. Oznacza to obowiązek korzystania z reprezentatywnych danych oraz nadrzędność nadzoru ludzkiego nad każdą formą automatyzacji procesów decyzyjnych. Skuteczność bez sprawiedliwości jest bowiem tylko wydajniejszą formą nierówności, a precyzja bez wyjaśnialności bywa zaledwie bardziej elegancką postacią ślepoty. Transparentność urządzeń ML, czyli stopień komunikowania informacji o ich logice działania, jest bezpośrednio powiązana z bezpieczeństwem pacjenta i możliwością nadzoru.

Przyszłość medycyny nie rozstrzygnie się ostatecznie w pojedynku między człowiekiem a maszyną, lecz między dwiema wizjami samej technologii i jej roli. Pierwsza z nich traktuje AI jako instrument całkowicie podporządkowany godności pacjenta, rygorowi dowodu naukowego i demokratycznej kontroli instytucjonalnej. Druga wizja postrzega ją jako autonomiczny porządek efektywności, do którego człowiek ma się po prostu bezrefleksyjnie dostosować w imię optymalizacji. Pierwsza z tych dróg buduje cywilizację zaufania, w której technologia służy rozszerzaniu ludzkich możliwości, empatii oraz wolności wyboru. Druga natomiast prowadzi prosto w stronę nowoczesnego feudalizmu danych, gdzie władza algorytmów jest nieprzejrzysta, a odpowiedzialność za błędy całkowicie rozmyta. Musimy zatem z całą mocą odrzucić metafizykę venture capital przebraną za metodologię nauki, by ocalić to, co w medycynie najcenniejsze. Wybór, przed którym stoimy, zdefiniuje nie tylko kształt przyszłych szpitali, ale przede wszystkim naszą podmiotowość w erze wszechobecnych algorytmów.

Europa między mocarstwem normatywnym a wyścigiem technologicznym

Jeżeli mamy zachować powagę jako naukowcy, lekarze czy obywatele, konkluzja może być tylko jedna. Nie wolno nam oddać autonomii poznawczej i biologicznej w ręce systemów, których inteligencja paruje przy pierwszym trudnym pytaniu. Prawdziwie wiarygodna AI, czyli koncepcja oparta na jakości danych i jasnej odpowiedzialności, zaczyna się tam, gdzie technologia godzi się na rozliczalność. Jest to twardy warunek cywilizacyjnej dopuszczalności tych narzędzi w obszarach krytycznych dla życia i zdrowia. Bez tego mamy do czynienia ze zjawiskiem, które można określić jako kapitulacja rozumu ubrana w język postępu. Prawdziwe zaufanie nie jest wszak marketingowym sloganem, lecz rygorystycznym żądaniem przejrzystości wobec algorytmicznych wyroków.

Europa opowiada dziś o sztucznej inteligencji tak, jakby opisywała własne przyszłe imperium, choć robi to z pozycji mocarstwa normatywnego. Ta strategia, polegająca na budowaniu przewagi

poprzez rygorystyczne standardy etyczne, zderza się jednak z brutalną rzeczywistością technologiczną. Istnieje bowiem zasadnicza sprzeczność naszej epoki, w której Bruksela chce regulować technologię o obcym kapitale i infrastrukturze. AI Act, który wszedł w życie 1 sierpnia 2024 roku, zaprojektowano jako system budowania bezpieczeństwa i nadzoru człowieka. Komisja przyznaje jednak otwarcie, że pełna stosowalność tych przepisów jest procesem rozłożonym na długie lata. To ambitna próba cywilizowania cyfrowego świata, którego rdzeń obliczeniowy oraz kapitałowy pozostaje wciąż poza realną europejską kontrolą.

Skoro Europa nie kontroluje źródeł mocy, próbuje przynajmniej kontrolować warunki wejścia do własnej przestrzeni prawnej. To pokazuje istotę europejskiego projektu, w którym brak własnych procesorów nadrabia się precyzją kodeksów i regulacji. W praktyce pytanie nie brzmi już, czy Europa chce sztucznej inteligencji, gdyż tutaj panuje rzadka polityczna zgoda. Prawdziwy spór dotyczy tego, czyje prawa mają pierwszeństwo i kto ostatecznie zapłaci rachunek za błędy systemów. Musimy wiedzieć, kto poniesie odpowiedzialność, gdy algorytm skrzywdzi pacjenta lub zdestabilizuje strukturę społeczną. To nie są wszak jedynie błahe usterki techniczne, lecz fundamentalne pytania o granice naszej zbiorowej suwerenności.

Realne różnice zdań zaczynają się wewnątrz Parlamentu Europejskiego, gdzie wybory z 2024 roku utrwaliły skomplikowaną mozaikę interesów. Największą grupą pozostała chadecka EPP, a za nią uplasowali się socjaldemokraci oraz rosnące w siłę nurty suwerenistyczne. Przyszłość unijnej AI nie będzie zatem dziełem jednej doktryny, lecz efektem nieustannego przeciągania liny między centrum a lewicą. EPP mówi o odpowiedzialności, ale myśli przede wszystkim w kategoriach siły gospodarczej i technologicznego bezpieczeństwa. W ich dokumentach powraca postulat budowy suwerennej i konkurencyjnej Unii, zdolnej do walki o rynkowy prymat. Chcą oni sprawnej implementacji AI Act, która nie zdusi innowacji nadmiarem biurokratycznych procedur.

Takie podejście jest charakterystyczne dla europejskiej chadecji, która traktuje przejrzystość modeli przede wszystkim jako warunek stabilnego rynku. Nie odrzucają oni idei wiarygodnej technologii, lecz widzą w niej narzędzie legitymizacji przemysłowego wzrostu i ekspansji. Pierwsze pytanie chadeków nie brzmi, jak ograniczyć władzę algorytmu, lecz jak zbudować ekosystem, którego nie zmiażdżą giganci. Dla EPP prawa obywatela są istotne, niemniej równie ważne pozostają zdolności inwestycyjne i strategiczna autonomia kontynentu. Dążą oni do stworzenia warunków, w których europejskie firmy będą mogły konkurować bez paraliżującego strachu przed regulacjami. Jest to wizja postępu, w której bilans zysków i strat waży tyle samo co etyka.

Socjaldemokraci z grupy S&D ustawiają akcenty zupełnie inaczej, widząc w AI przede wszystkim problem władzy nad człowiekiem. W ich narracji AI Act ma być tarczą chroniącą pracowników i obywateli przed arbitralnością zautomatyzowanych systemów. Podkreślają oni konieczność zakazu

rozpoznawania emocji w szkołach oraz sprzeciwiają się przewidywaniu przestępczości na podstawie profili. Dla lewicy wyjaśnialność modelu nie jest marketingiem, lecz warunkiem kontroli nad pracodawcą i administracją publiczną. Gdy centrum pyta o europejską siłę, socjaldemokracja pyta o europejską osłonę i godność jednostki w cyfrowym świecie. Uważają oni, że odpowiedzialność nie może rozplątać się w chmurze anonimowych kontraktorów, przy czym kładą nacisk na pełną przejrzystość łańcucha dostaw.

Ostatecznie europejska droga do sztucznej inteligencji pozostaje polem bitwy między marzeniem o wolnym rynku a lękiem przed algorytmiczną opresją. Ta wewnętrzna sprzeczność nie jest słabością, lecz wyrazem unikalnego DNA politycznego, które nie chce wybierać między innowacją a moralnością. Jesteśmy świadkami walki o definicję człowieczeństwa w czasach, gdy zamiast kapłana z kadzidłem mamy dewelopera z notebookiem. Czy ta ambicja normatywna wystarczy, by zrównoważyć brak własnej mocy obliczeniowej, pozostaje najważniejszym pytaniem tej dekady. Odpowiedź na nie określi, czy Europa stanie się muzeum mądrych regulacji, czy żywym laboratorium nowej technologii.

W świecie, w którym algorytmy coraz głębiej wnikają w naszą cielesność, pytanie o wiarygodność AI staje się pytaniem o granice ludzkiej wolności. Czy w pogoni za doskonałością medyczną nie ryzykujemy zamiany autentycznego leczenia na zautomatyzowaną iluzję troski? Musimy zdecydować, czy chcemy pozostać podmiotami własnego życia, czy jedynie danymi wejściowymi w systemie, którego logiki nikt z nas już nie potrafi wyjaśnić.

Inspiracje

[1]



"Artificial Intelligence and Brain Computer Interfaces in Healthcare" Chandra P. Sharma

Chyren Publication | ISBN: 9789371433471

Dr **Chandra P. Sharma** jest wybitnym naukowcem i ekspertem w dziedzinie technologii biomedycznej, specjalizującym się w biomateriałach. W latach 1980–2014 pełnił funkcję starszego naukowca i kierownika Wydziału Technologii Biomedycznych w Instytucie Nauk Medycznych i Technologii Sree Chitra Tirunal (SCTIMST) w Trivandrum. Obecnie jest profesorem nadzwyczajnym w Manipal College of Pharmaceutical Sciences na Uniwersytecie Manipal oraz profesorem honorowym emerytowanym w College of Biomedical Engineering & Applied Sciences na Uniwersytecie Purbanchal w Nepalu. Z wykształcenia fizyk ciała stałego w IIT Delhi, kontynuował specjalizację z biomateriałów na Uniwersytecie Utah i Uniwersytecie w Liverpoolu. Dr Sharma jest założycielem Towarzystwa Biomateriałów i Sztucznych Organów w Indiach (SBAOI) oraz Towarzystwa Inżynierii Tkankowej i Medycyny Regeneracyjnej w Indiach (STERMI). Jest autorem i redaktorem wielu wpływowych książek i prac naukowych na temat biomateriałów, inżynierii tkankowej i dostarczania leków.

Dr. Chandra P. Sharma is a distinguished scientist and expert in biomedical technology, specializing in biomaterials. He served as Senior Scientist G and Head of the Biomedical Technology Wing at the Sree Chitra Tirunal Institute for Medical Sciences & Technology (SCTIMST), Trivandrum, from 1980 to 2014. He is currently an Adjunct Professor at the Manipal College of Pharmaceutical Sciences, Manipal University, and an Honorary Emeritus Professor at the College of Biomedical Engineering & Applied Sciences, Purbanchal University, Nepal. Trained as a solid-state physicist at IIT Delhi, he pursued further specialization in biomaterials at the University of Utah and the University of Liverpool. Dr. Sharma is the founder of the Society for Biomaterials and Artificial Organs, India (SBAOI) and the Society for Tissue Engineering and Regenerative Medicine, India (STERMI). He has authored and edited numerous influential books and research papers on biomaterials, tissue engineering, and drug delivery.

Sree Chitra Tirunal Institute for Medical Sciences & Technology

Literatura uzupełniająca

Książki

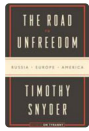
Literatura pokrewna (Google Scholar):



[1] "On Tyranny: Twenty Lessons from the Twentieth Century"

Timothy Snyder

2017 | Blurb | ISBN: 9781388091316



[2] "The Road to Unfreedom: Russia, Europe, America"

Timothy Snyder

2018 | Tim Duggan Books | ISBN: 9780525575405



[3] "Critical Theory of AI"

Simon Lindgren

2024 | John Wiley & Sons | ISBN: 9781509555789

[4] "The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power"

Shoshana Zuboff

2019 | PublicAffairs | ISBN: 9798507331451

Artykuły naukowe

Autorzy przywoływani:

[1] "A Value for n-person Games"

Lloyd S. Shapley

1953 | Contributions to the Theory of Games

[Google Scholar](#)

[2] "College Admissions and the Stability of Marriage"

David Gale, Lloyd S. Shapley

1962 | The American Mathematical Monthly

[Google Scholar](#)

[3] "Biases in artificial intelligence (AI) systems pose a range of ethical issues"

Hofmann, B.

2025 | Frontiers in Digital Health

[Google Scholar](#)

[4] "The Politics of Inevitability and the Politics of Eternity"

Timothy Snyder

2018 | Journal of Democracy | DOI: 10.1353/jod.2018.0040

[Google Scholar](#)

[5] "The Ethics of Algorithms: Mapping the Debate"

Brent Daniel Mittelstadt, Patrick Allo, Mariarosaria Taddeo, Sandra Wachter, Luciano Floridi

2016 | Big Data & Society

[Google Scholar](#)



Tekst opracowany z udziałem sztucznej inteligencji.

Redakcja i kontrola merytoryczna: Fundacja Dobre Państwo.

APA ONE Publishing System

Fundacja Dobre Państwo © 2026

Article ID: doa651a4-57de-44c2-b392-f062525506f8