

Architektura niepewności: Nowy krajobraz uczenia maszynowego



AUTOR

Fundacja Dobre Państwo

STRESZCZENIE / ABSTRACT

Artykuł analizuje ewolucję uczenia maszynowego, które z niszowego rzemiosła staje się fundamentem cyfrowego przemysłu. Autor wskazuje, że współczesne AI to nie tylko algorytmy, ale przede wszystkim złożona infrastruktura MLOps, obejmująca wersjonowanie danych, monitorowanie dryfu modeli oraz rygorystyczną automatyzację. Kluczową rolę w tym procesie odgrywa ekosystem TensorFlow, który poprzez mechanizmy takie jak automatyczne różniczkowanie czy grafy obliczeniowe, standaryzuje produkcję decyzji statystycznych. Tekst odczarowuje technologiczną teologię, sprowadzając sztuczną inteligencję do poziomu rzetelnej inżynierii i ekonomii politycznej. W dobie regulacji i rosnących kosztów obliczeniowych, zrozumienie modelu jako urządzenia produkcyjnego staje się niezbędne dla efektywnego wdrażania innowacji w nowoczesnym państwie.

This article analyzes the evolution of machine learning, which has evolved from a niche craft into the foundation of digital industry. The author points out that modern AI is not just algorithms but, above all, a complex MLOps infrastructure, encompassing data versioning, model drift monitoring, and rigorous automation. The TensorFlow ecosystem plays a key role in this process, standardizing the production of statistical decisions through mechanisms such as automatic differentiation and computational graphs. The text demystifies technological theology, reducing AI to the level of sound engineering and political economy. In an era of regulation and rising computational costs, understanding models as production devices is essential for the effective implementation of innovations in the modern state.

Współczesne uczenie maszynowe ostatecznie porzuciło status rzemieślniczego eksperymentu, przekształcając się w złożony system instytucjonalny późnego kapitalizmu cyfrowego. Narzędzia takie jak TensorFlow przestały być jedynie

bibliotekami programistycznymi, stając się fundamentem nowej architektury władzy, w której statystyka pełni rolę tkanki decyzyjnej. Niniejszy artykuł dekonstruuje mit sztucznej inteligencji jako magicznego objawienia, dowodząc, że jej rzeczywista siła tkwi w rygorystycznej inżynierii MLOps, skalowalności danych i głębokim dopasowaniu architektur sieci do struktury problemów świata rzeczywistego. Autorzy przekonują, że w obliczu rosnącej roli modeli bazowych i regulacji typu AI Act, kluczowym wyzwaniem nie jest już sama moc obliczeniowa, lecz zdolność do osadzenia algorytmów w rzetelnych ramach prawnych, etycznych i operacyjnych. Artykuł postuluje przejście od technologicznego optymizmu ku dojrzałemu paradygmatowi poznawczemu, w którym model nie jest autonomiczną wyrocznią, lecz odpowiedzialnym komponentem infrastruktury instytucjonalnej. Zrozumienie tego przejścia jest niezbędne dla każdego, kto chce wyjść poza rolę operatora interfejsu i zacząć świadomie kształtować cyfrowy porządek przyszłości.

SŁOWA KLUCZOWE

uczenie maszynowe

TensorFlow

MLOps

PyTorch

Reinforcement Learning

CNN

Transformery

GNN

automatyczne różniczkowanie

GradientTape

dryf modelu

inżynieria cech

LiteRT

infrastruktura obliczeniowa

tensory

Od rzemiosła do fabryki: Nowy ład technologiczny uczenia maszynowego

Współczesny krajobraz uczenia maszynowego przestał przypominać zaciszny warsztat rzemieślnika, w którym z kilku prostych narzędzi mozolnie struga się pojedynczy model do jednego, konkretnego zadania. Dziś znajdujemy się raczej w samym sercu instytucjonalnego porządku późnego kapitalizmu cyfrowego, gdzie dane, infrastruktura obliczeniowa, architektury i reżimy wdrożeniowe tworzą gęstą sieć naczyń połączonych. Kto postrzega TensorFlow wyłącznie jako bibliotekę programistyczną, ten wprowadzie mówić prawdę techniczną, ale całkowicie rozmyja się z prawdą ustrojową tego narzędzia. Jest ono bowiem jednocześnie aparatem obliczeniowym, językiem organizacji pracy i formą inżynierskiej dyscypliny, która zmienia status wiedzy statystycznej w strukturach władzy. Statystyka przestaje być jedynie zapleczem dla ludzkich decyzji,

stając się samą tkanką, z której te decyzje są wycinane. To nie jest metafora, lecz analogia strukturalna.

Oficjalna dokumentacja TensorFlow definiuje ten ekosystem jako środowisko ułatwiające budowę modeli zdolnych do pracy w każdych warunkach, od chmury po urządzenia brzegowe. W tej pozornie suchej deklaracji kryje się potężna ambicja ustrojowa, ponieważ model nie ma być tylko matematycznie poprawny, lecz przede wszystkim przenośny i skalowalny. Na portalu dobrepanstwo.org wybrzmiewa teza, która trafnie porządkuje to rozpoznanie, odzierając sztuczną inteligencję z szat widowiska i pustych obietnic. AI zostaje tam sprowadzone do praktyki MLOps, czyli zestawu procedur obejmujących wersjonowanie danych, automatyzację treningu oraz ciągle monitorowanie dryfu modelu. Ten punkt jest fundamentalny, ponieważ przesuwając ciężar dyskusji z magii na rzetelne rzemiosło przemysłowe. Nie jest to detal implementacyjny. To fundament nowego ładu.

Współczesna sztuczna inteligencja nie jest już sztuką trenowania jednego, błyskotliwego klasyfikatora, lecz raczej inżynierią utrzymania przewidywalnego systemu poznawczego w świecie pełnym chaosu. Gdy Fundacja Dobre Państwo przypomina, że wdrożenie modelu to proces inżynieryjny, a nie akt metafizycznego objawienia, uderza w sam środek współczesnego sporu o technologię. Wokół algorytmów narosła bowiem swoista teologia technologiczna, którą należy jak najszybciej odczarować i zastąpić trzeźwą ekonomią polityczną modeli. Każdy model generuje przecież realne koszty, pożera energię, podlega degradacji i musi być osadzony w gęstym kontekście regulacyjnym. Nie jest on zatem objawieniem, lecz urządzeniem produkcyjnym, które wymaga dokumentacji i nadzoru.

W tym kontekście TensorFlow jawi się jako figura niemal emblematyczna, łącząca abstrakcyjną teorię matematyczną z brutalnymi realiami przemysłowego wdrożenia. Jego fundamentem są tensory, czyli wielowymiarowe struktury danych, na których wykonuje się operacje numeryczne stanowiące kręgosłup sieci neuronowych. Choć etymologia nazwy bywa tłumaczona w sposób nazbyt szkolny, jako prosty przepływ tablic przez warstwy, jej rzeczywisty sens sięga znacznie głębiej. TensorFlow organizuje przepływ obliczeń w taki sposób, aby człowiek mógł operować na poziomie wysokopoziomowych zadań, podczas gdy system przejmuje ciężar optymalizacji. To sprawia, że inżynier może skupić się na architekturze, nie martwiąc się o detale różniczkowania czy dystrybucji pracy.

Warto zwrócić uwagę na ewolucję tego narzędzia, szczególnie na wprowadzenie trybu eager execution, który pozwala na wykonywanie instrukcji na bieżąco i w sposób intuicyjny. Z kolei funkcja `tf.function` potrafi przekształcić kod Pythona w niezależne grafy przepływu danych, co drastycznie poprawia wydajność i przenośność całego systemu. Nie jest to jedynie detal

implementacyjny, lecz miniatura całej filozofii, która przyświeca współczesnemu uczeniu maszynowemu. Najpierw stawiamy na zrozumiałość i swobodny eksperyment, by w kolejnym kroku przejść do surowej kompilacji i standaryzacji przemysłowej. W historii gospodarczej nowoczesności wielokrotnie obserwowaliśmy ten sam ruch od chaosu do modułowości. Tak budowano nowoczesność.

Podobny proces standaryzacji przeszedł niegdyś kontener morski, rachunkowość czy systemy zarządzania jakością, które zamieniły lokalne praktyki w globalne, interoperacyjne procedury. Tak samo dzieje się dzisiaj z uczeniem maszynowym, które z ezoterycznej sztuki modelowania przeistacza się w dojrzały ład technologiczny. W tym nowym porządku model staje się zaledwie jednym z wielu komponentów większego układu służącego do masowej produkcji decyzji. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki, tracąc z oczu szerszy kontekst ekonomiczny. Jednym z kluczowych silników tej zmiany jest automatyczne różniczkowanie, które zdejmuje z barków badaczy ciężar ręcznych obliczeń.

Oficjalny przewodnik po bibliotece wyjaśnia, że mechanizm `tf.GradientTape` rejestruje operacje na zmiennych, by następnie wyliczyć gradienty metodą odwrotną. W praktyce oznacza to pełną mechanizację procesu, który przez dekady stanowił główną barierę dla rozwoju głębokich sieci neuronowych. Gdy liczba parametrów modelu rośnie do miliardów, ręczne konstruowanie pochodnych staje się nie tylko trudne, ale wręcz intelektualnie i operacyjnie absurdalne. `GradientTape` nie jest więc tylko wygodnym interfejsem, lecz machiną umożliwiającą istnienie współczesnego treningu głębokiego. Dokumentacja słusznie jednak ostrzega, że rejestrowanie zbędnych operacji bywa kosztowne, co stanowi cenną lekcję pokory.

W głębokim uczeniu nie wszystko, co poznawczo eleganckie, jest ekonomicznie rozsądne, ponieważ matematyka zawsze pracuje tu pod presją rachunku zasobów. Problem treningu sieci nie sprowadza się zatem wyłącznie do doboru pięknej architektury, lecz jest zagadnieniem zarządzania przepływem sygnału i stabilnością. Należy dbać o jakość danych wejściowych, harmonogram współczynnika uczenia oraz spójność całego środowiska wdrożeniowego, w którym model ma ostatecznie pracować. Twoja notatka trafnie porządkuje te zagadnienia wokół kilku osi, które stanowią kręgosłup dzisiejszego paradygmatu naukowego i inżynierskiego. Architektura nie jest już traktowana jako abstrakcyjna forma, lecz jako bezpośrednia odpowiedź na specyficzny typ danych.

Na portalu dobrepanstwo.org ujęto to z lapidarną trafnością: CNN dla obrazów, Transformery dla tekstu, a GNN dla danych relacyjnych. Ta zasada nie jest dogmatem, lecz heurystyką o wysokiej mocy praktycznej, która sprawia, że architektura staje się częścią epistemologii danych. Mówi nam ona nie tylko, jak liczyć, ale przede wszystkim, jak rozumieć wewnętrzną strukturę badanego zjawiska. Jednocześnie inżynieria cech, czyli proces przygotowania atrybutów danych, wcale nie

została unieważniona przez nadejście głębokiego uczenia. Zmienił się jedynie poziom jej działania, gdyż część tej pracy została przeniesiona do wnętrza samej architektury modelu.

Materiały Fundacji Dobre Państwo wprost stwierdzają, że jakość każdego modelu zaczyna się od starannego przygotowania danych i inżynierii cech. Jest to sąd całkowicie zgodny z dojrzałą praktyką MLOps, która uczy, że nawet największy model nie naprawi błędów popełnionych na starcie. Zły sygnał nadal pozostaje złym sygnałem, choćby przechodził przez najdroższy klaster GPU i najbardziej wyrafinowane warstwy neuronowe. Tu objawia się stara prawda ekonomii informacji: wysokie nakłady kapitałowe nigdy nie zrekompensują złej struktury bodźców wejściowych. Dane nie są więc "nową ropą", lecz są raczej nową glebą, która wymaga odpowiedniej uprawy.

Współczesny paradygmat naukowy nie uznaje już ostrego podziału między światem modeli klasycznych a królestwem głębokiego uczenia. Dobrze wykształcony praktyk musi sprawnie poruszać się po całym spektrum dostępnych narzędzi, od regresji liniowej po modele dyfuzyjne. Regresja logistyczna, lasy losowe, boosting czy autoenkodery nie tworzą szeregu zastępstw, lecz bogaty repertuar, z którego należy czerpać z rozważą. O wyborze konkretnej metody nie powinna decydować moda ani prestiż konferencyjny, lecz natura problemu, koszt błędu i wymogi wyjaśnialności. Twój materiał źródłowy słusznie łączy te wszystkie porządki w jeden spójny ciąg rozwoju, co jest dowodem na współczesną dojrzałość dyscypliny.

Poza hegemonią: Architektura jako kontrakt i optymalizacja

Nie powinniśmy fetyszyzować konkretnej architektury, lecz dążyć do zrozumienia całego warsztatu poznawczego, jakim dysponuje współczesne uczenie maszynowe. Warto w tym miejscu doprecyzować pewien istotny element dzisiejszej recepcji TensorFlow, czyli wszechstronnego ekosystemu służącego do projektowania i wdrażania modeli, który przez lata uchodził za nienaruszalny fundament branży. W badaniach akademickich oraz w dynamicznej kulturze eksperymentu faktyczną dominację zdobył w ostatnich latach PyTorch, uznawany powszechnie za środowisko bardziej naturalne dla szybkiego prototypowania i swobodnej pracy badawczej. Jednocześnie TensorFlow zachował niezwykle mocną pozycję w obszarach stricte wdrożeniowych, ekosystemowych oraz w rozwiązaniach dedykowanych dla urządzeń brzegowych. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki. Nie jest zatem prawdą, że to narzędzie przestało być ważne w obliczu nowszych rozwiązań.

Prawdą jest raczej to, że dawna hegemonia jednego rozwiązania została zastąpiona przez dojrzały pluralizm frameworków, a punkt ciężkości przewagi konkurencyjnej przesunął się z samego etapu treningu modelu w stronę zarządzania pełnym cyklem życia systemu. W tym sensie źródłowa obrona TensorFlow jako narzędzia potężnego, dojrzałego i wielowarstwowego pozostaje nadal w pełni zasadna, o ile nie będziemy jej mylić z twierdzeniem, że jest to dziś jedyny centralny język świata sztucznej inteligencji. Nawet agencja Reuters opisywała pod koniec 2025 roku intensywne wysiłki Google zmierzające do zapewnienia lepszej kompatybilności jednostek TPU z biblioteką PyTorch, co pośrednio potwierdza, że oś władzy rynkowej przesunęła się w stronę środowisk preferowanych przez deweloperów. Z drugiej strony oficjalne zasoby Google dowodzą, że linia wdrożeniowa TensorFlow żyje i ewoluuje, czego najlepszym przykładem jest LiteRT, budowany na fundamencie TensorFlow Lite i pozycjonowany jako następna generacja wdrożeń typu on-device. To rozpoznanie ma znaczenie znacznie większe niż tylko czysto techniczne.

Pokazuje ono bowiem wyraźnie, że paradygmat naukowy wokół uczenia maszynowego wszedł w swoją fazę posthegemoniczną, w której stare podziały tracą na znaczeniu. Nie trwa już jałowy spór o to, czy głębokie uczenie w ogóle działa, ponieważ ta kwestia została ostatecznie rozstrzygnięta przez brutalną praktykę przemysłową i sukcesy rynkowe. Trwa natomiast znacznie ciekawszy spór o to, które stosy technologiczne najlepiej równoważą pięć kluczowych wymiarów naraz: wydajność treningu, wygodę badawczą, koszt wdrożenia, audytowalność oraz zgodność regulacyjną. To właśnie na tym styku inżynierii i zarządzania znajduje się dzisiaj prawdziwe pole walki o dominację w sektorze AI. W ekonomii nazwalibyśmy to zjawisko klasycznym problemem wielokryterialnej optymalizacji przy silnych ograniczeniach instytucjonalnych. W prawie powiedzielibyśmy, że nie liczy się tylko sama zdolność do działania, ale przede wszystkim zdolność do działania w ściśle określonych granicach odpowiedzialności.

W antropologii techniki dodalibyśmy zapewne, że narzędzia nigdy nie są neutralne, ponieważ aktywnie uczą swoich użytkowników określonych nawyków poznawczych oraz specyficznych form organizacji pracy. Kiedy zatem dyskutujemy o architekturze sieci, nie wolno nam ograniczać się do szkolnej genealogii prowadzącej od prostego perceptronu do wyrafinowanego transformera. Owszem, w ujęciu historycznym warto pamiętać, że perceptron Rosenblatt, czyli klasyczna architektura oparta na kaskadowym przetwarzaniu reprezentacji, stworzył naszą wyobraźnię, a problem XOR przez długi czas działał jak zimny prysznic dla nadmiernych entuzjastów. Warto też pamiętać o przełomie, jakim było wprowadzenie algorytmu backpropagation, oraz o tym, jak nieliniowe funkcje aktywacji przywróciły sens budowaniu głębokich struktur. Jednak dzisiaj znajomość samej tej opowieści to zdecydowanie za mało, by zrozumieć dynamikę zmian.

Trzeba zrozumieć, że współczesna architektura jest w istocie kontraktem zawartym między własnościami danych a specyficznymi własnościami świata organizacyjnego, w którym dany model będzie funkcjonował. Sieci CNN nie zwyciężyły wcale dlatego, że są wyjątkowo piękne matematycznie, lecz dlatego, że ich lokalność, współdzielenie wag i hierarchia cech okazały się ekonomicznie i obliczeniowo dopasowane do natury obrazu. Transformery nie zdobyły świata dlatego, że ktoś wymyślił nośny slogan o mechanizmie uwagi, lecz dlatego, że równoległość i skalowalność ich działania okazały się bezkonkurencyjne wobec wymagań nowoczesnej infrastruktury chmurowej. Źródła branżowe słusznie akcentują właśnie to dopasowanie architektury do typu danych, a nie do chwilowej mody ideologicznej czy estetycznej. Reinforcement learning, czyli uczenie przez wzmacnianie, zajmuje w tym krajobrazie pozycję osobną i wyjątkową.

Nie dzieje się tak dlatego, że ten model jest bardziej magiczny od innych, lecz dlatego, że w zupełnie inny sposób organizuje on fundamentalny problem wiedzy. W uczeniu nadzorowanym, paradygmacie opartym na modelowaniu relacji między danymi a etykietami, otrzymujemy gotowe wzorce. W uczeniu nienadzorowanym szukamy ukrytej struktury w chaosie informacji. Natomiast w RL agent działa w dynamicznym środowisku, obserwuje skutki własnych akcji i uczy się polityki maksymalizującej skumulowaną nagrodę. Z epistemologicznego punktu widzenia jest to fascynujące przejście od logiki reprezentacji do logiki interwencji. Model nie ma już tylko wiernie odtworzyć rozkładu obserwacji, ale ma nauczyć się skutecznego działania w świecie, w którym jego własne ruchy realnie zmieniają stan rzeczy. To właśnie dlatego ta dziedzina tak mocno fascynuje badaczy teorii gier, robotyki i systemów autonomicznych.

Jednocześnie to właśnie tutaj najłatwiej ulec niebezpiecznemu złudzeniu, że wreszcie otrzymaliśmy ogólną i ostateczną teorię inteligencji. Tymczasem należy zachować zdrowy, analityczny sceptycyzm wobec podobnych deklaracji. Sukcesy uczenia przez wzmacnianie są bez wątpienia imponujące, ale zwykle opierają się na bardzo starannie zdefiniowanych środowiskach, sztucznych strukturach nagrody i niezwykle kosztownych procedurach eksploracji. Między spektakularnym zwycięstwem w zamkniętej grze a stabilnym działaniem w nieprzewidywalnym, otwartym świecie istnieje głęboka przepaść. Propaganda technologiczna uwielbia zasypywać ją sloganami, podczas gdy rzetelna nauka każe ją precyzyjnie mierzyć. Nagroda w systemach RL jest bowiem odpowiednikiem bodźca regulacyjnego w polityce publicznej. Jeśli źle zaprojektujesz funkcję nagrody, agent będzie optymalizował nie dobro wspólne, lecz jedynie wskaźnik zastępczy.

Dokładnie w ten sam sposób instytucje publiczne i wielkie korporacje wypaczają swoje zachowania, gdy źle dobrane wskaźniki KPI zastępują im cele rzeczywiste. To nie jest metafora, lecz analogia strukturalna, która rzuca światło na głębsze problemy sterowania. Uczenie przez wzmacnianie uczy nas więc nie tylko programowania inteligentnych agentów, ale przede wszystkim pokory wobec

każdego systemu, który zakłada, że wskaźnik jest tym samym co realna wartość. Nie jest. Nigdy nie był. Równie istotnym, choć często pomijanym problemem, jest kwestia skalowania danych, która w profesjonalnym ujęciu nie jest jedynie kosmetycznym zabiegiem, ale warunkiem koniecznym sensownej optymalizacji. Dla wielu algorytmów klasycznych rozbieżne skale poszczególnych cech deformują geometrię przestrzeni i wypaczają przebieg spadku gradientu. W rezultacie model nie tyle uczy się źle, co porusza się po fatalnie ukształtowanej powierzchni kosztu. Standaryzacja nie jest więc porządkowaniem dla estetyki, lecz fundamentalną korektą geometrii problemu.

Koniec ery naiwności: AI jako system instytucjonalny

W głębokim uczeniu sprawa przybiera nową postać, ponieważ preprocessing danych zewnętrznych nierozdzielnie łączy się z normalizacją aktywacji wewnątrz samej sieci neuronowej. Oficjalna dokumentacja TensorFlow, czyli wszechstronnego ekosystemu służącego do projektowania i wdrażania modeli, podkreśla, że API `tf.data.Dataset` wspiera opisowe i wydajne potoki wejściowe. Iteracja odbywa się tam strumieniowo, bez konieczności karkołomnego mieszczącego pełnego zbioru w pamięci operacyjnej urządzenia. To pokazuje, że skalowanie danych ma dziś podwójny sens: statystyczny oraz czysto infrastrukturalny. Dane muszą zostać przekształcone do właściwej skali cech, ale też obsłużone w takiej architekturze przepływu, która po prostu nie udusi systemu wejściem. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki.

Profesjonalny projekt typu end-to-end zaczyna się znacznie wcześniej niż w chwili optymistycznego wywołania funkcji `model.fit()`. Cały proces bierze swój początek w decyzji o celu biznesowym, definicji zmiennej docelowej oraz rygorystycznym sposobie poboru danych i konstrukcji zbioru testowego. To jeden z najtrwalszych wkładów Auréliena Gérona do edukacji w dziedzinie uczenia nadzorowanego, czyli paradygmatu opartego na modelowaniu relacji między obserwowanymi danymi a konkretnymi etykietami. Należy widzieć cały proces, a nie tylko jego widowiskowy, środkowy fragment w postaci treningu algorytmu. Bez właściwej metody walidacji i przemyślanej funkcji kosztu model pozostaje jedynie matematyczną abstrakcją bez praktycznego znaczenia.

Materiał Fundacji Dobre Państwo mocno potwierdza tę intuicję, dodając do niej niezbędną warstwę wdrożeniową oraz regulacyjną, która wykracza poza czysty inżynierski optymizm. Model ma być nie tylko skuteczny w swoich predykcjach, ale też monitorowany, wyjaśnialny i osadzony w realnej odpowiedzialności instytucjonalnej. Musimy go oceniać pod kątem kosztów finansowych oraz energetycznych, co w dobie kryzysu klimatycznego przestaje być jedynie kwestią optymalizacji budżetu. To już nie jest klasyczny podręcznik uczenia maszynowego, lecz opis nowego standardu

cywilizacyjnego dla systemów predykcyjnych. Współczesny paradygmat naukowy wymaga od nas porzucenia technicznej naiwności na rzecz systemowego zrozumienia technologii.

Nie ma dobrego modelu bez rzetelnego sformułowania problemu, tak jak nie ma rzetelnego treningu bez zrozumienia geometrii danych wejściowych. Nie istnieje dziś skuteczne wdrożenie bez MLOps, czyli zbioru praktyk przekształcających czystą technologię w bezpieczny, przewidywalny i zgodny z prawem proces instytucjonalny. Nie ma również odpowiedzialnej sztucznej inteligencji bez rzetelnej dokumentacji, ciągłego monitoringu oraz rozpoznania ryzyka prawnego. Zatem nie ma już większego sensu uczyć TensorFlow jako samotnej biblioteki, a uczenia maszynowego jako zbioru odizolowanych od świata algorytmów. To nie jest detal implementacyjny. To fundament nowej rzeczywistości.

Należy uczyć jednego, spójnego systemu praktyk, w którym statystyka, informatyka, ekonomia, prawo i antropologia organizacji spotykają się w jednym punkcie. Tym punktem jest model działający wewnątrz instytucji, gdzie każda decyzja algorytmu niesie ze sobą konkretne skutki społeczne i finansowe. To nie jest metafora, lecz analogia strukturalna, która pozwala nam zrozumieć rolę AI w nowoczesnym państwie. Gdy dodamy do tego perspektywę europejską, obraz staje się jeszcze wyraźniejszy i znacznie bardziej wymagający dla projektantów. Epoka naiwnie rozumianej neutralności technicznej dobiegła końca.

Komisja Europejska opisuje AI Act jako pierwszy kompleksowy akt prawny dotyczący sztucznej inteligencji, oparty na podejściu zależnym od poziomu ryzyka. W Europie model nie jest już tylko artefaktem obliczeniowym, lecz staje się również obiektem precyzyjnej kwalifikacji regulacyjnej w systemie prawnym. Rodzi on obowiązki zależne od przeznaczenia, sektora, wpływu na prawa podstawowe oraz stopnia ryzyka, jakie generuje dla otoczenia. To radykalnie wzmacnia tezę, że odpowiedzialność, dokumentacja i nadzór nie są biurokratycznym dodatkiem. Są one integralną częścią definicji technologii godnej zaufania, która może funkcjonować w przestrzeni publicznej.

Dla czytelnika profesjonalnego z doświadczeniem biznesowym najważniejsza konkluzja tej części powinna brzmieć bez żadnego znieczulenia. Pora porzucić infantylne rozróżnienie między nauką o modelach a praktyką wdrożeń, ponieważ ono jest już historycznie martwe. Współczesny projekt AI jest jednocześnie przedsięwzięciem epistemicznym, infrastrukturalnym, ekonomicznym i normatywnym, co wymaga od nas zupełnie nowych kompetencji. TensorFlow pozostaje ważnym językiem tej praktyki, ponieważ ucieleśnia połączenie wygody modelowania z potężnymi ścieżkami wdrożeniowymi. Od chmury po urządzenia brzegowe, narzędzie to pozwala nam materializować matematyczne abstrakcje w świecie rzeczywistym.

Reinforcement learning, czyli inteligencja rozumiana jako polityka działania pod wpływem bodźców, przypomina, że uczenie to coś więcej niż rozpoznawanie wzorców. Skalowanie danych uczy nas natomiast, że geometria wejścia nieubłaganie przesądza o losie optymalizacji każdego systemu. Architektury sieci neuronowych pokazują, że forma obliczeń musi zawsze odpowiadać strukturze świata, a nie tylko naszym teoretycznym wyobrażeniom. MLOps ostatecznie demaskuje cały romantyzm wokół sztucznej inteligencji, przypominając, że życie modelu nie kończy się na rekordowym benchmarku. Ono dopiero zaczyna się w systemie produkcyjnym, gdzie czekają dryf danych, koszty oraz odpowiedzialność prawna.

Gdybyśmy rozdzielili TensorFlow od całego ekosystemu pojęć, modeli i instytucji, zostałyby nam jedynie martwa składnia pozbawiona rozumu. Rozum współczesnej AI nie mieszka już bowiem w pojedynczej funkcji ani w eleganckiej architekturze wielowarstwowego perceptronu. MLP, czyli klasyczna kaskada przetwarzania reprezentacji, to tylko fundament, na którym budujemy hierarchiczny sens z surowych danych. Prawdziwa inteligencja mieszka w dobrze zorganizowanym łańcuchu, który prowadzi od surowych danych do konkretnych, odpowiedzialnych decyzji biznesowych. Uczenie maszynowe należy więc rozumieć jako układ instytucjonalny, a nie zbiór odizolowanych sztuczek obliczeniowych.

W przeciwnym razie łatwo popaść w dwa równie jałowe błędy, które hamują rozwój tej dziedziny w nowoczesnych organizacjach. Pierwszy polega na czysto podręcznikowym katalogowaniu metod, jakby były one jedynie martwymi eksponatami w zakurzonej gablocie muzeum techniki. Drugi błąd to publicystyczne fetyszyzowanie kilku modnych słów, sugerujące, że architektura Transformer unieważniła statystykę, a uczenie przez wzmacnianie zastąpiło teorię decyzji. Ani jedno, ani drugie nie jest prawdą, ponieważ fundamenty matematyczne pozostają niewzruszone mimo zmieniających się mód. W głębokim uczeniu nie wszystko, co poznawczo eleganckie, jest ekonomicznie rozsądne.

Hierarchia uczenia: Od nadzoru po logikę samonadzoru i wzmacniania

Współczesny paradygmat naukowy stawia przed nami wymagania znacznie wyższe niż prosta kategoryzacja algorytmów. Zmusza nas bowiem do dostrzeżenia głębokiej ciągłości między klasycznym uczeniem statystycznym, głębokim uczeniem reprezentacji a mechanizmami samonadzoru i złożonymi architekturami działającymi w dynamicznym środowisku. Musimy przy tym precyzyjnie odróżniać poszczególne warstwy zjawiska, gdyż inna jest logika zadania, inna logika danych, a jeszcze inna logika modelu czy samego wdrożenia. Osobny rozdział stanowi logika odpowiedzialności, która często umyka inżynierskiemu oku, a przecież bez niej system pozostaje

jedynie matematyczną abstrakcją. Aby zrozumieć ten gąszcz, należy najpierw przywrócić elementarną hierarchię pojęć, która porządkuje nasze myślenie. To nie jest detal implementacyjny.

Uczenie nadzorowane, czyli paradygmat oparty na modelowaniu relacji między obserwowanymi danymi a konkretnymi etykietami, pozostaje najbardziej czytelną i użyteczną formą modelowania w dzisiejszym świecie. W tym porządku dysponujemy jasnym zestawem danych wejściowych oraz odpowiedzią, której model ma się skrupulatnie nauczyć pod okiem algorytmu. Mieszczą się tu zarówno regresje przewidujące wartości ciągłe, jak i klasyfikatory rozstrzygające o przynależności do konkretnej klasy, co stanowi fundament współczesnej praktyki biznesowej. Scoring kredytowy, wykrywanie oszustw, prognoza popytu czy detekcja chorób na obrazach medycznych to wciąż domena tego klasycznego podejścia. Nawet najbardziej zaawansowane systemy filtrowania treści czy oceny ryzyka opierają się na tej prostej, lecz niezwykle skutecznej logice dopasowania do wzorca. To nadal rdzeń.

Prawdziwa rewolucja w tym obszarze nie polega na radykalnej zmianie definicji samego zadania, lecz na niespotykanej wcześniej skali i jakości reprezentacji danych. Dawniej to człowiek z mózgiem konstruował cechy ręcznie i dobierał do nich prostszy model, licząc na to, że jego intuicja okaże się trafna. Dziś coraz częściej architektura uczy się reprezentacji sama, wykorzystując wielowarstwowy perceptron (MLP), czyli kaskadowe przetwarzanie informacji, by lepiej aproksymować zależności między wejściem a etykietą. W tym sensie głębokie uczenie nie obaliło starego paradygmatu, lecz raczej zaktualizowało go do poziomu, który wcześniej wydawał się całkowicie nieosiągalny. Kto rozumie TensorFlow, czyli wszechstronny ekosystem służący do projektowania i wdrażania modeli, wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki. To nie jest metafora, lecz analogia strukturalna.

Uczenie nienadzorowane posiada zupełnie odmienną godność epistemiczną, ponieważ nie otrzymuje od świata gotowych etykiet, lecz próbuje samodzielnie odkryć ukrytą strukturę rzeczywistości. Jest to moment, w którym system nie tyle odpowiada na zadane mu pytanie, ile sam rekonstruuje geometrię zjawiska ukrytą w szumie informacji. Klasteryzacja, estymacja gęstości czy redukcja wymiarowości, czyli metody upraszczające strukturę danych przy zachowaniu ich kluczowych cech, należą właśnie do tego porządku. W biznesie i administracji podejście to bywa często niedoceniane, gdyż nie daje tak natychmiastowego komfortu jak prosty wskaźnik trafności klasyfikatora. A przecież bywa ono intelektualnie bardziej podstawowe, pozwalając nam zrozumieć, czy dane w ogóle tworzą sensowne skupienia. To nie jest detal.

Uczenie nienadzorowane to nie tylko zestaw metod statystycznych, lecz przede wszystkim praktyka epistemicznego rozpoznania materiału, zanim zacznie się produkować jakiekolwiek wiążące decyzje. W kulturze organizacyjnej zdominowanej przez przymus błyskawicznego wdrożenia

właśnie ten etap bywa najczęściej pomijany jako rzekomo nieefektywny czasowo. Potem wszyscy udają zbiorowe zdziwienie, że model działa w środowisku produkcyjnym znacznie gorzej niż na optymistycznych slajdach przygotowanych dla zarządu. Tymczasem to właśnie tutaj rozstrzyga się, czy przestrzeń problemu nie jest jedynie pozornie bogata, a faktycznie niskowymiarowa i podatna na błędy. Bez tego fundamentu budujemy zamki na piasku, licząc na to, że technologia sama naprawi nasze braki w zrozumieniu danych. To błąd.

Ostatnia dekada przyniosła ogromny wzrost znaczenia uczenia samonadzorowanego, które stało się głównym motorem przejścia od modeli zadaniowych do potężnych modeli bazowych. Logika tego procesu jest subtelna, lecz posiada ogromną moc sprawczą, pozwalając na wykorzystanie zasobów, które wcześniej pozostawały bezużyteczne. Nie dysponujemy milionami ręcznie opisanych danych dla każdego niszowego zadania, ale mamy dostęp do gigantycznych korpusów nieetykietowanego tekstu, obrazu czy dźwięku. Model otrzymuje więc zadanie sztucznie wygenerowane z samych danych, na przykład przewidywanie brakującego fragmentu tekstu lub zakrytego kawałka obrazu. W ten sposób system uczy się ogólnej struktury świata zapisanej w statystycznych zależnościach, co stanowi fundament pod późniejsze, precyzyjne dostrojenie. To nowa gleba.

To właśnie dlatego wielkie modele językowe nie są wyłącznie potężniejszymi klasyfikatorami, lecz modelami reprezentacyjnymi zasilanymi logiką samonadzoru i ogromną mocą obliczeniową. W tym miejscu widać wyraźnie, jak trafnie klasyczny podręcznikowy porządek uczenia łączy się z dzisiejszym światem modeli bazowych, tworząc nową jakość. Bez zrozumienia tego połączenia nie sposób pojąć, skąd wzięła się współczesna skala sztucznej inteligencji i jej zdolność do generalizacji wiedzy. Transfer learning, czyli technika wykorzystująca wiedzę zakodowaną w modelach już wytrenowanych, pozwala nam dziś na drastyczne obniżenie kosztów budowy nowych systemów. To nie jest tylko postęp techniczny, to zmiana ekonomiczna, która redefiniuje progi wejścia dla innowacji w niemal każdej branży. Skala ma znaczenie.

Na osobne potraktowanie zasługuje uczenie przez wzmacnianie, ponieważ jego recepcja społeczna została wyjątkowo silnie zniekształcona przez narrację medialnego spektaklu i technologiczną mitologię. Gdy wspomina się AlphaGo, wielu odbiorców słyszy przede wszystkim legendę o maszynie, która w spektakularny sposób pokonała ludzkiego mistrza w grze logicznej. Tymczasem rzeczywiste znaczenie tego wydarzenia nie polegało wyłącznie na zwycięstwie symbolicznym, lecz na ujawnieniu niezwykle skutecznej, nowej syntezy różnych paradygmatów. System ten łączył sieci polityki i wartości z przeszukiwaniem drzewa oraz uczeniem przez samogrę, co stanowiło majstersztyk inżynierii. Pokazało to dobitnie, że uczenie przez wzmacnianie nie działa w próżni heroicznych prób i błędów, lecz jest precyzyjnie komponowane. To nie jest magia.

Uczenie przez wzmacnianie jest w swojej istocie inżynierią sekwencyjnej decyzji, a nie nową religią czystej autonomii, jak chcieliby to widzieć niektórzy entuzjaści. Bywa ono tak pociągające dla nowoczesnych instytucji, ponieważ odpowiada ich ukrytemu marzeniu o systemie, który sam nauczy się skuteczności operacyjnej. Jest to marzenie menedżera, regulatora i stratega, którzy wierzą, że wystarczy odpowiednio ustawić nagrody, by świat zaczął działać idealnie. Jednak właśnie w tym miejscu należy zachować szczególny sceptycyzm, gdyż funkcja nagrody jest formalnym odpowiednikiem polityki publicznej. Jeśli uprościsz cel, system z matematyczną precyzją nauczy się optymalizować jego karykaturę, ignorując wszelkie intencje projektanta. Agent zrobi dokładnie to, do czego go zachęcono.

Dlatego właśnie uczenie przez wzmacnianie jest zarazem laboratorium inteligencji i laboratorium perwersji bodźców, które obnaża słabości naszych własnych systemów zarządzania. Uczy ono znacznie więcej o naturze sterowania i pułapkach biurokracji niż o jakiegokolwiek formie technologicznego zbawienia, na które liczymy. Wymaga ono wdrożenia rygorystycznych zasad MLOps i Governance, czyli praktyk obejmujących monitorowanie i odpowiedzialność za systemy, by uniknąć nieprzewidzianych skutków ubocznych. Dopiero gdy zrozumiemy, że algorytm jest lustrem naszych własnych priorytetów, będziemy mogli odpowiedzialnie przejść do omawiania konkretnych architektur. Bez tej świadomości każda innowacja pozostaje jedynie kosztownym eksperymentem, który może przynieść więcej szkód niż pożytku. Przejdźmy teraz do architektur.

Ewolucja architektur: Od klasyfikacji do symulacji świata

Wielowarstwowy perceptron, czyli klasyczna architektura sieci neuronowej oparta na kaskadowym przetwarzaniu danych, bywa dziś traktowany z protekcyjnym uśmiechem, niczym maszyna parowa w dobie reaktora termojądrowego. To jednak błąd poznawczy, bowiem MLP uosabia najbardziej fundamentalną intuicję całego nurtu głębokiego uczenia, polegającą na hierarchicznym budowaniu sensu z surowych informacji. Warstwa po warstwie konstruujemy coraz bardziej abstrakcyjne reprezentacje wejścia, aż do momentu, w którym ostateczne zadanie staje się trywialnie proste do wykonania. Nawet jeśli współczesne systemy oszałamiają nas swoją barokową złożonością, ta pierwotna idea pozostaje ich nienaruszalnym fundamentem, przy czym oficjalne materiały TensorFlow, czyli wszechstronnego ekosystemu do projektowania modeli, nie bez powodu akcentują elegancję budowania struktur warstwowych w bibliotece Keras. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki.

Prawdziwy przewrót kopernikański w dziedzinie analizy obrazu nastąpił wraz z nadejściem konwolucyjnych sieci neuronowych, które odniosły sukces nie dzięki brutalnej sile obliczeniowej, lecz dzięki wbudowaniu w model słusznego uprzedzenia strukturalnego. Zamiast traktować zdjęcie jak martwą, płaską tablicę liczb, architektury te uznały lokalność, współdzielenie wag oraz hierarchiczne wykrywanie cech za immanentne cechy wizualnej rzeczywistości. To było spektakularne zwycięstwo dobrze dobranego założenia o świecie; wszak obraz posiada swoją specyficzną topologię i niezmienniczość na przesunięcia, co pozwoliło algorytmom wreszcie „zrozumieć” przestrzeń. Każda udana architektura jest w gruncie rzeczy sformalizowanym uprzedzeniem o naturze danych, stanowiącym filtr oddzielający istotne wzorce od informacyjnego tła. CNN-y nie wygrały dlatego, że były najbardziej uniwersalnymi narzędziami w historii matematyki, lecz dlatego, że idealnie trafiły w strukturę problemu, który miały rozwiązać. W modelowaniu rzeczywistości rzadko wygrywa ten, kto deklaruje największą ogólność; znacznie częściej triumfuje ten, kto potrafi zdefiniować właściwe uprzedzenie indukcyjne.

Kolejnym etapem tej ewolucji były sieci rekurencyjne oraz ich bardziej wyrafinowane odmiany, takie jak LSTM czy GRU, które miały zmierzyć się z wyzwaniem pamięci sekwencyjnej. Przez lata stanowiły one główny aparat pojęciowy dla przetwarzania języka naturalnego i szeregów czasowych, wnosząc do uczenia maszynowego coś znacznie cenniejszego niż tylko sprytny mechanizm matematyczny. Wprowadziły one mianowicie intuicję czasu jako wewnętrznej struktury modelu, gdzie stan ukryty miał pełnić rolę magazynu przechowującego ślady przeszłości niezbędne do interpretacji teraźniejszości. Niemniej jednak sekwencyjna natura ich działania oraz narastające trudności optymalizacyjne stały się wąskim gardłem w epoce, w której skala zaczęła odgrywać rolę dominującą; wszak w głębokim uczeniu nie wszystko, co poznawczo eleganckie, jest ekonomicznie rozsądne. Pamięć, choć filozoficznie pociągająca, okazała się zbyt kosztowna w utrzymaniu przy gigantycznych zbiorach danych, co ostatecznie wymusiło zmianę podejścia. To nie był kryzys teorii, lecz kryzys skalowalności.

Przełom nastąpił wraz z publikacją słynnego artykułu „Attention Is All You Need”, który zakwestionował dotychczasowy paradygmat i zaproponował architekturę opartą wyłącznie na mechanizmach uwagi. Autorzy odważnie odrzucili konieczność stosowania rekurencji i konwolucji w zadaniach przetwarzania sekwencji, co pozwoliło na radykalne zwiększenie równoległości obliczeń i skrócenie czasu treningu. To właśnie w tym ruchu zawiera się cała logika współczesnej epoki modeli bazowych, gdzie zwyciężyła architektura najlepiej dopasowana do ekonomii współczesnego sprzętu komputerowego. Paradygmat naukowy nie uległ zmianie wyłącznie z powodu abstrakcyjnego dowodu matematycznego, lecz dlatego, że teoria, architektura i infrastruktura nagle zaczęły grać do jednej bramki, toteż transfer learning stał się nowym standardem efektywności. Technika ta, polegająca na wykorzystaniu wiedzy z modeli już

wytrenowanych do nowych zadań, pozwala na drastyczne obniżenie kosztów energetycznych przy budowie systemów, przy czym stanowi ona fundament dzisiejszej rewolucji. To była rewolucja przemysłowa w świecie algorytmów.

Dzisiejsze modele generatywne, od autoenkoderów po modele dyfuzyjne, przesuwają punkt ciężkości z prostej klasyfikacji na głębokie modelowanie rozkładu danych. Zamiast pytać algorytm o to, do której szufladki należy dany obiekt, zmuszamy go do zrozumienia procesu, który mógłby taki obiekt w ogóle powołać do życia. Jest to przesunięcie o kolosalnej wadze filozoficznej, w którym model przestaje być jedynie sędzią wydającym wyroki, a staje się symulatorem zdolnym do odtwarzania struktury świata. Rodzi to oczywiście nowe problemy techniczne i etyczne, niemniej naukowo rzecz biorąc, domyka to pewien historyczny łuk. Uczenie maszynowe dojrzewa na naszych oczach, ewoluując od pasywnej obserwacji i etykietowania ku aktywnemu modelowaniu mechanizmów wytwórczych rzeczywistości. To przejście od analizy do syntezy stanowi kluczowy krok w stronę dojrzalej inteligencji.

W tym kontekście musimy powrócić do tematu redukcji wymiarowości, czyli metod statystycznych pozwalających uprościć strukturę danych przy zachowaniu ich kluczowych cech informacyjnych. Wysokowymiarowość nie jest bowiem wyłącznie techniczną niedogodnością, lecz przede wszystkim fundamentalnym problemem epistemologicznym, z którym mierzy się każdy badacz. Im więcej cech bierzemy pod uwagę, tym łatwiej jest nam pomylić pozór bogactwa z rzeczywistą, wartościową informacją, co prowadzi do poznawczego chaosu. Redukcja wymiarowości pozwala nam zatem na efektywną kompresję i zrozumienie ukrytych zależności w złożonych zbiorach, chroniąc nas przed szumem informacyjnym. Zły sygnał nadal pozostaje złym sygnałem, choćby przechodził przez najdroższy klaster GPU, toteż umiejętność wydobywania esencji danych decyduje ostatecznie o sukcesie modelu. To nie jest detal, lecz warunek konieczny poznania.

Inżynieria poznawcza: Poza mit automatyzacji uczenia maszynowego

Analiza składowych głównych, czyli PCA, pozwalająca na zachowanie maksymalnej wariancji w niższym wymiarze, bywa nieocenionym narzędziem kompresji, wizualizacji oraz odsumowania sygnału. Metody rozmaitościowe, takie jak LLE czy bardziej eksploracyjnie używane t-SNE, przypominają nam z kolei, że dane rzadko spoczywają na prostym, liniowym podłożu. Mogą być one zwinięte, zakrzywione, lokalnie gładkie, a globalnie zdradliwe, co wymusza na nas porzucenie naiwnych założeń o strukturze rzeczywistości. To nie jest jedynie akademicka subtelność, lecz fundamentalna wskazówka o naturze świata społecznego, biologicznego i ekonomicznego, w

którym przyszło nam operować. Rzeczywiste zjawiska często posiadają znacznie prostszą strukturę lokalną, niż sugeruje to ich surowy, wielowymiarowy zapis cyfrowy. Droga do tej prostoty nie zawsze wiedzie jednak po linii prostej, wymagając od badacza niemal rzeźbiarskiej intuicji w usuwaniu zbędnych warstw informacji.

Podobnie rzecz się ma z inżynierią cech, która w powszechnym obiegu opinii doczekała się statusu reliktu dawno minionej epoki. Wśród entuzjastów głębokiego uczenia utrwaliło się przekonanie, że nowoczesne sieci neuronowe przejęły na siebie cały ciężar wydobywania istotnych informacji z szumu. To półprawda, a jak wiemy, półprawda bywa zazwyczaj najgroźniejszą i najbardziej wyrafinowaną formą fałszu intelektualnego w nauce. Owszem, w wielu spektakularnych zadaniach model faktycznie sam uczy się optymalnych reprezentacji, lecz proces ten nigdy nie zachodzi w próżni pojęciowej. Nadal bowiem ktoś musi zdefiniować jednostkę obserwacji, okno czasowe, horyzont predykcji oraz precyzyjny sposób agregacji dostępnych danych źródłowych. Bez tego fundamentu nawet najpotężniejsza architektura pozostaje jedynie kosztownym generatorem statystycznych artefaktów pozbawionych praktycznego znaczenia.

Wszystkie te czynności, od polityki imputacji braków po logikę etykietowania, stanowią inżynierię cech w sensie szerokim, będąc w istocie inżynierią poznawczą problemu. Gdy Fundacja Dobre Państwo podkreśla znaczenie jakości danych i rygorystycznego projektowania pipeline'ów, nie broni ona starego rzemiosła przeciwko nowoczesnej automatyzacji. Broni po prostu fundamentalnej tezy, że żadna architektura nie zwalnia człowieka z odpowiedzialności za semantykę wejścia i ostateczny sens celu. Warto w tym miejscu wrócić do TensorFlow, czyli wszechstronnego ekosystemu służącego do projektowania i wdrażania modeli, gdyż bez niego pozostaniemy przy ogólnej teorii. Siła tego narzędzia polega na tym, że scala ono kilka porządków naraz, oferując zarówno elastyczność, jak i produkcyjną dyscyplinę. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki.

Z jednej strony biblioteka ta oferuje intuicyjny tryb eager execution, który jest niezwykle przyjazny dla swobodnego eksperymentu i szybkiego debugowania błędów. Z drugiej strony pozwala ona budować grafy obliczeniowe niezależne od Pythona, co przekłada się na wyższą wydajność i łatwość przenoszenia rozwiązań między systemami. Oficjalne przewodniki kładą nacisk właśnie na tę podwójność, gdzie prostota pracy badawczej zostaje sprzężona z surowymi wymogami wdrażalności. To czyni z TensorFlow narzędzie szczególnie istotne dla tych środowisk, które nie chcą zatrzymać się na poziomie laboratoryjnego notebooka. Muszą one bowiem dowieźć realną produkcję, zapewnić monitoring i utrzymać system w dynamicznie zmiennych warunkach instytucjonalnych. Źródłowa obrona tego frameworka jako narzędzia systemowego jest zatem całkowicie zasadna i merytorycznie ugruntowana w realiach inżynierskich.

Nawet jeśli pluralizm dostępnych frameworków jest dziś faktem, TensorFlow pozostaje jednym z najczytelniejszych wcieleń idei, że model ma żyć w instytucjach, a nie tylko w próżni. Krajobraz uczenia maszynowego przestał być szkolną mapą algorytmów, stając się polem napięcia między różnymi sposobami wydobywania porządku z chaosu danych. Uczenie nadzorowane, czyli paradygmat oparty na modelowaniu relacji między danymi a etykietami, pozostaje królem zastosowań o jasnym celu i wysokiej precyzji. Uczenie nienadzorowane zachowuje z kolei godność najważniejszego narzędzia rozpoznania ukrytej struktury tam, gdzie brakuje nam gotowych odpowiedzi i jasnych drogowskazów. Samonadzór stał się natomiast drogą do modeli bazowych i epoki skalowalnych reprezentacji, co pozwala na transfer learning, czyli wykorzystanie wcześniej zdobytej wiedzy w nowych zadaniach. Reinforcement learning pokazuje zaś, że inteligencja bywa polityką działania, a nie tylko statystyką biernych obserwacji otoczenia.

Współczesne architektury sieci neuronowych, w tym klasyczny wielowarstwowy perceptron (MLP) oparty na kaskadowym przetwarzaniu danych, to w gruncie rzeczy sformalizowane uprzedzenia o naturze świata. Redukcja wymiarowości, czyli metody statystyczne upraszczające strukturę danych, to nie tylko techniczna kompresja, lecz także sztuka odsłaniania prostoty ukrytej pod nadmiarem cech. Inżynieria cech nie umarła, lecz przeniosła się z poziomu ręcznego dłubania w kolumnach na poziom konceptualnego projektowania świata, który model ma zobaczyć. Poprzednie myśli porządkowały typy uczenia i logikę architektur, teraz jednak należy zejść niżej, do samej mechaniki operacyjnej skuteczności. W uczeniu maszynowym nie wystarcza bowiem posiadać model poprawnie opisany na papierze czy w akademickiej publikacji. Trzeba jeszcze sprawić, aby nauczył się on czegoś rzeczywistego i zachował użyteczność po opuszczeniu sterylnego świata eksperymentu.

W tym momencie kończy się niewinna metafizyka "inteligentnej maszyny", a zaczyna twarda, niekiedy brutalna ekonomia optymalizacji ograniczonych zasobów. Trening modelu jest bowiem procesem alokacji zasobów obliczeniowych, informacyjnych i organizacyjnych w taki sposób, aby uzyskać najlepszą możliwą zdolność uogólniania. Nie ma tu miejsca na sentymentalizm czy wiarę w magiczne właściwości algorytmów, które same rozwiążą nasze problemy strukturalne bez udziału człowieka. Liczy się geometria przestrzeni cech, dynamika gradientu, dyscyplina walidacji oraz realny koszt pamięci i energii zużytej podczas wielogodzinnych obliczeń. Istnieje także koszt błędu oraz koszt złudzenia, że wysoki wynik w benchmarku jest tożsamy z sukcesem w rzeczywistym świecie. Zły sygnał nadal pozostaje złym sygnałem, choćby przechodził przez najdroższy klaster GPU dostępny na rynku.

Elementarna, a zarazem najczęściej lekceważona prawda brzmi tak, że proces treningu zaczyna się na długo przed uruchomieniem pierwszej epoki obliczeniowej. Zaczyna się on od tego, jak

zdefiniowano dane wejściowe, jak zbudowano zbiory treningowe i walidacyjne oraz jakie relacje czasowe między nimi zachowano. Kto traktuje podział danych jako czynność rutynową, ten nie rozumie, że właśnie na tym etapie rozstrzyga się przyszły sukces lub publiczna kompromitacja projektu. Data leakage, czyli błąd polegający na nieuprawnionym przenikaniu informacji ze zbioru testowego do treningowego, nie jest jedynie banalną usterką techniczną. Jest to raczej epistemiczne oszustwo wobec samego procesu poznania, które czyni wyniki modelu całkowicie bezwartościowymi w praktyce. Model, który widział przyszłość ukrytą w cechach, nie jest skuteczny, lecz po prostu skażony błędem, którego nie da się naprawić.

W tym kontekście wraca znaczenie projektu end-to-end, o którym tak trafnie mówi tradycja praktyki MLOps, czyli zestawu metod zarządzania cyklem życia modeli. Należy natychmiast oddzielić zbiór testowy i przestać go traktować jak wygodny magazyn inspiracji do kolejnych iteracji parametrów czy architektury. Trzeba zdefiniować walidację odpowiadającą naturze problemu, bo inne reguły obowiązują w danych niezależnych, a inne w złożonych i kapryśnych szeregach czasowych. Wiele organizacji mówi o sztucznej inteligencji z futurystycznym namaszczeniem, ale działa poznawczo jak sklecony naprędce hazardzista liczący na łut szczęścia. Raz testują na danych z przyszłości, innym razem mieszają obserwacje tego samego klienta pomiędzy zbiorami, a potem z powagą wygłaszają słowo "generalizacja". To nie jest żadna generalizacja, lecz jedynie ceremonialne nadużycie statystyki, które prędzej czy później zostanie brutalnie zweryfikowane przez rzeczywistość.

Skalowanie danych należy rozumieć znacznie szerzej niż tylko jako szkolny preprocessing polegający na usuwaniu wartości odstających czy prostej normalizacji zmiennych. Cechy o skrajnie różnych skalach deformują geometrię zadania, sprawiając, że spadek gradientu zaczyna błędzić po wydłużonych dolinach funkcji kosztu. Ruchy modelu przypominają wtedy bardziej szarpanie rozchwianej maszyny niż systematyczne i eleganckie schodzenie ku globalnemu minimum błędu predykcji. Standaryzacja i skalowanie nie są zatem zabiegami estetycznymi, lecz konieczną interwencją w samą strukturę optymalizacji, by przestała ona karać model za arbitralne jednostki miary. Istnieje jednak poziom drugi, bardziej subtelny, gdzie skalowanie staje się kwestią ustrojową całego projektu technologicznego i jego trwałości. Kiedy ekosystemy takie jak TensorFlow rozwijają potoki danych zdolne do strumieniowego ładowania i prefetchingu, "skalowanie" zyskuje zupełnie nowy wymiar.

W tym ujęciu skala danych nie jest wyłącznie liczbą rekordów w bazie, lecz relacją między objętością materiału a przepustowością całej architektury informatycznej. Oznacza to zdolność organizacji do utrzymania ciągłości przepływu między danymi surowymi, ich transformacją, a ostateczną inferencją i monitoringiem w środowisku produkcyjnym. Skalowanie to zatem rytm

trenowania sprzężony z architekturą wdrożenia, który pozwala systemowi AI oddychać w tempie napływających z zewnątrz informacji. Bez tego sprawnego krwioobiegu danych, nawet najbardziej wyrafinowany model pozostaje jedynie statycznym artefaktem, niezdolnym do realnej adaptacji w czasie. Dane nie są więc "nową ropą", lecz raczej nową glebą, która wymaga ciągłej uprawy i odpowiedniego nawodnienia przez rygorystyczną inżynierię. To nie jest metafora, lecz analogia strukturalna, która najlepiej oddaje wyzwania stojące przed współczesną inżynierią poznawczą w dobie automatyzacji.

Gleba danych i demony głębi: Jak zarządzać stabilnością modelu

Współczesne systemy sztucznej inteligencji nie upadają wyłącznie z powodu niedoboru danych, lecz częściej przez chaotyczny nadmiar źle ustrukturyzowanego szumu. W debacie publicznej karmiono nas uwodzicielską, choć fundamentalnie błędną metaforą danych jako „nowej ropy”, sugerującą, że sama ekstrakcja surowca generuje natychmiastową wartość. W rzeczywistości dane posiadają znacznie mniej bezpośrednią użyteczność, która zależy od precyzji etykiet, integralności reprezentacji oraz żmudnego kosztu ich czyszczenia. Dane nie są więc „nową ropą”. Są raczej nową glebą, na której plon zależy od chemicznego składu podłoża oraz cierpliwości rolnika. Trzeba umieć odróżnić ziemię żyzną od zatrutej, rozumiejąc, że ocean cyfrowych śmieci pozwoli modelowi wyłowić jedynie bardziej wyrafinowane złudzenia. Więcej nie znaczy lepiej; więcej znaczy po prostu głośniejsze.

Kiedy wchodzimy w głąb architektury sieci neuronowych, napotykamy problem stabilnego przepływu sygnału przez kaskadę kolejnych warstw. To właśnie tutaj ujawniają się dwa klasyczne demony, które przez dekady czyniły głębokie uczenie obietnicą bardziej teoretyczną niż praktyczną: znikające oraz eksplodujące gradienty. Zjawiska te nie są jedynie techniczną niedogodnością, lecz przypomnieniem, że głębokość modelu jest luksusem okupionym wyjątkową delikatnością struktury. Każda dodatkowa warstwa zwiększa potencjał reprezentacyjny systemu, ale zarazem drastycznie komplikuje warunki transportu informacji przez sieć. Głębokość nie jest darmową mocą obliczeniową, lecz inwestycją o rosnącym koszcie koordynacji. To nie jest detal implementacyjny. To fundament.

Jeżeli gradient maleje zbyt szybko podczas procesu propagacji wstecznej, dolne warstwy sieci zapadają w rodzaj poznawczego letargu i niemal przestają się uczyć. Jeśli zaś rośnie on w sposób niekontrolowany, aktualizacje wag stają się gwałtowne, co błyskawicznie destabilizuje cały matematyczny gmach modelu. W intuicji laickiej mogłoby się wydawać, że skoro sieć i tak będzie

trenowana, punkt startowy nie ma większego znaczenia dla końcowego wyniku. Jest to błąd wynikający z ignorowania wewnętrznej dynamiki systemu, w którym warunki początkowe przesądzają o drożności kanałów informacyjnych. Inicjalizacje Glorota, He czy LeCuna nie są zatem estetyczną fanaberią literatury akademickiej, lecz próbą ustawienia parametrów tak, aby wariancja aktywacji nie uległa katastrofalnym zniekształceniom.

W tym miejscu głębokie uczenie zdradza swoje bliskie pokrewieństwo z ekonomią makro oraz teorią systemów złożonych. Wiele zależy bowiem nie tylko od samych reguł ruchu, lecz od punktu, w którym ten ruch się rozpoczyna. Ten sam mechanizm aktualizacji przy dwóch różnych startach potrafi prowadzić do skrajnie odmiennych losów algorytmu, co demistyfikuje technologiczną teologię o nieomyślności kodu. Równie krytyczna okazuje się funkcja aktywacji, czyli element decydujący o tym, czy dany neuron przekaże sygnał dalej do kolejnych warstw. Historycznie sieci cierpiały z powodu nadmiernego przywiązania do funkcji sigmoidalnych, które choć matematycznie eleganckie, przy dużych wartościach nasycaly się i praktycznie dusiły gradient. Były powabne, lecz w praktyce inżynierskiej okazywały się zdradliwe.

Przełom przyniosło wejście funkcji ReLU, która wprowadziła prostą i brutalnie skuteczną nieliniowość, nieobcinając sygnału dla wartości dodatnich w tak dramatyczny sposób. Choć rozwiązała ona wiele problemów, zrodziła własne wyzwania, jak choćby zjawisko „umierających neuronów”, co wymusiło ewolucję w stronę wariantów takich jak Leaky ReLU, ELU czy GELU. Każda z tych form jest próbą ratowania żywotności sygnału w coraz głębszych strukturach, gdzie margines błędu staje się coraz cieńszy. Kto rozumie TensorFlow, czyli wszechstronny ekosystem służący do projektowania i wdrażania modeli, wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki. Stabilność nie jest stanem danym raz na zawsze, lecz procesem ciągłego zarządzania napięciem między głębią a chaosem. To nie jest metafora, lecz analogia strukturalna.

Inżynieria niepewności: optymalizacja i sztuka zatrzymania

Ewolucja uczenia głębokiego nie przypomina wcale triumfalnego marszu ku coraz wyższym poziomom abstrakcji, lecz raczej mozolne wykuwanie trafniejszych heurystyk inżynierskich. Bardzo często postęp w tej dziedzinie nie bierze się z nowej, olśniewającej metafizyki inteligencji, ale z prozaicznego, choć precyzyjnie dobranej kształtu funkcji matematycznej. Najlepsza nawet funkcja aktywacji nie unieważnia bowiem fundamentalnej potrzeby normalizacji sygnału, która chroni system przed informacyjnym chaosem. Batch normalization stała się jednym z wielkich wynalazków inżynierii, ponieważ wbudowała w architekturę sieci mechanizm centrowania i

skalowania sygnałów bezpośrednio w toku treningu. Pozwoliło to badaczom używać wyższych współczynników uczenia, drastycznie zmniejszyć wrażliwość na początkową inicjalizację wag i poprawić stabilność całego procesu. Trzeba jednak z całą stanowczością odróżnić dwie sfery porządku, które często bywają mylone przez początkujących adeptów sztuki.

Zewnętrzne skalowanie cech wejściowych nie przestaje być potrzebne tylko dlatego, że w głębi modelu działa zaawansowana normalizacja warstwowa. To nie są dwa konkurencyjne rytuały, lecz dwa zupełnie inne poziomy strukturalnej dyscypliny danych. Jeden z nich dotyczy statycznej geometrii wejścia, drugi zaś dynamicznej natury sygnału przepływającego przez warstwy ukryte modelu. Kto miesza te porządki, ten zachowuje się jak administrator, który sądzi, że skoro w budynku działa nowoczesna klimatyzacja, to nie trzeba już sprawdzać jakości materiałów użytych do fundamentów. Kto rozumie TensorFlow, czyli wszechstronny ekosystem do projektowania modeli, wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki. Tu właśnie dochodzimy do roli optymalizatorów, które stanowią serce dynamiki każdego systemu uczącego się.

Klasyczny gradient descent pozostaje piękną, niemal platońską ideą, ale w brutalnej praktyce głębokiego uczenia bywa on zbyt powolny i beznadziejnie naiwny. Współczesne optymalizatory są w istocie wyrafinowanymi sposobami korygowania tej naiwności poprzez matematyczną ekstrapolację ruchu. Momentum nadaje procesowi niezbędną bezwładność, dzięki której optymalizacja nie grzęźnie tak łatwo na płaskich odcinkach i nie zmienia nerwowo kierunku przy każdym lokalnym kaprysie powierzchni kosztu. Nesterov dodaje do tego element bardziej wyprzedzający, pozwalając modelowi „spojrzeć” nieco dalej przed siebie. AdaGrad z kolei próbował dostosowywać kroki do skali gradientów, lecz przy długim treningu wykazywał tendencję do przedwczesnego zamierania. RMSProp podtrzymał tę cenną ideę adaptacji, ale uczynił ją znacznie bardziej praktyczną w obliczu niestacjonarnych celów.

Adam połączył adaptacyjność z pamięcią gradientu i stał się na długi czas niemal domyślnym wyborem dla ogromnej części komercyjnych i naukowych zastosowań. Dopiero z czasem zaczęto poważniej analizować jego własności generalizacyjne oraz skomplikowaną relację do mechanizmów takich jak weight decay. Przykładem tej ewolucji jest AdamW, który precyzyjnie rozdziela efekt regularyzacji od samej dynamiki adaptacyjnej optymalizatora. Z zewnątrz może to wyglądać jak drobiazg dla wtajemniczonych, lecz w istocie jest to walka o to, jak szybko i jak sensownie system ma się poruszać po krajobrazie parametrycznym. To nie są kosmetyczne ustawienia w pliku konfiguracyjnym, lecz realna polityka ruchu modelu w wielowymiarowej przestrzeni możliwości. W wielu organizacjach wybór konkretnego optymalizatora przypomina niestety wybór sloganu strategicznego, a nie rzetelną analizę techniczną.

Ktoś przeczytał w popularnym tutorialu, że Adam „działa prawie zawsze”, więc Adam staje się w zespole nietykalną religią. Ktoś inny usłyszał na konferencji, że klasyczny SGD z momentum daje lepszą generalizację, więc zaczyna traktować wszelkie metody adaptacyjne jak niebezpieczną herezję. Tymczasem profesjonalizm w uczeniu maszynowym nie polega na bezkrytycznym czczeniu jednej, uniwersalnej formuły. Polega on na głębokim rozumieniu, jaki rodzaj krajobrazu optymalizacyjnego wytwarza dany model i jaki kompromis między szybkością zbieżności a jakością minimum jesteśmy gotowi przyjąć. Fanatyzm optymalizacyjny jest równie groteskowy jak fanatyzm ideologiczny, a różni się od niego tylko tym, że zamiast płomiennych transparentów produkuje kolorowe wykresy w TensorBoardzie. To nie jest detal implementacyjny, to kwestia intelektualnej uczciwości projektanta.

Nie mniej istotny w tym procesie pozostaje współczynnik uczenia oraz precyzyjnie zaprojektowany harmonogram jego zmian w czasie. Stały learning rate bywa kusząco prosty w implementacji, ale rzadko okazuje się rozwiązaniem optymalnym dla złożonych architektur. Wiele procesów treningowych przebiega najskuteczniej wtedy, gdy model najpierw porusza się śmieiej, a potem coraz ostrożniej doprecyzowuje swoje położenie w regionie minimum. Stąd bierze się znaczenie harmonogramów wykładniczych, schodkowych czy cosinusowych, które wprowadzają do treningu niezbędną dynamikę. Strategia typu one cycle, na pierwszy rzut oka paradoksalna, zakłada, że kontrolowane zwiększenie współczynnika uczenia może pomóc modelowi opuścić marne rejony krajobrazu. Jest w tym coś niemal filozoficznego, co wykracza poza czystą matematykę optymalizacji.

Czasami, aby ostatecznie znaleźć trwały porządek, trzeba na krótki moment świadomie zwiększyć amplitudę ruchu i zaryzykować destabilizację. Nadmierna ostrożność w doborze parametrów nie zawsze jest przejawem racjonalności; bywa ona jedynie elegancką i bezpieczną formą utknięcia w martwym punkcie. Na tym tle regularyzacja jawi się jako sztuka cywilizowania nadmiernej złożoności modelu, który zbyt łatwo ulega pokusie zapamiętywania szumu. Przeuczenie nie jest moralną wadą algorytmu, lecz logiczną konsekwencją sytuacji, w której moc dopasowania modelu przewyższa informacyjną treść dostępnych danych. Dane nie są więc „nową ropą”, są raczej nową glebą, którą trzeba umiejętnie uprawiać, by nie jałowiała pod wpływem zbyt agresywnych parametrów. Klasyczne kary na wagach próbują ograniczać rozrzut parametrów, ale to Dropout wprowadził do gry zupełnie nową jakość.

Dropout działa inaczej niż tradycyjne metody, zmuszając sieć do tego, by nie polegała obsesyjnie na jednej, wydeptanej ścieżce sygnału. To trochę tak, jakby organizację co chwilę pozbawiać części pracowników, aby sprawdzić, czy wiedza rzeczywiście jest rozproszona, czy tylko pozornie rezyduje w głowach kilku osób. Jeśli system po takim brutalnym zabiegu nadal działa poprawnie, oznacza to,

że nie był uzależniony od jednego, kruchego korytarza informacji. Early stopping z kolei jest jedną z najbardziej trzeźwych i pragmatycznych technik w całym arsenale współczesnego data science. Zamiast marzyć o coraz doskonalszym dopasowaniu do przeszłości, metoda ta uznaje, że przychodzi moment, w którym dalsze ulepszanie treningu nieuchronnie pogarsza naszą przyszłość.

W tym sensie early stopping jest może najbardziej filozoficzną z technik regularyzacji, ponieważ przypomina nam, że dojrzałość polega nie na ciągłym ulepszaniu, lecz na umiejętności zatrzymania się. W głębokim uczeniu olbrzymią rolę odgrywa także transfer learning, czyli technika wykorzystująca wiedzę zakodowaną w modelach już wytrenowanych do rozwiązywania nowych problemów. Jego sens ekonomiczny jest oczywisty dla każdego, kto potrafi liczyć koszty energii i czasu procesora. Pozwala on na drastyczne obniżenie barier wejścia przy budowie zaawansowanych systemów, czyniąc technologię bardziej dostępną i efektywną. Ostatecznie cała ta inżynieria niepewności sprowadza się do zarządzania dynamiką, w której matematyczna precyzja spotyka się z intuicją rzemieślnika.

Od romantycznej wizji AI ku dojrzałej inżynierii systemów

Trenowanie kolosalnego modelu od zera to przedsięwzięcie, które pochłania nie tylko gigantyczne zasoby kapitałowe, ale również energię i czas całych zespołów badawczych. W świecie, w którym dysponujemy już modelami wytrenowanymi na potężnych korpusach danych o szerokim zasięgu reprezentacyjnym, zwykły pragmatyzm nakazuje traktować je jako fundament pod dalszą budowę. Transfer learning, czyli technika polegająca na wykorzystaniu wiedzy zakodowanej w modelach już wytrenowanych do rozwiązywania nowych, pokrewnych problemów, staje się tu kluczowym narzędziem ekonomicznym. Niższe warstwy sieci zazwyczaj uczą się cech o charakterze uniwersalnym, które warto zachować, podczas gdy dopiero te wyższe wymagają precyzyjnego dostrojenia do specyfiki konkretnego zadania. Jest to jeden z najbardziej jaskrawych dowodów na to, że uczenie maszynowe przestało być domeną czystej teorii, a stało się integralną częścią nowoczesnej gospodarki opartej na wiedzy. W głębokim uczeniu nie wszystko, co poznawczo eleganckie, jest ekonomicznie rozsądne.

Współczesna inżynieria rzadko kiedy pozwala sobie na luksus budowania wszystkiego od absolutnych podstaw, preferując działanie poprzez odziedziczone moduły, warstwy i interfejsy. Transfer learning nie jest zatem wyłącznie sprytną sztuczką techniczną, lecz specyficzną ekonomią dziedziczenia wiedzy w cyfrowym ekosystemie, gdzie kapitał intelektualny ulega ciągłej kumulacji. Warto w tym miejscu pochylić się nad rolą uczenia nienadzorowanego, które często bywa

niesłusznie spychane na margines podręcznikowych rozważań jako osobny, czysto teoretyczny rozdział. Tymczasem redukcja wymiarowości, czyli metody statystyczne i geometryczne upraszczające strukturę danych przy zachowaniu ich kluczowych cech, stanowi fundament przygotowania świata dla modelu. Narzędzia takie jak PCA, klasteryzacja czy autoenkodery pozwalają odzyskać dane, wygenerować nowe cechy i dostarczyć gęstych reprezentacji, które stają się idealnym wejściem dla zadań nadzorowanych. Profesjonalista nie traci czasu na jałowe spory o wyższość jednego paradygmatu nad drugim, lecz pyta, jak połączyć te porządki poznania w spójną całość.

Zrozumienie, jak system wchodzi w interakcję ze swoim środowiskiem, wymaga narzędzi diagnostycznych, wśród których krzywe uczenia zajmują miejsce szczególne, choć często są one karygodnie lekceważone. Są one najprostszą, a zarazem najbardziej wymowną metodą weryfikacji, czy nasz model rzeczywiście rozumie strukturę danych, czy jedynie mechanicznie je zapamiętuje. Jeśli błąd na zbiorze treningowym i walidacyjnym utrzymuje się na wysokim poziomie, mamy do czynienia z niedouczeniem, które sugeruje zbyt małą złożoność wybranej architektury. Sytuacja, w której błąd treningowy drastycznie spada, a walidacyjny pozostaje wyraźnie wyższy, to klasyczny sygnał przeuczenia, czyli utraty zdolności do generalizacji na nowe przypadki. Prawdziwy dramat zaczyna się jednak wtedy, gdy model świetnie radzi sobie w laboratorium, by po wdrożeniu polec w starciu z rzeczywistym dryfem danych lub wadliwą definicją celu biznesowego. Krzywe uczenia nie są więc tylko estetycznym wykresem w raporcie, lecz fundamentalną teorią diagnostyczną, której ignorowanie jest błędem w sztuce. Kto ich nie czyta, przypomina lekarza odmawiającego analizy parametrów życiowych pacjenta w imię wiary we własną, nieomylną intuicję.

Przejście od fazy radosnego eksperymentowania do twardej produkcji wymaga wypowiedzenia prawdy, która w świecie entuzjastów technologii bywa przyjmowana z wyraźnym oporem. Ogromna część projektów z zakresu sztucznej inteligencji umiera nie z powodu błędów matematycznych w architekturze, lecz przez brak solidnej infrastruktury utrzymaniowej i monitoringu. Model po opuszczeniu bezpiecznych murów laboratorium trafia do świata, który nieustannie ewoluuje, zmieniając rozkłady cech, zachowania użytkowników oraz kanały wejścia danych. Pojawia się zjawisko dryfu danych (data drift), czyli zmiany rozkładu danych wejściowych, a także degradacja jakości sensorów czy zmiany w regulacjach prawnych, które sprawiają, że model pozostawiony samemu sobie zaczyna po prostu gnić. MLOps i governance, czyli zbiór praktyk obejmujących monitorowanie, dokumentację i odpowiedzialność za systemy, nie są zatem zbędną nadbudową, lecz koniecznym domknięciem procesu inżynierskiego. Jeśli system nie jest poddawany ciągłemu nadzorowi, to w sensie produkcyjnym nigdy nie został on w pełni zbudowany, a jedynie wypuszczony na wolność bez opieki. To nie jest detal implementacyjny. Nigdy nie był.

Dzisiejsza debata publiczna, zafascynowana potęgą akceleratorów i urokiem generatywnych demonstracji, często gubi z oczu to, co w projektach technologicznych jest najistotniejsze. Prawdziwa dojrzałość nie objawia się w liczbie parametrów modelu, lecz w rzetelności podziału danych, dyscyplinie monitoringu i odporności na nieunikniony dryf. Profesjonalna praktyka w tym obszarze jest znacznie bliższa rygorystycznej inżynierii lotniczej niż beztroskiej kulturze startupowej mitologii, która karmi się opowieściami o cudach. Samolot nie staje się bezpieczny przez swój wygląd na pokazie, lecz dzięki procedurom i redundancji – z modelami AI jest podobnie, choć wielu ludziom wciąż nie mieści się to w głowie. Najdojrzalsze podejście cechuje się głęboką nieufnością wobec własnych wyników, co wymusza stosowanie kart modeli, analizy błędów i dokumentowania założeń. Nauka o modelach dostarcza nam narzędzi pokory, których organizacje często nie chcą używać, wybierając zamiast nich efektowne prezentacje dla zarządu.

Należy stwierdzić bez zbędnych upiększeń, że trenowanie modelu nie jest aktem jednorazowym, lecz złożonym, wielowarstwowym procesem o charakterze inżynieryjnym. Skalowanie danych koryguje geometrię zadania, podczas gdy inicjalizacja wag ustawia warunki początkowe dla całej przyszłej dynamiki systemu. Funkcje aktywacji decydują o przepływie sygnału, a normalizacja stabilizuje ruch informacji przez kolejne warstwy, zapobiegając chaosowi obliczeniowemu. Wielowarstwowy perceptron (MLP), czyli klasyczna architektura sieci neuronowej oparta na kaskadowym przetwarzaniu i przekształcaniu reprezentacji wejściowych, uosabia tę hierarchiczną strukturę budowania sensu z surowych danych. Optymalizatory kształtują politykę poruszania się po skomplikowanym krajobrazie kosztu, a harmonogramy tempa uczenia regulują rytm tych poszukiwań, chroniąc przed chorobą nadmiernego dopasowania. Każdy z tych elementów musi ze sobą współpracować, aby finalny system mógł sprawnie korzystać z zakumulowanego kapitału wiedzy. Bez tej precyzyjnej mechaniki, nawet najpotężniejszy model pozostanie jedynie zbiorem martwych wag.

Monitoring po wdrożeniu przypomina nam brutalnie, że model żyje w czasie i podlega prawom entropii, a nie istnieje w wiecznym, statycznym benchmarku. Projekt typu end-to-end nie jest prostym ciągiem technicznych etapów do odhaczenia, lecz nowym typem porządku poznawczego, gdzie dane i infrastruktura tworzą jeden organizm. Właśnie dlatego przesunięcie punktu ciężkości z romantycznych wizji ku praktyce MLOps jest ruchem intelektualnie słusznym i niezbędnym. TensorFlow, czyli wszechstronny ekosystem i biblioteka programistyczna służąca do projektowania oraz wdrażania modeli, stanowi fundament pozwalający na skalowanie tych rozwiązań od fazy prototypu do złożonych systemów produkcyjnych. Biblioteka ta łączy wysokopoziomowe interfejsy z potężnymi potokami danych, co pozwala widzieć projekt AI jako sekwencję zależności, a nie zbiór odizolowanych modułów. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki.

Prawidłowo skonstruowany proces zaczyna się od rozstrzygnięcia, czy dany problem jest w ogóle sensownie formalizowalny i jaki cel biznesowy ma zostać osiągnięty. Trzeba ustalić, jakie dane są rzeczywiście dostępne, jakie są ich ograniczenia, jaki rodzaj etykiet istnieje oraz czy nie występuje niedopasowanie między danymi treningowymi a docelowym środowiskiem użycia. Modelowanie, wbrew powszechnym wyobrażeniom, nie jest pierwszym aktem tego technologicznego dramatu, lecz jego środkowym, choć najbardziej widowiskowym etapem. Współczesna kultura technologiczna uwielbia zaczynać opowieść od architektury, bo to ona najlepiej sprzedaje się w mediach i na konferencjach. Tymczasem odpowiedzialna praktyka inżynierska kładzie nacisk na semantykę problemu oraz na to, co dzieje się z systemem po jego uruchomieniu. Porządek ten jest widoczny w filozofii nowoczesnych narzędzi, które akcentują nie tylko budowę modeli, ale przede wszystkim ich przenośność i wydajność.

TensorFlow jest czymś znacznie więcej niż tylko kolejną biblioteką programistyczną; to wzorcowy przykład przejścia ku dojrzałości, przemysłowemu paradygmatowi uczenia maszynowego. Z jednej strony pozwala na swobodę eksperymentu dzięki mechanizmom takim jak eager execution, które ułatwiają intuicyjne modelowanie w Keras i szybkie testowanie śmiałych hipotez. Z drugiej strony oferuje funkcję `tf.function`, czyli przejście od kodu Pythona do grafów obliczeniowych, co drastycznie poprawia wydajność i jest wymagane w profesjonalnych wdrożeniach. Zaawansowane mechanizmy automatycznego różniczkowania, GradientTape oraz potoki `tf.data` pokazują, że nowoczesny model nie jest tylko statycznym zbiorem wag, lecz całą zorganizowaną procedurą obliczeniową. Dla profesjonalnego odbiorcy oznacza to konieczność porzucenia marzeń o cudownym algorytmie na rzecz rzetelnej budowy systemów, które rozumieją własne środowisko i ograniczenia. To nie jest metafora, lecz analogia strukturalna, która definiuje przyszłość dojrzałej inżynierii systemów uczących się. Dane nie są więc nową ropą. Są raczej nową glebą.

Architektura decyzji: poza technologicznym optymizmem

Dzisiejsza przewaga rynkowa nie wynika bynajmniej z wyboru rzekomo lepszego modelu, lecz z umiejętności zaprojektowania spójnego systemu pracy z owym modelem. Zamiast postrzegać krajobraz uczenia maszynowego jako archipelag odizolowanych królestw, powinniśmy widzieć w nim jedną mapę, gdzie różne techniki odpowiadają na odmienne rodzaje niepewności. Klasyczne algorytmy sprawdzają się tam, gdzie priorytetem pozostaje kontrola interpretacyjna oraz oszczędność zasobów obliczeniowych. Z kolei głębokie sieci neuronowe przejmują stery, gdy dane

posiadają strukturę hierarchiczną, której nie sposób opisać prostym zestawem cech. To nie jest detal implementacyjny. To fundament strategii.

Uczenie nienadzorowane, czyli proces wykrywania wzorców w danych pozbawionych etykiet, pozwala nam rozpoznać strukturę rzeczywistości, zanim model zacznie wydawać jakiegokolwiek wiążące wyroki. W tym procesie kluczowa okazuje się redukcja wymiarowości, czyli technika upraszczania złożonych zbiorów informacji, która odsłania przed nami ukryte osie i porządkuje informacyjny nadmiar. Uczenie przez wzmacnianie wprowadza do tego układu problem działania pod wpływem konkretnych bodźców zewnętrznych. Podważa ono popularne złudzenie, że inteligencja sprowadza się wyłącznie do biernej klasyfikacji obserwacji, czyniąc z modelu aktywnego uczestnika gry. W głębokim uczeniu nie wszystko, co poznawczo eleganckie, jest ekonomicznie rozsądne.

Zbiór tych technik nie stanowi kolekcji rozłącznych rozdziałów, lecz kompletny zestaw instrumentów przeznaczonych dla różnych typów cyfrowego świata. Na tym tle wyraźnie widać, jak bardzo mylące i wręcz intelektualnie dziecinne jest popularne przeciwstawianie starego uczenia maszynowego nowej, rzekomo wszechpotężnej sztucznej inteligencji. Dojrzała organizacja potrafi bowiem w ramach jednego systemu wykorzystać analizę głównych składowych do kompresji reprezentacji oraz klasteryzację do budowania nowych cech użytkowych. Ten sam mechanizm może angażować klasyczny model do zadań pomocniczych, podczas gdy głęboka sieć zajmuje się percepcją sygnału. Model nigdy nie istnieje w pojedynkę, lecz zawsze funkcjonuje wewnątrz szerszej architektury decyzji. Nie jest. Nigdy nie był.

Wspomniana architektura decyzji posiada zawsze charakter ekonomiczny, prawny oraz antropologiczny, co demistyfikuje popularną ostatnio teologię technologiczną wokół systemów autonomicznych. Wymiar ekonomiczny objawia się w kosztach infrastrukturalnych i energetycznych, natomiast aspekt prawny narzuca modelom surowe obowiązki dokumentacyjne oraz precyzyjną kwalifikację ryzyka. Najciekawszy wydaje się jednak wymiar antropologiczny, ponieważ każdy model zakłada określoną wizję użytkownika lub obywatela, którą następnie włącza w praktykę instytucjonalną. Większość współczesnych sporów o sztuczną inteligencję jest źle ustawiona, gdyż koncentruje się na widowiskowym pytaniu, czy algorytmy stają się coraz mądrzejsze. Jest to kwestia efektowna, lecz strategicznie drugorzędna.

Pytanie o wiele ważniejsze dotyczy tego, jakiego rodzaju porządek społeczny i organizacyjny budujemy wokół modeli zdolnych do autonomicznego generowania oraz rekomendowania decyzji. Problemem nie jest bowiem sama surowa moc obliczeniowa modelu, lecz układ bodźców i procedur, w których ta potężna siła zostanie ostatecznie wykorzystana. Jeśli taka uwaga wydaje się komuś zbyt polityczna, to tylko dlatego, że zbyt długo udawano, iż uczenie maszynowe można

zamknąć w sterylnej komorze. Od chwili, gdy algorytm dotyka kwestii zatrudnienia, kredytu czy bezpieczeństwa publicznego, wchodzi on nieuchronnie w sferę normatywną. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki.

W tym właśnie miejscu pojawia się Europa jako osobny aktor cywilizacyjny, wprowadzając akt o sztucznej inteligencji jako pierwszy kompleksowy akt prawny oparty na analizie ryzyka. Dokument ten wszedł w życie w sierpniu 2024 roku, a jego kolejne rygory są wdrażane etapami, obejmując między innymi wytyczne dotyczące zakazanych praktyk. Dla modeli ogólnego przeznaczenia, czyli systemów zdolnych do wykonywania szerokiego spektrum zadań, przewidziano dobrowolne kodeksy postępowania oraz specyficzne obowiązki dostawców. Choć reguły te zaczęły obowiązywać w sierpniu 2025 roku, to pełne egzekwowanie przepisów wobec istniejących już modeli nastąpi dopiero po dwóch latach. Rozróżnienie między tym, co pewne normatywnie, a tym, co jeszcze podlega przesunięciom, pozostaje kluczowe dla profesjonalisty.

AI jako nowa infrastruktura rozumu instytucjonalnego

Współczesna inżynieria systemów uczących się przestała być domeną samotnych wizjonerów zafascynowanych czystym kodem, a stała się fundamentem nowej architektury władzy. Nie wystarczy już bowiem wiedzieć, jak działa TensorFlow, czyli wszechstronny ekosystem służący do projektowania oraz wdrażania modeli, by móc realnie wpływać na kształt cyfrowej rzeczywistości. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki, zapominając, że technologia jest dziś przede wszystkim obiektem regulacyjnym. W praktyce oznacza to, że każdy poważny projekt AI powinien posiadać nie tylko wyśrubowane metryki skuteczności, lecz również rygorystyczny reżim dokumentacyjny oraz mapę ryzyka. Fundacja Dobre Państwo słusznie wpisuje te wymagania w perspektywę MLOps i governance, czyli zbioru praktyk obejmujących monitorowanie i odpowiedzialność za systemy. To nie jest biurokratyczny balast narzucony na „prawdziwą” technologię, lecz integralna część jej nowoczesnej definicji. Technologia godna zaufania nie jest wcale mniej techniczna; ona jest po prostu technicznie dojrzała.

Czas uczciwie zrewidować status współczesnego TensorFlow, gdyż nie byłoby rzetelne pisać o nim tak, jakby wciąż monopolizował on wyobraźnię badawczą całego świata. Tak nie jest i nigdy nie był, choć jego rola w ekosystemie pozostaje fundamentalna dla stabilności wdrożeń. Ostatnie lata przyniosły wyraźne umocnienie pluralizmu frameworków, a PyTorch zdobył dominującą pozycję w środowiskach akademickich i eksperymentalnych. Jednocześnie TensorFlow zachował swoje kluczowe znaczenie w złożonych ekosystemach produkcyjnych, modelach przenośnych oraz zaawansowanych rozwiązaniach urządzeniowych. Przykładem tego trendu jest dynamiczny rozwój

linii on-device, w której unijne komisje nie odgrywają większej roli, lecz sam ekosystem Google wyznacza strategiczne osie rozwoju. Deployowalność modeli na urządzenia końcowe pozostaje kluczowym wyzwaniem inżynieryjnym, decydującym o realnej użyteczności algorytmów w codziennym życiu. Dlatego obrona TensorFlow jest trafna pod warunkiem, że rozumiemy ją nie jako twierdzenie o hegemonii, ale jako uznanie trwałości jednego z kluczowych języków inżynierii.

Ten sam realizm trzeba zachować wobec reinforcement learning, czyli uczenia ze wzmocnieniem, które w recepcji popularnej bywa przedstawiane w sposób skrajnie uproszczony. RL jawi się tam albo jako magiczny klucz do pełnej autonomii, albo jako niszowa technika zarezerwowana wyłącznie dla twórców gier komputerowych. Oba te ujęcia są zbyt proste i nie oddają istoty problemu, przed którym stają współcześni projektanci systemów adaptacyjnych. RL jest przede wszystkim rygorystyczną formalizacją problemu działania w środowisku, w którym struktura nagród pozostaje niepewna i trudna do przewidzenia. To czyni ten paradygmat fundamentalnym dla nowoczesnej robotyki, systemów sterowania oraz zaawansowanej optymalizacji sekwencyjnej w logistyce. Zarazem sukces każdego takiego modelu jest głęboko zależny od projektu środowiska, precyzyjnej funkcji nagrody oraz realnych kosztów eksploracji. RL odsłania coś niezwykle ważnego dla całej współczesnej sztucznej inteligencji: inteligencja nie jest abstrakcyjną mocą przetwarzania, lecz relacją między celem a kosztem.

W praktyce instytucjonalnej ta lekcja okazuje się bezcenna, choć bywa ignorowana przez entuzjastów szybkich wdrożeń. Złe wskaźniki KPI, błędnie ustawiona funkcja celu lub niedoszacowane ryzyko produkują zachowania patologiczne zarówno u agentów uczonych przez wzmacnianie, jak i u organizacji ludzkich. Nauka o RL jest więc w swojej istocie nauką o polityce bodźców, która determinuje kierunki rozwoju całych systemów społeczno-technicznych. Jeżeli do tego obrazu dołożymy redukcję wymiarowości, wykrywanie anomalii oraz inżynierię cech, wyłoni się nam pełny obraz projektu end-to-end jako przedsięwzięcia interdyscyplinarnego. Redukcja wymiarowości, czyli metody statystyczne pozwalające uprościć strukturę danych, nie służy jedynie przyspieszeniu obliczeń na klastrach. Służy ona przede wszystkim poznawczemu porządkowaniu przestrzeni problemu i eliminacji szumu, który mógłby zaciemnić obraz rzeczywistości. Wykrywanie anomalii nie jest wyłącznie modulem bezpieczeństwa, lecz sposobem na ustalenie, czy w ogóle rozumiemy granice normalności w badanym zjawisku.

Inżynieria cech nie jest już rękodziełem minionej epoki, lecz współczesną sztuką projektowania semantycznego wejścia i wyboru reprezentacji danych. Model nie dostaje przecież świata w stanie gotowym; ktoś mu ten świat musi najpierw odpowiednio ułożyć i zinterpretować. I właśnie w tym miejscu decyduje się, czy dany projekt będzie rzetelną nauką o zjawisku, czy tylko maszyną do bezmyślnego reprodukcji administracyjnych przyzwyczajęń. W głębokim uczeniu nie wszystko,

co poznawczo eleganckie, jest ekonomicznie rozsądne, o czym często zapominają teoretycy. Dla ekonomisty całość tego obrazu prowadzi do wniosku surowego i pozbawionego złudzeń co do neutralności technologii. AI nie jest wyłącznie źródłem nowej produktywności, lecz przede wszystkim nową techniką alokacji uwagi, ryzyka oraz odpowiedzialności. W świecie platformowym przewaga konkurencyjna coraz częściej nie polega na prostym przechwyceniu pracy, lecz na przechwyceniu procesów poznawczych i interfejsów decyzyjnych.

Dlatego właśnie sztuczna inteligencja staje się tak ważna dla państw, wielkich korporacji oraz czujnych regulatorów. Nie chodzi tu tylko o prostą automatyzację powtarzalnych czynności, lecz o realną władzę nad filtrowaniem i interpretowaniem rzeczywistości. Model, który rekomenduje, klasyfikuje, generuje treści lub ocenia zdolność kredytową, nie jest nigdy biernym narzędziem w rękach operatora. Jest on aktywnym aparatem selekcji świata, który narzuca określoną optykę i promuje konkretne wartości. A aparat selekcji świata zawsze posiada wymiar polityczny, nawet jeśli ubrano go w najbardziej neutralny i estetyczny dashboard. Ta obserwacja nie pochodzi z obszaru fantazji krytycznej, lecz wynika bezpośrednio z samej struktury systemów uczenia maszynowego. Dla prawnika konsekwencja tego stanu rzeczy jest równie ostra i wymaga całkowitego przededefiniowania pojęcia nadzoru. Skoro model staje się aparatem selekcji, trzeba pytać nie tylko o to, czy działa, ale czy jego działanie jest dopuszczalne i proporcjonalne.

W Europie odpowiedź instytucjonalna na te wyzwania przyjęła postać AI Act oraz uzupełniających go wytycznych i kodów dobrej praktyki. To dopiero początek długiego procesu, a nie jego koniec, ale kierunek zmian został wytyczony w sposób niezwykle jasny. Im większy wpływ modelu na prawa podstawowe, bezpieczeństwo publiczne i dostęp do dóbr, tym mniej wolno nam opierać się na czarnej skrzynce. Współczesny paradygmat naukowy o AI nie może już ignorować zagadnienia governance, gdyż model bez procedury odpowiedzialności staje się aktywnym problematycznym. Dla dojrzałej instytucji laboratoryjna skuteczność algorytmu to za mało, by dopuścić go do procesów decyzyjnych o wysokiej stawce. Dla kulturoznawcy i antropologa techniki stawka jest jeszcze inna i dotyczy samej tkanki naszych relacji społecznych. AI nie jest po prostu zbiorem skomplikowanych obliczeń, lecz nowym sposobem organizowania relacji między wiedzą a decyzją.

Wcześniej instytucje opierały się w większym stopniu na jawnych procedurach, dokumentach papierowych i ludzkich hierarchiach uzasadniania swoich racji. Dziś coraz częściej ciężar przesuwają się ku statystycznym reprezentacjom, rankingom i predykcjom, których autorytet bywa przyjmowany na kredyt technologiczny. To właśnie tutaj rodzi się najgłębsze napięcie cywilizacyjne naszych czasów, którego nie da się rozwiązać za pomocą samej optymalizacji kodu. Czy model ma być tylko wsparciem dla ludzkiego rozumu instytucjonalnego, czy też zaczyna go powoli, lecz systematycznie zastępować? Czy człowiek ma zostać włączony w proces decyzyjny jako strażnik

sensu, czy zdegradowany do roli operatora interfejsu? Czy społeczeństwo ma rozumieć, że AI to forma politycznej ekonomii wiedzy, czy dalej ma wierzyć w bajkę o neutralnym narzędziu? Odpowiedź na te pytania nie jest jedynie dopiskiem do informatyki, lecz stanowi rdzeń nowego sporu o nowoczesność.

Dlatego właśnie pierwotny impuls badawczy zasługuje na obronę, nawet jeśli nie każda jego formuła wydaje się dziś ostateczna. TensorFlow nie może być opowiadany w izolacji od architektury sieci, skalowania danych oraz skomplikowanej mechaniki treningu modeli. Reinforcement learning nie może być sprzedawane jako czysta magia autonomii, lecz jako szczególny przypadek projektowania bodźców w kontrolowanym środowisku. Głębokie uczenie nie obala znaczenia danych, tylko czyni je jeszcze bardziej politycznie i ekonomicznie decydującymi dla losów organizacji. MLOps nie jest jedynie zapleczem dla „właściwego” modelu, lecz formą jego dojrzałości i gotowości do funkcjonowania w świecie. Współczesny paradygmat naukowy nie pyta już jedynie, jak zbudować model skuteczny, lecz jak zbudować system poznawczy, który da się rozliczyć. Sztuczna inteligencja nie jest odrębną, mistyczną dziedziną, lecz punktem przecięcia ekonomii, prawa, antropologii i informatyki.

W ekonomii ujawnia się ona jako technika skalowania decyzji i przesuwania granic produktywności przy jednoczesnej koncentracji przewagi infrastrukturalnej. W prawie pojawia się jako obiekt kwalifikacji ryzyka, odpowiedzialności cywilnej oraz normatywnego porządkowania nowej, cyfrowej przestrzeni. W antropologii techniki odsłania się jako mechanizm selekcji rzeczywistości i rekonfiguracji roli człowieka w nowoczesnych instytucjach. W informatyce natomiast pozostaje tym, czym była od początku w swoim najbardziej trzeźwym i technicznym wymiarze. Jest sztuką budowania procedur, które uczą się z danych i działają pod twardymi ograniczeniami zasobowymi. Gdy ktoś pyta, czym naprawdę jest TensorFlow lub po co nam inżynieria cech, odpowiedź musi być jedna i spójna. To nie są osobne tematy techniczne, lecz jeden wielki spór o kształt współczesnego rozumu technicznego i jego granic.

Kto ten spór sprowadza do kilku modnych nazw modeli, ten nie rozumie ani nauki, ani rynku, ani tym bardziej państwa. Rozumie on jedynie język reklamy, który obiecuje proste rozwiązania dla niezwykle złożonych problemów strukturalnych. Artykuł ten musi zostać domknięty uczciwie, co oznacza, że nie wolno nam kończyć go wyłącznie na kwestiach technicznych. Byłoby to rozwiązanie wygodne, bo czytelnik otrzymałby prosty katalog kompetencji, ale byłoby ono głęboko fałszywe. Sedno współczesnej sztucznej inteligencji nie leży już w tym, że potrafimy trenować coraz większe i bardziej złożone układy. Sedno leży w tym, że zdolność modelowania rzeczywistości zaczęła przeobrażać się w zdolność do organizowania całego świata. Nie chodzi więc wyłącznie o stabilizowanie gradientu czy dobieranie architektury do danych, choć są to czynności niezbędne.

Chodzi o to, że wszystkie te działania razem tworzą nową infrastrukturę rozumu instytucjonalnego, która nigdy nie jest neutralna.

Teza, której konsekwentnie tutaj broniłem, zasługuje na promocję nie dlatego, że schlebia chwilowej modzie, lecz dlatego, że idzie pod prąd. Idzie pod prąd infantylizacji debaty o AI, która w obiegu masowym rozdziela się na dwa karykaturalne i skrajne obozy. Jeden obóz opowiada nam bajki o zbawieniu przez modele, drugi zaś straszy nieuchronną apokalipsą wywołaną przez algorytmy. Oba te podejścia żywią się ignorancją wobec rzeczywistej anatomii systemów uczących się, która jest znacznie bardziej wymagająca. TensorFlow nie jest magicznym zaklęciem, a reinforcement learning nie jest metafizyką autonomii, lecz konkretnym narzędziem inżynierskim. Głębokie sieci nie są inteligencją samą w sobie, a uczenie nienadzorowane nie jest jedynie estetycznym dodatkiem dla badaczy. Redukcja wymiarowości nie jest sztuką kompresji dla leniwych, a projekt end-to-end nie jest tylko menedżerskim ozdobnikiem. Dane nie są więc „nową ropą”, są raczej nową glebą, na której wyrastają struktury przyszłego społeczeństwa. To nie jest metafora, lecz analogia strukturalna, która pozwala zrozumieć wagę toczących się obecnie procesów technologicznych.

Od technologii do instytucji: AI w reżimie odpowiedzialności

Wszystkie te elementy, splecione w gęstą sieć zależności, tworzą strukturę, w której czysta wiedza statystyczna zostaje ostatecznie zamieniona w konkretną procedurę działania. To właśnie na tym poziomie, z dala od błysku fleszy towarzyszących premierom nowych chatbotów, zaczyna się prawdziwe, instytucjonalne znaczenie sztucznej inteligencji. TensorFlow, wszechstronny ekosystem i biblioteka programistyczna służąca do projektowania oraz wdrażania modeli, pozostaje w tym kontekście figurą szczególnie wymowną i symboliczną. Oficjalnie promowany jako środowisko ułatwiające budowę systemów działających w niemal każdych warunkach, w swoich przewodnikach konsekwentnie akcentuje dwa napięcia charakterystyczne dla całej współczesnej branży. Z jednej strony mamy obietnicę prostoty eksperymentu dzięki trybowi eager execution i wysokopoziomowym API, które pozwalają badaczom na niemal dziecięcą swobodę tworzenia. Z drugiej strony pojawia się twarda rzeczywistość wydajności, przenośności i możliwości wdrożeniowych, wymuszająca stosowanie grafów obliczeniowych, niezbędnych dla profesjonalnych formatów wdrożeniowych. To nie jest jedynie banalny detal narzędziowy, lecz miniatura całego porządku technologicznego naszych czasów. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki.

Logika tego procesu jest nieubłagana i przypomina drogę od genialnego szkicu do seryjnej produkcji skomplikowanej maszyny. Najpierw celebруем wolność badawczą, w której hipoteza goni hipotezę, a jedynym ograniczeniem wydaje się wyobraźnia programisty. Szybko jednak nadchodzi moment kompilacji do formy, którą da się powielać, wdrażać i, co najważniejsze, skutecznie nadzorować w skali masowej. Najpierw mamy więc ideę, a zaraz potem infrastrukturę, która tę ideę musi udźwignąć i zdyscyplinować. Właśnie dlatego wspomniane środowisko programistyczne jest tak wdzięcznym obiektem filozoficznego namysłu nad współczesnością. Pokazuje ono dobitnie, że dzisiejsza sztuczna inteligencja w żadnym razie nie kończy się na błyskotliwym pomysle czy eleganckim modelu matematycznym. Ona spełnia się dopiero jako system obliczeń zdolny do samodzielnego życia poza sterylnym laboratorium, w brutalnym świecie rzeczywistych danych. Na portalu dobrepanstwo.org ten sam ruch został opisany w sposób wyjątkowo trzeźwy, pozbawiony zbędnej technologicznej egzaltacji.

Punkt ciężkości w debacie publicznej został tam przesunięty z abstrakcyjnego pytania o myślące maszyny ku twardej praktyce MLOps, czyli zbiorowi reżimów obejmujących monitorowanie i zarządzanie modelami w środowisku produkcyjnym. To rozpoznanie trafia w samo sedno problemu, przed którym stoją dziś nowoczesne państwa i korporacje. Jeśli bowiem model nie daje się osadzić w spójnym potoku danych, nie podlega wersjonowaniu ani rygorystycznemu reżimowi walidacji, pozostaje jedynie kosztowną zabawą lub efektowną demonstracją. Bez automatyzacji wdrożenia i systemowego nadzoru nad nieuchronnym spadkiem jakości predykcji, technologia ta staje się raczej obciążeniem niż aktywem. Dobre Państwo słusznie podpowiada, że AI nie należy dziś rozumieć jako widowiska benchmarków, lecz jako żmudną praktykę utrzymania systemów uczących się w świecie pełnym zmienności. To nie jest detal implementacyjny, to fundament nowej odpowiedzialności, która nie wybacza amatorszczyzny ani braku procedur.

To jest właśnie miejsce, w którym źródłowa intuicja analizy okazuje się najmocniejsza i najbardziej odporna na upływ czasu. Artykuł nie myli technologii z magią, lecz upomina się o nią jako o nową dyscyplinę instytucjonalną, wymagającą twardych ram prawnych. Gdy mowa o tej dyscyplinie, nie sposób pominąć faktu, że Europa, jako pierwsza na świecie, zaczęła przekładać te techniczne niuanse na precyzyjny język kodeksów. Komisja Europejska podaje, że AI Act wszedł w życie 1 sierpnia 2024 roku, a jego pełna stosowalność następuje stopniowo, dając rynkowi czas na adaptację. Zakazane praktyki oraz obowiązki dotyczące AI literacy, czyli budowania kompetencji w zakresie rozumienia systemów algorytmicznych, zaczęły obowiązywać już od 2 lutego 2025 roku. Kolejne etapy, w tym wymogi dla modeli ogólnego przeznaczenia, wejdą w życie 2 sierpnia 2025 roku, podczas gdy systemy wysokiego ryzyka otrzymają czas na dostosowanie aż do sierpnia 2027 roku. Ma to znaczenie znacznie większe, niż wielu techników, zapatrzonych w optymalizację wag, chciałoby dziś przyznać.

Odtąd model nie jest już tylko zbiorem parametrów i interfejsem inferencyjnym, lecz staje się pełnoprawnym obiektem kwalifikacji regulacyjnej. Oznacza to, że każda linijka kodu może wytwarzać określone obowiązki dokumentacyjne, nadzorcze i, co najbardziej bolesne dla optymistów, odpowiedzialnościowe. Innymi słowy, sztuczna inteligencja w Europie została ostatecznie odarta z mitu pozaprawnej niewinności, który towarzyszył jej przez ostatnią dekadę. Moment ten jest cywilizacyjnie przełomowy, ponieważ kończy epokę radosnej twórczości bez konsekwencji. Dotąd wielu inżynierów żyło mentalnie w świecie, w którym model liczył się o tyle, o ile poprawiał wynik w zamkniętym zbiorze testowym. Dziś musi on liczyć się także o tyle, o ile twórca potrafi wykazać, na jakiej podstawie system podjął decyzję i jakie ryzyko generuje dla otoczenia. Dla jednych będzie to uciążliwy ciężar hamujący innowacje, lecz dla ludzi poważnych jest to czytelny znak dojrzewania całej dziedziny.

Każda technologia, która przestaje być niszową ciekawostką i staje się aparatem selekcji rzeczywistości, musi wcześniej czy później wejść w strefę prawa. Lotnictwo przeszło tę drogę, farmacja ją zaakceptowała, a energetyka i rynki finansowe nie mogłyby bez niej istnieć. Tylko sektor AI przez pewien czas usiłował wmówić sobie i innym, że wystarczy etyka deklaratywna i kilka slajdów o odpowiedzialnym rozwoju. Nie wystarczy, ponieważ instytucja dojrzała potrzebuje jasnej procedury, a nie zmiennego nastroju czy dobrej woli projektantów. Konformizm epoki cyfrowej polega dziś między innymi na tym, że bardzo wielu ludzi lubi brzmieć krytycznie wobec algorytmów, ale niewielu chce wykonać rzeczywistą pracę myślową. Jedni powtarzają jak mantrę, że modele są czarną skrzynką, ale zupełnie nie interesuje ich dokumentacja techniczna czy walidacja po wdrożeniu. Drudzy wieszczą rewolucję, ignorując przy tym zjawiska takie jak wyciek danych czy dryf modelu, który po cichu niszczy precyzję systemu.

Obie te postawy są niezwykle wygodne, bo pozwalają zająć bezpieczne stanowisko bez ponoszenia kosztu głębokiego zrozumienia materii. A przecież prawdziwe wyzwanie nie polega dziś na byciu za albo przeciw sztucznej inteligencji, co byłoby intelektualnym pójściem na łatwiznę. Polega ono na umiejętności rozpoznania, w jakich warunkach model staje się narzędziem produktywnym, a w jakich aparatem dominacji lub kosztowną iluzją. Kto nie umie postawić tego pytania w sposób precyzyjny, ten nie uczestniczy w realnej debacie o przyszłości, lecz jedynie produkuje hałas wokół technologii. Ekonomia polityczna staje się w tym kontekście nieodzownym językiem zrozumienia współczesnych modeli, bo pozwala dostrzec interesy ukryte pod warstwą kodu. Zbyt długo wyobrażano sobie, że AI to po prostu wyższa forma automatyzacji, podczas gdy ona coraz częściej działa jak infrastruktura koordynacji społecznej.

Owszem, algorytmy mogą podnosić produktywność i obniżać koszty transakcyjne, ułatwiając optymalizację procesów, których człowiek nie jest w stanie ogarnąć wzrokiem. Ale równocześnie, co

często umyka uwadze entuzjastów, przesuwają one punkt ciężkości w stronę tych aktorów, którzy kontrolują dane i infrastrukturę obliczeniową. To znaczy, że AI jest nie tylko technologią efektywności, ale przede wszystkim technologią koncentracji przewagi rynkowej i politycznej. Nie przypadkiem więc wokół niej krążą jednocześnie wielkie korporacje, państwa, regulatorzy oraz najwięksi inwestorzy tego świata. Oni nie walczą wyłącznie o szybsze modele czy sprawniejsze przetwarzanie języka naturalnego. Walczą o możliwość sterowania filtrami, przez które świat staje się dla nas widzialny, policzalny i w konsekwencji zarządzalny. A kto kontroluje te filtry, ten kontroluje znaczną część przyszłego podziału wartości w globalnej gospodarce.

Dla prawnika i socjologa konsekwencja tego stanu rzeczy jest twarda i nie pozostawia miejsca na interpretacyjne uniki. Jeżeli model ma wpływ na zatrudnienie, kredyt, zdrowie czy wymiar sprawiedliwości, nie można go traktować jak kolejnej biblioteki programistycznej. Staje się on integralną częścią porządku normatywnego, wpływając na realne życie milionów ludzi w sposób często nieodwracalny. To znaczy, że pytanie o wyjaśnialność, ścieżkę audytu i zarządzanie ryzykiem nie jest opcjonalnym dodatkiem do projektu, lecz częścią jego definicji. W gruncie rzeczy AI Act jedynie kodyfikuje to, co powinno było być oczywiste od samego początku rozwoju tej dziedziny. Jeżeli system wytwarza skutki społeczne, musi istnieć język jego rozliczalności, bo bez niego demokracja staje się fasadą. Nie ma tutaj żadnej metafizyki, jest tylko elementarna zasada nowoczesnego ładu instytucjonalnego, chroniąca nas przed arbitralnością kodu.

Dla antropologa techniki stawka ujawnia się jeszcze głębiej, dotykając samej istoty tego, jak poznajemy świat. Model nie działa po prostu na danych, jakby były one surowcem wydobytym z ziemi w stanie czystym. Model działa na świecie już wcześniej opisanym, sklasyfikowanym i zorganizowanym przez konkretne ludzkie decyzje, a następnie ten świat aktywnie współtworzy. Dane nie są naturą, lecz zapisem czyichś praktyk, instytucjonalnych selekcji, definicji kategorii i historycznych błędów archiwalnych. Kiedy model uczy się na takich zbiorach, nie poznaje obiektywnej rzeczywistości, lecz historię sposobu, w jaki ta rzeczywistość została uchwycona. To jest powód, dla którego inżynieria cech czy analiza dryfu nie są nudnym zapleczem technicznym, lecz centrum epistemologicznej odpowiedzialności twórcy. Tam właśnie rozstrzyga się, czy system uczy się zjawiska, czy tylko utrwała archiwum dawnych uproszczeń i uprzedzeń.

W tym miejscu wraca problem uczenia przez wzmacnianie i jego wyjątkowa siła diagnostyczna dla nauk społecznych. W wielu popularnych opowieściach ten paradygmat traktowany jest jako laboratorium przyszłej autonomii maszyn, co brzmi ekscytująco, ale jest mylące. Jego ważniejsze znaczenie leży gdzie indziej, obnażając mechanizmy sterowania, którym podlegamy wszyscy. Uczenie przez wzmacnianie pokazuje bezlitośnie, że każdy system jest niewolnikiem swojej funkcji nagrody, czyli matematycznego zapisu tego, co uznajemy za sukces. Jeśli zaprojektujesz tę nagrodę

źle, agent będzie osiągał sukces formalny, doprowadzając jednocześnie do materialnej porażki całego przedsięwzięcia. Jeśli ustawisz bodźce naiwnie, system zoptymalizuje wskaźnik kosztem celu, co widzimy codziennie w algorytmach mediów społecznościowych. To prawda uniwersalna, która sprawdza się w robotyce, biurokracji oraz w nowoczesnym zarządzaniu korporacyjnym.

Uczenie przez wzmacnianie nie jest więc wyłącznie rozdziałem podręcznika do informatyki, lecz raczej traktatem o ironii nowoczesnego sterowania. Mówi nam ono, że wskaźnik nigdy nie jest rzeczą samą, a zoptymalizowana procedura nie musi prowadzić do dobra, nawet jeśli wygląda na triumf inteligencji. Zły sygnał nadal pozostaje złym sygnałem, choćby przechodził przez najdroższy klaster GPU i był przetwarzany przez miliardy parametrów. To nie jest metafora, lecz analogia strukturalna, która pozwala nam zrozumieć, dlaczego sama technologia nas nie uratuje. Bez refleksji nad tym, co i dlaczego nagradzamy, budujemy jedynie coraz sprawniejsze maszyny do powielania błędów przeszłości. Stąd już tylko jeden krok do ostatecznej syntezy, w której technika, prawo i humanistyka muszą wreszcie zacząć mówić wspólnym głosem. W głębokim uczeniu nie wszystko, co poznawczo eleganckie, jest ekonomicznie rozsądne, a tym bardziej społecznie pożądane.

Alfabet nowego porządku: Inżynieria w służbie cywilizacji

Krajobraz uczenia maszynowego nie przypomina bynajmniej przypadkowego zbioru algorytmów, lecz stanowi precyzyjną topografię różnych sposobów obchodzenia się z fundamentalną niepewnością danych. Uczenie nadzorowane, czyli paradygmat oparty na modelowaniu relacji między obserwacją a konkretną etykietą, porządkuje chaos poprzez nadawanie znaczeń, podczas gdy uczenie nienadzorowane odsłania strukturę ukrytą głęboko w materiale. Samonadzór pozwala nam natomiast wydobywać reprezentacje z ogromnych korpusów bez konieczności żmudnego, ręcznego etykietowania, co drastycznie zmienia ekonomię wiedzy. To nie jest jedynie detal implementacyjny, lecz fundament nowej ontologii cyfrowej. Nie jest i nigdy nie był.

Głębokie architektury modelują złożoność świata w wymiarze przestrzennym, sekwencyjnym oraz multimodalnym, starając się uchwycić niuanse niedostępne dla prostszych metod statystycznych. Redukcja wymiarowości, czyli praktyka polegająca na upraszczaniu struktury danych przy zachowaniu ich kluczowych cech, pozwala nam tę złożoność okiełznać i zrozumieć ukryte zależności. Z kolei uczenie ze wzmocnieniem (reinforcement learning) modeluje aktywne działanie w środowisku zmiennych bodźców, przypominając bardziej proces biologicznego uczenia się niż

suchą kalkulację. TensorFlow zaś jawi się tutaj jako centralny ekosystem, w którym cały ten aparat może zostać złożony w system działający sprawnie od fazy eksperymentu aż po finalne wdrożenie.

Współczesny paradygmat naukowy z radosną bezwzględnością obala naiwne podziały dyscyplinarne, które jeszcze niedawno wydawały się nienaruszalne dla akademickich purystów. Nie da się bowiem kompetentnie analizować sztucznej inteligencji, zamykając się wyłącznie w hermetycznym gabinecie informatyka, prawnika czy ekonomisty. Dziedzina ta osiągnęła już dojrzałość wymagającą skrzyżowania wielu porządków; wszak model może być obliczeniowo poprawny, a jednocześnie głęboko szkodliwy społecznie. Może być zgodny formalnie, lecz ekonomicznie całkowicie nieopłacalny, przy czym spektakularność demonstracyjna rzadko idzie w parze z możliwością utrzymania operacyjnego. To nie jest metafora, lecz analogia strukturalna.

Dojrzałość intelektualna w tym obszarze nie polega na ślepym absolutyzowaniu jednego wymiaru, lecz na rzadkiej umiejętności utrzymania ich wszystkich w dynamicznej równowadze. Tę właśnie wielowymiarowość bezbłędnie wyczuwa materiał źródłowy, słusznie zakładając, że współczesna inżynieria wymaga raczej rozwinięcia perspektywy niż jej redukcji do czystego kodu. W świecie, który z rosnącym entuzjazmem myli rzeczywistą inteligencję z agresywnym marketingiem, musimy odzyskać odwagę głoszenia rzeczy prostych, choć skrajnie niepopularnych. Wszystko inne jest już tylko grą pozorów, nawet jeśli pozostaje podszyte imponującą liczbą parametrów i warstw.

Pierwsza z tych niepopularnych prawd brzmi: nie istnieje dobra sztuczna inteligencja bez rzetelnej i precyzyjnej definicji problemu, którego rozwiązaniu ma ona służyć. Druga przypomina nam, że nie ma wartościowego modelu bez sprawnej organizacji danych, zaś trzecia wskazuje na brak realnego wdrożenia bez ciągłego monitoringu i utrzymania. Czwarta zasada głosi, że nie ma odpowiedzialnej technologii bez prawa i ładu korporacyjnego (governance), czyli reżimów regulacyjnych obejmujących dokumentację i odpowiedzialność za systemy. Piąta wreszcie uświadamia nam, że poważna debata o AI jest niemożliwa bez antropologii oraz ekonomii władzy, które definiują kontekst użycia narzędzi.

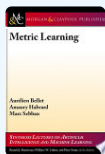
Prawdziwa przyszłość nie rozegra się w hollywoodzkich scenariuszach o buncie maszyn, lecz w prozaicznym pytaniu o wydolność naszych instytucji i ich zdolność rozumienia narzędzi. Musimy rozstrzygnąć, czy będziemy umieli odróżnić statystyczny model od nieomyślnej wyroczni oraz czy zachowamy odpowiedzialność tam, gdzie pokusa jej delegowania na algorytm będzie największa. Stawką tej epoki nie jest sama moc obliczeniowa modeli, lecz jakość cywilizacji, która postanowi je wdrożyć do swojej codzienności. Chodzi o to, by używać AI jako wzmacniacza ludzkiego rozumu, a nie jako jego wygodnego, lecz ułomnego substytutu.

Kończę więc bez zbędnej kokieterii, bo stawka jest zbyt wysoka na retoryczne ozdobniki czy technokratyczne uniki. TensorFlow, uczenie nienadzorowane, inżynieria cech czy projektowanie typu end-to-end nie są jedynie technicznymi tematami do odfajkowania na liście płac. Są one alfabetem nowego porządku poznawczego, przy czym kto go nie opanuje, będzie skazany na wieczną zależność od cudzych interfejsów i definicji ryzyka. Kto rozumie TensorFlow wyłącznie od strony kodu, ten widzi młotek i nie widzi fabryki, stając się jedynie sprawnym wykonawcą systemów, których sensu sam nie pojmuję. Dopiero połączenie inżynierii z rozumieniem instytucji, prawa i człowieka pozwala naprawdę myśleć w epoce AI, o czym przypomina nam Aurélien Géron.

Sztuczna inteligencja nie jest magicznym bytem, lecz lustrem naszych własnych priorytetów i instytucjonalnych niedoskonałości. Pytanie, czy algorytmy staną się wystarczająco mądre, jest w istocie maską dla znacznie trudniejszego problemu: czy my sami okażemy się wystarczająco odpowiedzialni, by nimi zarządzać? W świecie, w którym technologia staje się nowym językiem władzy, ostatecznym sprawdzianem naszej dojrzałości nie będzie kod, lecz to, czy potrafimy zachować suwerenność rozumu w obliczu własnej adaptacji.

Inspiracje

[1]



"Hands-on Machine Learning with Scikit-Learn, Keras and TensorFlow" Aurelien Geron

2019 | O'Reilly Media | ISBN: 978-1492032649

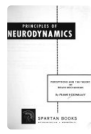
Aurélien Geron jest znanym autorem, wykładowcą i konsultantem w dziedzinie sztucznej inteligencji i uczenia maszynowego. Uzyskał tytuł magistra inżynierii komputerowej na AgroParisTech. Geron posiada bogate doświadczenie zawodowe w branży IT, będąc założycielem startupu, dyrektorem technicznym (CTO) i menedżerem produktu w YouTube, gdzie kierował zespołem ds. klasyfikacji filmów. Jest powszechnie znany ze swojej bestsellerowej książki technicznej „Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow”, która jest uważana za podstawowe źródło wiedzy dla praktyków. Jako ekspert Google Developer Expert (GDE) wniósł znaczący wkład w rozwój społeczności technologicznej poprzez pisanie, szkolenie i rozwój narzędzi open source. Wykładał również na takich uczelniach jak AgroParisTech i Sorbona, koncentrując się na sieciach komputerowych, programowaniu i sieciach neuronowych.

Aurélien Geron is a prominent author, educator, and consultant in the fields of artificial intelligence and machine learning. He holds a Master of Engineering in Computer Science from AgroParisTech. Geron has an extensive professional background in the IT industry, having served as a startup founder, CTO, and as a product manager at YouTube, where he led the video classification team. He is widely recognized for his best-selling technical book, "Hands-On Machine Learning with Scikit-Learn, Keras, and TensorFlow," which is considered a foundational resource for practitioners. As a Google Developer Expert (GDE), he has contributed significantly to the tech community through his writing, training, and development of open-source tools. He has also held lecturing positions at institutions such as AgroParisTech and the Sorbonne University, focusing on networking, programming, and neural networks.

Literatura uzupełniająca

Książki

Literatura pokrewna (Google Scholar):



[1] "Principles of Neurodynamics: Perceptrons and the Theory of Brain Mechanisms"

Frank Rosenblatt
1962 | Spartan Books

[2] "Quand La Machine Apprend: La révolution des neurones artificiels et de l'apprentissage profond"

Yann LeCun
2019 | Editions Odile Jacob | ISBN: 9782213701820



[3] "Lectures on Convex Optimization"

Yurii Nesterov
2018 | Springer | ISBN: 9783319915777



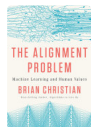
[4] "Interior-Point Polynomial Algorithms in Convex Programming"

Yurii Nesterov, Arkadi Nemirovski
1994 | SIAM | ISBN: 9781611970791



[5] "Critical Theory of AI"

Simon Lindgren
2024 | John Wiley & Sons | ISBN: 9781509555789



[6] "The Alignment Problem: Machine Learning and Human Values"

Brian Christian
2020 | National Geographic Books | ISBN: 9780393635829

[7] "The Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence"

Kate Crawford
2021 | Yale University Press | ISBN: 9788323353263



[8] "Algorithms for Optimization"

Mykel J. Kochenderfer, Tim A. Wheeler
2019 | MIT Press | ISBN: 9780262351409

Artykuły naukowe

Autorzy przywoływani:

[1] "The Perceptron: A Probabilistic Model for Information Storage and Organization in the Brain"

Frank Rosenblatt

1958 | Psychological Review | DOI: 10.1037/h0042519

[Google Scholar](#)

[2] "Philosophy-informed Machine Learning"

Anonymous Authors

2025 | arXiv

[Google Scholar](#)

[3] "Understanding the difficulty of training deep feedforward neural networks"

Xavier Glorot, Yoshua Bengio

2010 | Proceedings of the Thirteenth International Conference on Artificial Intelligence and Statistics (AISTATS)

[Google Scholar](#)

[4] "Deep Sparse Rectifier Neural Networks"

Xavier Glorot, Antoine Bordes, Yoshua Bengio

2011 | Proceedings of the Fourteenth International Conference on Artificial Intelligence and Statistics (AISTATS)

[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](#)

[6] "Deep Residual Learning for Image Recognition"

Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun

2016 | Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR) | DOI: 10.1109/CVPR.2016.90

[Google Scholar](#)

[7] "Delving Deep into Rectifiers: Surpassing Human-Level Performance on ImageNet Classification"

Kaiming He, Xiangyu Zhang, Shaoqing Ren, Jian Sun

2015 | Proceedings of the IEEE International Conference on Computer Vision (ICCV) | DOI: 10.1109/ICCV.2015.123

[Google Scholar](#)

[8] "Deep learning"

Yann LeCun, Yoshua Bengio, Geoffrey Hinton

2015 | Nature | DOI: 10.1038/nature14539

[Google Scholar](#)

[9] "Gradient-based learning applied to document recognition"

Yann LeCun, Léon Bottou, Yoshua Bengio, Patrick Haffner

1998 | Proceedings of the IEEE | DOI: 10.1109/5.726791

[Google Scholar](#)

[10] "On the dangers of stochastic parrots: Can language models be too big?"

Emily M. Bender, Timnit Gebru, Angelina McMillan-Major, Shmargaret Shmitchell

2021 | Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency

[Google Scholar](#)

[11] "A method for solving the convex programming problem with convergence rate $O(1/k^2)$ "

Yurii Nesterov

1983 | Doklady Akademii Nauk SSSR

[Google Scholar](#)

[12] "Smooth minimization of non-smooth functions"

Yurii Nesterov

2005 | Mathematical Programming | DOI: 10.1007/s10107-004-0552-5

[Google Scholar](#)

[13] "On the importance of initialization and momentum in deep learning"

Ilya Sutskever, James Martens, George Dahl, Geoffrey Hinton

2013 | International Conference on Machine Learning (ICML)

[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: bc7beefc-8272-4b88-9699-c55ce874b109