



Politechnika
Wrocławska

Sztuczna inteligencja

Technologie Informacyjne

Wojciech Myszka

Katedra Mechaniki, Inżynierii Materiałowej i Biomedycznej

wrzesień/październik 2025



unite! University Network for Innovation
Technology and Engineering



HR EXCELLENCE IN RESEARCH

Wstęp

1. Właściwie nie warto o niej mówić — jest na ustach wszystkich
2. Powszechnie (???) używana, zazwyczaj bez głębszej refleksji
3. Nieodłącznie związana z komputerami (co nawet miało odzwierciedlenie we wczesnych nazwach komputera):
 - ▶ mózg elektronowy
 - ▶ a po angielsku *electronic brain* lub czasami *giant brain*

Trochę historii I

1. Golem *istota utworzona z gliny na kształt człowieka, ale pozbawiona duszy rozumiejącej nieszama, a zatem również zdolności mowy.*

Pojawia się w Starym Testamencie, nieodłącznie związany z Żydowskim folklorem i Talmudem

2. 1738 Jaques Vaucanson prezentuje trzy roboty: *Dobosz, Flecista i Trawiąca kaczką.*
3. 1769 **Wolfgang von Kempelen** zaprezentował *Mechanicznego Turka* — maszynę, która miała grać w szachy (na podobnej zasadzie działał późniejszy *Ajeeb* czy *Mephisto*).
4. 1818 **Mary Shelley** publikuje powieść *Frankenstein albo: Współczesny Prometeusz*

Trochę historii II

5. 1843 **Ada Lovelace** tłumaczy artykuł „**Szkic maszyny analitycznej**” wynalezionej przez Charlesa Babbage’a. W notatkach do tłumaczenia opracowuje to, co uważa się za **pierwszy algorytm**, dzięki któremu maszyna Babbage’a mogłaby obliczać pełny ciąg liczb Bernoulliego.
6. 1912 **El Ajedrecista** automat potrafiący rozgrywać z człowiekiem końcówki z trzema figurami (działał algorytmicznie).
7. 1913 **Bertrand Russell** i **Alfred North Whitehead** publikują *Principia Mathematica*, przetłomowe dzieło na temat logiki formalnej i rozumowania.
8. 1923 W sztuce R.U.R. Karla Čapka pojawia się termin „robot”.

Trochę historii III

9. (lata 40 XX w.) Isaac Asimov publikuje opowiadania SF wydane później jako zbiór pod tytułem **Ja, Robot**. Sformułowane tam są *Prawa robotyki*.
10. 1943 pierwszy artykuł autorstwa Arturo Rosenbluetha, Norberta Wienera i Juliana Bigelowa *Behavior, Purpose and Teleology*, który inicjuje dziedzinę **cybernetyka**.
11. 1948–1949 **William Grey Walter** buduje autonomicznego robota w kształcie żółwia.
12. (1945–1950)
Alan Turing pracuje w Bletchley Park nad automatyzacją procesu dekodowania szyfrów używanych przez Niemców.

Trochę historii IV

W 1950 roku Alan Turing wygłasza wykład „Inteligentne maszyny, teoria heretycka” na temat potencjalnych problemów automatyzacji procesów racjonalnych. Alan Turing publikuje książkę „Maszyny obliczeniowe i inteligencja” oraz wprowadza test Turinga.

13. (lata 50-te) **John McCarthy** opracowuje język **LISP** dla komputera IBM 704. W 1960 roku produkowana jest seria IBM 700-7000, kluczowe systemy w informatyce i wczesnym programowaniu sztucznej inteligencji.

Trochę historii V

14. 1956

Claude Shannon opracowuje **Tezeusza**, cybernetyczną mysz rozwiązującą labirynt. W latach 1948 i 1958 Claude Shannon przedstawia artykuł na temat teorii informacji.

W 1956 r. konferencja w Dartmouth, zorganizowana przez Johna McCarthy'ego, Marviną Minsky'ego i Claude'a Shannona, dała początek narodzinom sztucznej inteligencji jako odrębnej dyscypliny naukowej.

15. 1957 Frank Rosenblatt opierając się na idei sieci neuronowej (opracowanej w 1943 roku) projektuje sieć (**Perceptron**), która będzie mogła się uczyć i symuluje jej działania na komputerze. Później powstaje fizyczne urządzenie...

Trochę historii VI

16. 1959 **John McCarthy** i **Marvin Minsky** zakładają Projekt Sztucznej Inteligencji w ramach Laboratorium Badawczego Elektroniki i Centrum Obliczeniowego na MIT.
17. 1960 **James L. Adam** buduje *Stanford Cart*, prototyp samochodu autonomicznego.
18. 1961 W General Motors zaprezentowano *Unimate*, pierwszego robota przemysłowego.
19. 1966 *Shakey*, pierwszy robot ogólnego przeznaczenia, został zaprezentowany w Laboratorium Sztucznej Inteligencji Uniwersytetu Stanforda. Prowadząc badania nad Shakeyem, Helen Chan Wolf dokonuje przełomowych odkryć w dziedzinie kartografii i rozpoznawania obrazów, będąc pionierem w dziedzinie rozpoznawania twarzy.

Trochę historii VII

20. 1966 powstaje program **ELIZA** czat symulujący psychoanalityka napisany przez **Josepha Weizenbauma**.
21. 1968 **Stanisław Lem** w opowiadaniu SF *Rozprawa* opisuje sytuację związaną z wykorzystaniem humanoidalnego robota wyposażonego w sztuczną inteligencję jako pilota statku kosmicznego.
22. 1969 Film **Stanleya Kubricka** *2001: Odyseja kosmiczna* przedstawia HAL-a 9000, złowrogą sztuczną inteligencję. Marvin Minsky był konsultantem przy tym filmie.
23. 1971 **Nolan Bushnell** i **Ted Dabney** tworzą *Computer Space*, pierwszą grę zręcznościową wykorzystującą sztuczną inteligencję.

Trochę historii VIII

24. 1971 **Stanisław Lem** w opowiadaniu SF *Ananke* przewiduje zjawisko zwane dziś *bias* (stronniczość) związane z uczeniem (trenowaniem) sztucznej inteligencji (które powoduje katastrofę rakiety lądującej na Marsie).

25. 1981

13% gospodarstw domowych w Wielkiej Brytanii posiada komputer domowy. Do 2017 roku odsetek ten gwałtownie wzrośnie do 88%.

IBM wypuszcza pierwszy komputer osobisty IBM PC.

Trochę historii IX

26. 1982 **Ridley Scott** adaptuje powieść Philipa K. Dicka „Czy androidy śnią o elektrycznych owcach?” (1968) w filmie „**Łowca androidów**” z Harrisonem Fordem w roli głównej. Książka odzwierciedla rosnący wpływ maszyn i informacji na koncepcję tożsamości i człowieczeństwa.
27. 1989 Angielski naukowiec **Tim Berners-Lee** wymyślił sieć WWW w 1989 roku. Uruchomienie sieci WWW znacząco zwiększyło ilość danych dostępnych do badań nad uczeniem maszynowym.
28. 1990 **Rodney Brooks** publikuje artykuł „Słonie nie grają w szachy” i zostaje dyrektorem MIT AI LAB. Brooks ożywia badania nad sieciami neutralnymi.

Pierwsze udane testy rozpoznawania mowy w Instytucie Badawczym Stanforda.

Trochę historii X

29. 1999 **Bracia Wachowscy**/Siostry Wachowskie wyreżyserowali/ły film *Matrix*, przedstawiający walkę o władzę między maszynami a ludzkością.
30. 1999 Sony wprowadza na rynek Aibo, pierwszego dostępnego w sprzedaży robota-psa.
31. 2000 **Cynthia Breazeal** opracowuje *Kismet*, program i robota, które potrafią rozpoznawać ludzkie emocje poprzez analizę twarzy.
32. 2010 **Hiroshi Ishiguro** i Intelligent Robots Lab przedstawiają androida Geminoid.
33. 2015

Trochę historii XI

- ▶ Future of Life Institute publikuje list otwarty „Priorytety badawcze w zakresie solidnej i korzystnej sztucznej inteligencji” podpisany przez 8000 osób, w tym Steve’a Wozniaka, Stephena Hawkinga i Elona Muska.
 - ▶ Future of Life Institute publikuje dokument „Broń autonomiczna: list otwarty od badaczy sztucznej inteligencji i robotyki”, poparty przez prawie 4000 badaczy zajmujących się sztuczną inteligencją i robotyką.
34. 2015 ERIKA, androidowa prezenterka wiadomości, jest prezentowana przez **Hiroshi Ishiguro** i Intelligent Robots Lab z Uniwersytetu w Osace.
35. 2016 Program AlphaGo firmy Google DeepMind pokonuje mistrza Lee Sedola w grze Go Go.

Trochę historii XII

36. 2018

- ▶ Unia Europejska powołuje ELLIS, Europejskie Laboratorium Systemów Uczenia Się i Inteligentnych.
- ▶ Sztuczna inteligencja przetwarzająca język firmy Alibaba uzyskała lepsze wyniki niż człowiek w teście czytania i rozumienia tekstu na Uniwersytecie Stanforda.

Co się składa na Sztuczną Intelgencję?

1. Uczenie maszynowe (ML *Machine Learning*):
 - 1.1 Uczenie z trenerem — *nadzorowane (supervised learning)*;
 - 1.2 Uczenie nienadzorowane (*unsupervised learning*);
 - 1.3 Uczenie posiltkowane — *przez wzmacnianie (reinforcement learning)*.
2. Generatywna SI (*Generative AI*).
3. Przetwarzanie języków naturalnych (NLP *Natural Language Processing*).
4. Systemy ekspertowe (*Expert Systems*).

Na czym polega nadzorowane nuczanie maszynowe I

Nadzorowane uczenie maszynowe (*Supervised Learning*) to metoda w uczeniu maszynowym, w której algorytm jest trenowany na **etykietowanych danych**. Oznacza to, że każdy **wejściowy punkt danych** (np. obraz, fragment tekstu, rekord klienta) ma przypisaną **poprawną odpowiedź** (etykietę lub wartość wyjściową), która działa jak „nauczyciel” nadzorujący proces uczenia.

Jak to działa?

Proces ten polega na tym, że model uczy się **mapować funkcję** z danych wejściowych na dane wyjściowe.

1. **Dane Etykietowane:** Algorytm otrzymuje duży zestaw danych, w którym każda próbka jest parą: (dane wejściowe, oczekiwany wynik).
 - ▶ *Przykład:* (zdjęcie kota, etykieta „kot”), (cechy mieszkania, cena sprzedaży).
2. **Trening:** Model analizuje te pary i iteratively dostosowuje swoje wewnętrzne parametry (wagi), aby zminimalizować różnicę (błąd) między swoimi **przewidywaniami** a **poprawnymi etykietami** (oczekiwanymi wynikami). Jest to ciągłe korygowanie, przypominające ucznia, który sprawdza swoje odpowiedzi z kluczem.
3. **Generalizacja:** Celem jest nauczenie modelu **ogólnych wzorców** i zależności, a nie tylko zapamiętywania danych treningowych. Po

Główne rodzaje zadań

Nadzorowane uczenie maszynowe jest stosowane do rozwiązywania dwóch głównych typów problemów:

Typ Zadań	Opis	Przykład
Klasyfikacja	Przewidywanie dyskretnej (kategorycznej) etykiety lub klasy.	Określenie, czy e-mail to SPAM czy NIE SPAM (klasyfikacja binarna), lub rozpoznanie, czy na zdjęciu jest pies , kot czy królik (klasyfikacja wieloklasowa).
Regresja	Przewidywanie ciągłej wartości numerycznej .	Przewidywanie ceny domu na podstawie jego wielkości i lokalizacji, lub prognozowanie temperatury na jutro

Na czym polega nienadzorowane nauczanie maszynowe?

Nienadzorowane uczenie maszynowe (*Unsupervised Learning*) to metoda, w której algorytm uczy się na podstawie **nieetykietowanych danych**, czyli danych bez z góry określonych poprawnych odpowiedzi. Jego celem jest samodzielne **odkrywanie ukrytych wzorców, struktur i relacji** w zbiorze danych, aby go uporządkować i zrozumieć.

Jest to jak sytuacja, w której dajesz algorytmowi stos zdjęć (np. zwierząt, owoców, budynków) i prosisz, by je posegregował, ale nie mówisz mu, co przedstawiają. Algorytm musi sam znaleźć podobieństwa i różnice, aby stworzyć naturalne grupy.

Jak działa nienadzorowane uczenie?

W odróżnieniu od uczenia nadzorowanego, gdzie model koryguje swoje błędy na podstawie znanych etykiet, w uczeniu nienadzorowanym nie ma „nauczyciela”. Model sam eksploruje dane:

1. **Brak etykiet:** Algorytm otrzymuje surowe dane wejściowe (np. dane demograficzne klientów, wyniki pomiarów). Nie ma żadnej kolumny mówiącej, do jakiej grupy ten klient należy ani co oznaczają te pomiary.
2. **Odkrywanie wzorców:** Algorytm stosuje techniki matematyczne i statystyczne, aby zidentyfikować podobieństwa i różnice między punktami danych. Szuka naturalnych skupisk lub sposobów na uproszczenie danych.
3. **Tworzenie struktury:** Model nie przewiduje wartości, ale raczej transformuje dane lub je organizuje w logiczne struktury.

Główne zadania I

Nienadzorowane uczenie maszynowe jest wykorzystywane do rozwiązywania problemów eksploracyjnych:

1. Klasteryzacja (grupowanie)

Polega na grupowaniu podobnych do siebie punktów danych w **klastry** (grupy).

- ▶ **Cel:** Znalezienie naturalnych podziałów w danych, o których wcześniej nie wiedzieliśmy.
- ▶ **Przykład: Segmentacja klientów.** Algorytm analizuje zachowania zakupowe klientów i automatycznie grupuje ich w segmenty (np. „oszczędni łowcy okazji”, „klienci premium”), co pozwala firmie lepiej targetować reklamy. Najpopularniejsze algorytmy to **K-Means** i **Klasteryzacja Hierarchiczna**.

Główne zadania II

2. Redukcja wymiarowości

Polega na zmniejszeniu liczby cech (zmiennych) w zbiorze danych przy zachowaniu jak największej ilości istotnych informacji.

- ▶ **Cel:** Uproszczenie skomplikowanych danych, redukcja szumu i przyspieszenie obliczeń.
- ▶ **Przykład:** Analiza obrazu. Obraz cyfrowy może mieć tysiące pikseli (wymiarów). Redukcja wymiarowości pozwala zidentyfikować tylko te kluczowe cechy pikseli, które są ważne do rozpoznania obiektu, jednocześnie eliminując zbędny szum. Najczęściej używany algorytm to **PCA** (Principal Component Analysis).

Główne zadania III

3. Reguły asocjacyjne

Polegają na znajdowaniu zależności i reguł między zmiennymi w dużych zbiorach danych transakcyjnych.

- ▶ **Cel:** Odkrywanie, które elementy są często kupowane razem.
- ▶ **Przykład: Analiza koszyków zakupowych.** Odkrycie reguły typu: „Jeśli klient kupuje chleb i masło, to z 80% prawdopodobieństwem kupi też ser”. Pozwala to na optymalizację układu towarów w sklepie lub tworzenie spersonalizowanych rekomendacji online.

Na czym polega reinforcement learning? I

Uczenie ze wzmocnieniem (*Reinforcement Learning, RL*) to paradygmat uczenia maszynowego, w którym **agent** (program komputerowy) uczy się podejmowania optymalnych **decyzji** poprzez **interakcję z otoczeniem** w celu **maksymalizacji skumulowanej nagrody** w długim okresie.

Proces ten jest wzorowany na psychologii behawioralnej i uczeniu się metodą prób i błędów, podobnie jak szkolenie zwierzęcia lub nauka gry wideo.

Podstawowa zasada działania (pętla)

RL polega na ciągłej interakcji między agentem a środowiskiem:

1. **Obserwacja stanu (State):** Agent obserwuje bieżący stan otoczenia (np. położenie w grze, dane z czujników robota).

Na czym polega reinforcement learning? II

2. **Podjęcie działania (Action):** Agent podejmuje decyzję, jakie działanie wykonać na podstawie swojej aktualnej **polityki** (strategii).
3. **Nowy stan i nagroda (Reward):** Otoczenie reaguje na działanie, przechodząc do nowego stanu i zwracając **sygnał nagrody** (pozytywny) lub **kary** (negatywny).
4. **Uczenie:** Agent wykorzystuje otrzymaną nagrodę do **poprawy swojej polityki**, ucząc się, które działania w danym stanie prowadzą do wyższych nagród w przyszłości.

Celem agenta jest nie tylko zdobycie natychmiastowej nagrody, ale **maksymalizacja sumy nagród w perspektywie długoterminowej** (skumulowanej nagrody).

Kluczowe elementy I

System uczenia ze wzmocnieniem składa się z następujących elementów:

- ▶ **Agent** (Learner/Decyzja): To jest algorytm lub model, który uczy się i podejmuje decyzje.
- ▶ **Środowisko** (Environment): Świat, w którym agent działa (np. gra komputerowa, symulacja giełdy, świat fizyczny dla robota). Środowisko określa możliwe stany i zasady reakcji.
- ▶ **Stan** (State): Bieżąca sytuacja w środowisku, którą agent obserwuje.
- ▶ **Akcja** (Action): Wybór, jakiego agent dokonuje w danym stanie.
- ▶ **Nagroda** (Reward): Liczba, która jest natychmiastową informacją zwrotną od środowiska, wskazującą, jak dobre było ostatnie działanie.

Kluczowe elementy II

- ▶ **Polityka** (Policy): Strategia agenta; funkcja, która mapuje stany na akcje (określa, jakie działanie agent powinien podjąć w danym stanie).
- ▶ **Funkcja Wartości** (Value Function): Estymacja przewidywanej **skumulowanej nagrody**, którą agent może uzyskać, zaczynając od danego stanu i postępując zgodnie z polityką. Pomaga agentowi ocenić długoterminową opłacalność stanów, nie tylko natychmiastową nagrodę.

Eksploracja vs. Eksploatacja

Kluczowym wyzwaniem w RL jest znalezienie równowagi między:

- ▶ **Eksploatacja:** Wybieranie działania, które agent uważa za najlepsze (najbardziej nagradzające) na podstawie dotychczasowej wiedzy.
- ▶ **Eksploracja:** Wybieranie nowego, nieznanego działania w celu zebrania dodatkowych informacji o środowisku, co może prowadzić do odkrycia lepszej strategii.

Agent musi wystarczająco dużo eksplorować, aby znaleźć optymalną ścieżkę, ale też wykorzystywać zdobytą wiedzę (eksploatować), aby maksymalizować nagrodę.

Przykłady zastosowań

Uczenie ze wzmocnieniem jest idealne do problemów, które wymagają podejmowania **sekwencyjnych decyzji**:

- ▶ **Gry komputerowe:** Agenci AlphaGo i AlphaZero, grający w Go, szachy.
- ▶ **Robotyka:** Nauka robotów chwytania, chodzenia lub nawigacji w złożonych środowiskach.
- ▶ **Autonomiczne pojazdy:** Podejmowanie decyzji o przyspieszeniu, hamowaniu i skręcaniu w dynamicznym ruchu.
- ▶ **Zarządzanie zasobami:** Optymalizacja zużycia energii w centrach danych (np. DeepMind użyło RL do optymalizacji chłodzenia w Google).
- ▶ **Finanse:** Tworzenie algorytmów handlowych na giełdzie.

Podstawy działania generatywnej AI I

Generatywna sztuczna inteligencja to rodzaj AI, która jest zdolna do tworzenia nowych treści, takich jak tekst, obrazy, muzyka, a nawet kod. W przeciwieństwie do tradycyjnych systemów AI, które są projektowane do wykonywania konkretnych zadań (np. rozpoznawania obrazów czy tłumaczenia języków), generatywna AI potrafi „rozumieć” i naśladować wzorce danych, na których została wytrenowana, aby następnie generować nowe, unikalne i realistyczne wyniki.

Podstawy jej działania:

1. **Dane treningowe:** Generatywna AI wymaga ogromnych ilości danych treningowych. Na przykład, jeśli ma tworzyć obrazy, potrzebuje tysięcy, a nawet milionów zdjęć. Jeśli ma pisać teksty, potrzebuje obszernego korpusu tekstów. Te dane uczą model „jak wyglądają” lub „jak brzmią” rzeczy, które ma generować.

Podstawy działania generatywnej AI II

2. **Architektura modelu:** Istnieje wiele różnych architektur generatywnych modeli AI, ale dwie najpopularniejsze to:

- ▶ **Generatywne sieci przeciwstawne (GANs – Generative Adversarial Networks):** Składają się z dwóch współpracujących (i rywalizujących) sieci neuronowych:
 - ▶ **Generator:** Jego zadaniem jest tworzenie nowych danych (np. obrazów) na podstawie szumu wejściowego. Stara się tworzyć dane tak realistyczne, jak to tylko możliwe, aby oszukać dyskryminator.
 - ▶ **Dyskryminator:** Jego zadaniem jest odróżnianie prawdziwych danych treningowych od danych wygenerowanych przez generator. Im lepiej generator oszukuje dyskryminator, tym lepsze stają się generowane dane.

Ten proces rywalizacji trwa, aż generator będzie w stanie tworzyć dane, których dyskryminator nie potrafi odróżnić od prawdziwych.

Podstawy działania generatywnej AI III

- **Transformery (Transformers):** Chociaż początkowo zaprojektowane do zadań związanych z językiem naturalnym (NLP), okazały się niezwykle skuteczne w generowaniu tekstu, a także zostały zaadaptowane do innych modalności, w tym obrazów (np. w modelach DALL-E, GPT-3/4). Kluczowym elementem Transformerów jest mechanizm uwagi (*attention mechanism*), który pozwala modelowi ważyć znaczenie różnych części danych wejściowych podczas generowania wyjściowych. Modele językowe bazujące na Transformerach uczą się przewidywać następne słowo w sekwencji, co pozwala im tworzyć spójne i kontekstowe teksty.

Podstawy działania generatywnej AI IV

3. **Proces uczenia:** Podczas treningu model generatywny uczy się ukrytych wzorców i zależności w danych. Nie „rozumie” w ludzkim sensie, ale tworzy statystyczną reprezentację tych danych. Na przykład, model generujący obrazy kotów nauczy się, jakie cechy wizualne (futro, uszy, oczy) są typowe dla kotów i jak są one ułożone względem siebie.
4. **Generowanie nowych treści:** Po wytrenowaniu, model może być użyty do generowania nowych, wcześniej nieistniejących treści. W przypadku GANów, generator otrzymuje losowy „szum” jako wejście i przekształca go w nowe dane. W przypadku modeli językowych, użytkownik podaje „prompt” (monit), a model kontynuuje tekst, przewidując kolejne słowa.

Kluczowe cechy generatywnej AI:

Podstawy działania generatywnej AI V

- ▶ **Kreatywność:** Potrafi tworzyć unikalne i oryginalne treści.
- ▶ **Realizm:** Dąży do tworzenia treści, które są trudne do odróżnienia od prawdziwych.
- ▶ **Wielomodalność:** Zdolność do generowania różnych typów treści (tekst, obrazy, dźwięk).
- ▶ **Kontrolowana generacja:** Wiele nowoczesnych modeli pozwala użytkownikowi wpływać na proces generowania poprzez podawanie opisów tekstowych, stylów lub innych parametrów.

Generatywna AI ma szerokie zastosowania, od wspomagania artystów i projektantów, przez automatyczne generowanie treści marketingowych, po tworzenie realistycznych symulacji i wspieranie innowacji w wielu dziedzinach.

Co to jest przetwarzanie języków naturalnych?

Przetwarzanie języka naturalnego (NLP) (*Natural Language Processing*) to interdyscyplinarna dziedzina na pograniczu **sztucznej inteligencji (AI)**, **informatyki** i **lingwistyki**, która umożliwia komputerom **rozumienie, interpretowanie i generowanie ludzkiego języka naturalnego** (mowy i pisma).

Celem NLP jest wypełnienie luki komunikacyjnej między człowiekiem a maszyną, dając komputerom zdolność do logicznej i kontekstowej analizy języka, tak jak robią to ludzie.

Jak działa NLP?

NLP wykorzystuje zaawansowane algorytmy, uczenie maszynowe (w tym głębokie sieci neuronowe) i reguły lingwistyczne do rozbicia języka na zrozumiałe dla maszyny elementy:

1. **Analiza syntaktyczna (Składniowa):** Skupia się na strukturze zdania, gramatyce i zasadach budowania fraz. Pomaga to komputerowi zrozumieć, jak słowa są ze sobą powiązane.
 - ▶ *Przykład:* Rozpoznanie, które słowo jest podmiotem, a które orzeczeniem.
2. **Analiza semantyczna (Znaczeniowa):** Skupia się na znaczeniu słów, fraz i całego tekstu, często w kontekście.
 - ▶ *Przykład:* Rozróżnienie, czy słowo „zamek” w zdaniu oznacza budowlę, czy mechanizm w drzwiach.
3. **Analiza pragmatyczna:** Zrozumienie, jak kontekst zewnętrzny wpływa na znaczenie i intencję komunikatu (np. ironia, cel wypowiedzi).

Kluczowe zadania i zastosowania NLP I

NLP jest siłą napędową wielu technologii, których używamy na co dzień:

Zadanie NLP	Opis	Przykład zastosowania
Tłumaczenie maszynowe	Automatyczny przekład tekstu lub mowy między językami.	Google Translate.
Analiza sentymentu	Określenie emocjonalnego wydźwięku tekstu (pozytywny, negatywny, neutralny).	Analiza opinii klientów w mediach społecznościowych.
Rozpoznawanie bytów nazwanych (NER)	Identyfikacja i klasyfikacja nazw własnych w tekście (osoby, miejsca).	Automatyczne tworzenie indeksów i baz danych z dokumentów.

Kluczowe zadania i zastosowania NLP II

**Rozpoznawanie
mowy**

Konwersja mowy na
tekst pisany.

Asystenci głosowi (Siri,
Alexa) i transkrypcja
nagrań.

**Generowanie
języka
naturalnego
(NLG)**

Tworzenie spójnego,
czytelного tekstu z
danych strukturalnych.

Automatyczne pisanie
podsumowań sportowych
lub raportów
finansowych.

**Systemy
Q&A/Chatboty**

Odpowiadanie na
pytania użytkowników
lub prowadzenie
konwersacji w języku
naturalnym.

Wirtualni asystenci
obsługi klienta.

Kluczowe zadania i zastosowania NLP III

Podsumowywanie tekstu	Automatyczne tworzenie skróconych wersji długich dokumentów, przy zachowaniu kluczowych informacji.	Streszczenia artykułów naukowych lub prasowych.
------------------------------	---	---

Co to są systemy ekspertowe?

Systemy ekspertowe to programy komputerowe należące do dziedziny sztucznej inteligencji (AI), które mają zdolność **naśladowania procesu decyzyjnego i wiedzy ludzkiego eksperta** w bardzo wąskiej i specjalistycznej dziedzinie.

Zamiast polegać na tradycyjnym, proceduralnym kodzie, systemy te wykorzystują skodyfikowaną wiedzę do rozwiązywania złożonych problemów, proponowania rozwiązań, diagnozowania lub doradzania.

Główne komponenty systemu ekspertowego I

Architektura systemu ekspertowego składa się z kilku kluczowych, współpracujących ze sobą elementów:

1. Baza wiedzy (Knowledge Base)

To serce systemu, które zawiera całą specjalistyczną wiedzę z danej dziedziny. Wiedza ta jest zazwyczaj reprezentowana w postaci:

- ▶ **Faktów:** Podstawowych informacji i danych.
- ▶ **Reguł (If-Then):** Reguł decyzyjnych (np. *JEŻELI* pacjent ma gorączkę i kaszel, *TO* istnieje prawdopodobieństwo grypy).
- ▶ **Heurystyk:** Nieformalnych zasad i „reguł kciuka” używanych przez ekspertów.

Główne komponenty systemu ekspertowego II

2. Silnik wnioskowania (Inference Engine)

To „mózg” systemu, który przetwarza i interpretuje wiedzę z Bazy Wiedzy. Jego zadaniem jest:

- ▶ Analizowanie danych wejściowych od użytkownika.
- ▶ Stosowanie reguł i faktów w celu wyciągnięcia logicznych wniosków.
- ▶ Wykorzystywanie metod wnioskowania (np. **w przód** – z faktów do wniosku; **wstecz** – z hipotezy do faktów ją potwierdzających).

3. Interfejs użytkownika (User Interface)

Umożliwia użytkownikowi wprowadzanie danych (np. objawów, wyników badań, parametrów) i otrzymywanie diagnozy lub porady w przystępnej formie.

Główne komponenty systemu ekspertowego III

4. Moduł Uzasadnienia (Explanation Facility)

Opcjonalny, ale kluczowy komponent, który odróżnia system ekspertowy od zwykłego programu. Wyjaśnia on użytkownikowi, **dlaczego** system podjął daną decyzję lub zasugerował konkretny wniosek, zwiększając zaufanie do jego rekomendacji.

Zastosowania systemów ekspertowych

Systemy ekspertowe są wysoce wyspecjalizowane, co pozwala im osiągać wysoką dokładność w wąskiej dziedzinie. Stosowane są m.in. w:

- ▶ **Medycynie:** Diagnostyka chorób (np. słynny system MYCIN), rekomendowanie planów leczenia.
- ▶ **Finansach:** Ocena zdolności kredytowej, wykrywanie oszustw, doradztwo inwestycyjne.
- ▶ **Inżynierii:** Diagnostyka awarii maszyn i sprzętu, planowanie i projektowanie złożonych systemów.
- ▶ **Zarządzaniu:** Wspomaganie decyzji strategicznych, zarządzanie ryzykiem.

Jakie są podstawy działania LLM? I

Podstawą działania **dużych modeli językowych (LLM)** jest mechanizm **przewidywania następnego słowa** (lub, bardziej precyzyjnie, **tokenu**) w sekwencji, oparty na analizie ogromnych zbiorów danych tekstowych i zaawansowanej architekturze sieci neuronowej zwanej **Transformerem**.

Oto kluczowe filary ich funkcjonowania:

1. Architektura transformera (Transformer architecture)

I

Sercem każdego nowoczesnego LLM jest architektura Transformer, wprowadzona w 2017 roku. Umożliwiła ona przełom w przetwarzaniu języka dzięki:

- **Tokenizacji (Tokenization):** Tekst wejściowy (np. zdanie) jest dzielony na mniejsze jednostki zwane **tokenami** (mogą to być całe słowa, części słów lub pojedyncze znaki). Te tokeny są następnie zamieniane na **wektory liczbowe (osadzanie słów, word embeddings)**, ponieważ komputery operują na liczbach, a nie na słowach. Wektory te kodują znaczenie i kontekst.

1. Architektura transformera (Transformer architecture)

II

- ▶ **Mechanizm uwagi (Attention mechanism/Self-attention):** To kluczowy element Transformera. Pozwala on modelowi na **ważenie istotności** różnych tokenów w tekście wejściowym względem tokenu, który aktualnie przetwarza. Dzięki temu model potrafi zrozumieć dalekie zależności w zdaniu i zachować kontekst, co jest niezbędne do poprawnego rozumienia języka naturalnego (np. w zdaniu „Pies jadł kość, bo był głodny,” mechanizm uwagi wiąże „był głodny” z „Pies”, a nie „kość”).

2. Uczenie na ogromnych danych (Pre-training on vast datasets) I

Modele LLM są wstępnie trenowane (pre-trained) na **gigantycznych korpusach tekstowych**, które obejmują miliardy, a nawet biliony tokenów, pochodzących z różnych źródeł, takich jak:

- ▶ Książki
- ▶ Artykuły i strony internetowe (np. Wikipedia)
- ▶ Kod źródłowy
- ▶ Dialogi i transkrypcje

Proces uczenia jest zazwyczaj **nienadzorowany** (unsupervised learning) i polega głównie na:

2. Uczenie na ogromnych danych (Pre-training on vast datasets) II

- ▶ **Przewidywaniu następnego tokenu:** Model uczy się, jaki token jest najbardziej prawdopodobny, że wystąpi po danym ciągu tokenów.
- ▶ **Uzupełnianiu brakujących tokenów (maskowanie):** W niektórych architekturach model jest trenowany, aby odgadnąć, jaki token został celowo ukryty.

W wyniku tego procesu model ustala ogromną liczbę **parametrów** (wag w sieci neuronowej), które reprezentują statystyczne wzorce i zależności w języku.

3. Generowanie tekstu jako przewidywanie sekwencji (text generation as sequence prediction) I

Kiedy użytkownik wprowadza zapytanie (tzw. **prompt**), LLM generuje odpowiedź w następujących krokach:

1. **Analiza Promptu:** Prompt jest tokenizowany i przetwarzany przez Transformera, aby ustalić kontekst.
2. **Generowanie Pierwszego Tokenu:** Model oblicza **prawdopodobieństwo** wystąpienia każdego możliwego następnego tokenu, bazując na całym kontekście promptu. Wybierany jest token o najwyższym (lub losowo ważonym) prawdopodobieństwie.
3. **Iteracja:** Nowo wygenerowany token jest dodawany do sekwencji, a model używa całego nowego ciągu (prompt + pierwszy token) jako kontekstu do przewidywania **kolejnego tokenu**.

3. Generowanie tekstu jako przewidywanie sekwencji (text generation as sequence prediction) II

4. **Kontynuacja:** Proces powtarza się, token po tokenie, aż do wygenerowania kompletnej i spójnej odpowiedzi (np. osiągnięcia limitu tokenów lub wygenerowania tokenu końca zdania).

LLM nie „rozumie” tekstu w ludzkim sensie. Działa jak niesamowicie zaawansowana **maszyna do przewidywania**, generując najbardziej logiczne i prawdopodobne ciągi słów, które są zgodne ze statystycznymi wzorcami, których nauczył się podczas treningu.

4. Dostrajanie (Fine-tuning) i RLHF I

Po wstępnym treningu (pre-training) podstawowy model jest często **dostrajany**, aby lepiej służyć konkretnym celom i zwiększyć jego użyteczność:

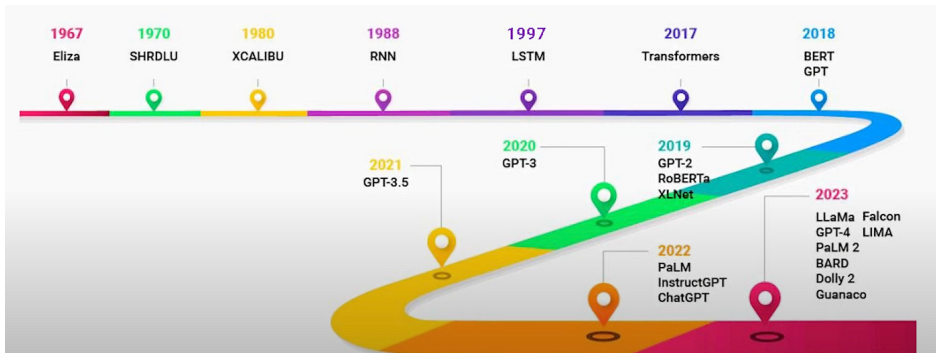
- ▶ **Dostrajanie instruktażowe (Instruction Tuning):** Model jest uczony, aby lepiej odpowiadać na instrukcje i zapytania użytkownika.
- ▶ **Uczenie ze Wzmacnianiem na podstawie Informacji Zwrotnej od Człowieka (Reinforcement Learning from Human Feedback – RLHF):** Ten etap jest kluczowy dla modeli konwersacyjnych (jak ChatGPT). Polega na:
 1. Generowaniu wielu odpowiedzi na jeden prompt.
 2. Ocenianiu tych odpowiedzi przez ludzi-oceniających pod kątem jakości, pomocności i bezpieczeństwa.

4. Dostrajanie (Fine-tuning) i RLHF II

3. Trenowaniu **modelu nagród** (Reward Model) na podstawie tych ludzkich ocen.
4. Używaniu modelu nagród do dalszego optymalizowania i ulepszania głównego LLM za pomocą algorytmów uczenia ze wzmacnianiem.

Dzięki RLHF, model uczy się nie tylko **co** mówić, ale także **jak** się komunikować, aby jego odpowiedzi były bardziej naturalne, trafne i zgodne z oczekiwaniami użytkowników.

Historia LLM



Matematyczne podstawy Sztucznej Inteligencji I

1. Algebra liniowa

Dane w AI są reprezentowane jako struktury matematyczne – wektory, macierze i tensory – które są domeną algebry liniowej.

- ▶ *Wektory i macierze*: Dane wejściowe (np. cechy obrazu, słowa w tekście, wartości pomiarowe) są kodowane jako wektory i macierze. Operacje takie jak iloczyn macierzy są podstawą działania sieci neuronowych (każda warstwa to w zasadzie operacja macierzowa).
- ▶ *Transformacje*: Algebra liniowa opisuje, jak dane są przekształcane w wielowymiarowej przestrzeni, co jest istotą procesów uczenia, np. redukcji wymiarowości (PCA).
- ▶ *Wartości i wektory własne*: Używane do analizy głównych składowych i kompresji danych.

Matematyczne podstawy Sztucznej Inteligencji II

2. Rachunek prawdopodobieństwa i statystyka

AI i ML opierają się na danych, a dane w świecie rzeczywistym są z natury niepewne i losowe. Rachunek prawdopodobieństwa i statystyka pozwalają modelować i zarządzać tą niepewnością.

- ▶ *Rachunek prawdopodobieństwa*: Służy do modelowania losowości i określania prawdopodobieństwa zdarzeń.
- ▶ *Statystyka*: Używana do analizy danych, oceny i walidacji modeli (np. testowanie hipotez), a także do mierzenia wydajności algorytmów (np. błąd średniokwadratowy).

3. Rachunek różniczkowy i całkowy

Jest niezbędny do optymalizacji, czyli kluczowego procesu, w którym algorytmy ML uczą się, minimalizując błąd (funkcję kosztu).

Matematyczne podstawy Sztucznej Inteligencji III

- ▶ *Pochodne i gradienty*: Najważniejszym pojęciem jest gradient – wektor pochodnych cząstkowych, który wskazuje kierunek najszybszego wzrostu funkcji kosztu albo albo minimalizacji błędu
- ▶ *Rachunek wektorowy*: Rozszerza rachunek różniczkowy i całkowy na funkcje wielu zmiennych, co jest typowe dla operacji na wielowymiarowych danych.

4. Optymalizacja

Metody optymalizacji są esencją procesu uczenia maszynowego.



?????!



AI slop I

„**AI slop**” (w wolnym tłumaczeniu „pomyje AI”, „breja AI” lub „cyfrowy ściek”) to **niskiej jakości, masowo generowanych treści** (tekstu, obrazów, filmów, audio) za pomocą narzędzi generatywnej sztucznej inteligencji, które:

1. **Są tworzone z minimalnym wysiłkiem i bez ludzkiej refleksji.** Priorytetem jest ilość i szybkość, a nie zawartość czy jakość.
2. **Są powtarzalne, chaotyczne, bezwartościowe lub pozbawione głębszego sensu.**
3. **Zalewają internet**, zwłaszcza media społecznościowe, blogi i wyniki wyszukiwania, utrudniając znalezienie wartościowych, oryginalnych treści. Często mają na celu nabicie wyświetleń, kliknięć lub przychodów z reklam (tzw. „content farming” lub „SEO slop”).

AI slop II

4. **Mogą zawierać subtelne błędy, nieścisłości, uproszczenia lub zniekształcenia** (tzw. „halucynacje” AI), które są podawane jako prawda.
5. **Bywają łatwe do rozpoznania** (np. nienaturalnie wyglądające ręce na obrazach, dziwne proporcje, chaotyczne tła, błędy gramatyczne i stylistyczne w tekście), ale ich masowość i coraz lepsza jakość utrudniają to.

Krótko mówiąc: AI slop to cyfrowe „śmieci” lub „szum”, który jest ubocznym i niepożądanym efektem masowego stosowania generatywnej sztucznej inteligencji. Zjawisko to prowadzi do obniżenia ogólnej jakości internetu i zaufania do treści online.

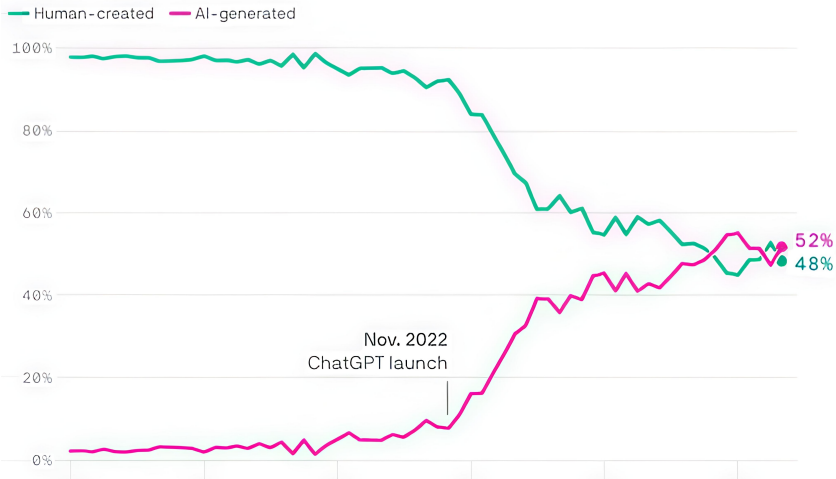
Problemy I

1. Zużycie energii
 - ▶ Szkolenie
 - ▶ Czat
2. „Wolna wola”
3. Świat rzeczywisty
 - ▶ Ludzie działają w fizycznej rzeczywistości
 - ▶ Sztuczna Inteligencja działa w...?
4. Ludzie współdziałają, pracują w grupach, wymieniają się poglądami — SI „rzuca kostką”...
5. Źródła „wiedzy”

Problem 11

Share of articles that were written by humans or generated by AI

Monthly, January 2020 to May 2025; Based on a sample of 65,000 English-language articles published online



...

Jakiś startup zajmujący się podcastami opartymi na sztucznej inteligencji publikuje ponad 3000 odcinków tygodniowo, wszystkie prowadzone przez prezenterów tworzonych przez sztuczną inteligencję.

???

Paint glass full of red wine



Here you go! A glass full of red wine:



Here is a fancy clock showing half past six:



Kolofon

Prezentacja została przygotowana korzystając z języka znaczników [Markdown](#); tekst źródłowy został przekonwertowany do LaTeXa korzystając z programu [pandoc](#).

Program [pdfTeX](#) został użyty do wygenerowania slajdów w formacie PDF. Wykorzystano klasę [beamer](#) i [szablon prezentacji](#) Politechniki Wrocławskiej. Użyto fontu [iwona](#).

Zawartość merytoryczna powstała częściowo w głowie autora, a częściowo została wygenerowana przez Sztuczną Inteligencję (Google Gemini, Perplexity,...).

W trakcie przygotowania wykładu nie ucierpiała żadna SI (czego nie można powiedzieć o składającym wszystko w całość niżej podpisanym).