

**AI is coming,
czyli swobodne rozważania nad sztuczną inteligencją**

Adam Piórkowski*, Mirosław Gajer

Akademia Górniczo-Hutnicza im. St. Staszica w Krakowie, Polska

Streszczenie

Niniejszy artykuł przedstawia swobodne rozważania nad sztuczną inteligencją w kontekście odbioru społecznego i pokładanych w niej nadziei. Prezentowane są różne aspekty, przede wszystkim dotyczące edukacji i nauki. W sposób nawiązujący do tradycji i popkultury wyjaśniono wybrane zagadnienia związane z działaniem sztucznych sieci neuronowych, ze szczególnym wskazaniem tego, co jest pomijane w dyskursie medialnym: braków i niedociągnięć ze strony tej technologii. To, co oferują obecnie istniejące systemy sztucznej inteligencji jest bardzo dalekie od tego, co mogłoby być dopiero ewentualnie postrzegane jako prawdziwa sztuczna inteligencja. W szczególności obecnie nie ma absolutnie żadnych szans, aby można było się spodziewać, że jakikolwiek system sztucznej inteligencji będzie w stanie udowodnić przykładowo hipotezę Riemanna. Podobnie istniejące obecnie systemy komputerowego przekładu są również dalekie od pożądanego w tym zakresie ideału, a samo zastosowane w ich przypadku uczenie maszynowe nie jest bynajmniej w stanie rozwiązać skutecznie wszelkich pojawiających się tutaj problemów.

Słowa kluczowe – sztuczna inteligencja, sieci neuronowe, uczenie głębokie, uczenie maszynowe

* E-mail: pioro@agh.edu.pl

1. Sztuczna inteligencja w domu i zagrodzie

Świat ogarnęła wszechobecna euforia, że oto pojawiło się narzędzie, które jest w stanie rozwiązać problemy ludzkości – sztuczna inteligencja (ang. *Artificial Intelligence*, AI). „Pascal i Machiavelli w jednej osobie” – warto tu zacytować jednego z bohaterów filmu *Hydrozagadka* (1971, reż. Andrzej Kondratiuk). Wszelkie dziedziny nauki doznają raptownie przyrastającego udziału badań, których głównym, wyróżnionym autem jest fakt, że do wytworzenia ich zawartości użyto sztucznej inteligencji. Wszelkie poprzednie podejścia nie są warte uwagi wobec sztucznej inteligencji. A skoro „tworzący naukę” przerzucili się niejako na ten kierunek, to co dopiero się dzieje na ulicach i pod strzechami. A tym bardziej – w szkołach.

Sztuczna inteligencja w postaci np. czatu GPT pokazała uczniom szkół podstawowych, że to, do czego mozolnie dochodzą i czego się uczą latami, można wygenerować łatwo i bez trudu prostym zapytaniem w smartfonie. Zdecydowanie podważyło to sens nauki na poziomie szkół – już nie tylko podstawowych. Ta fascynacja propaguje dalej, a z nią podąża demotywacja do dalszego samodoskonalenia, co może przyszłe pokolenia uzależnić od tych rozwiązań. Niedługo młodzież nawet pole koła będzie wyznaczać przy pomocy sieci neuronowych. Stanowi to kolejne potencjalne niebezpieczeństwa, min. bezrefleksyjne, bezkrytyczne i bezgraniczne zaufanie do sztucznej, a nie własnej, naturalnej oraz redukcję wiedzy i umiejętności [1], co tym bardziej uniemożliwi weryfikację podstawianych przez dostarczane sztuczną inteligencję wyników, a z tego wynika bezpośrednie zagrożenie manipulacją czy wręcz sterowaniem odbiorcami [2].

Fenomen sztucznej inteligencji wiąże się z kilkoma cechami, jakie posiada proponowane rozwiązanie. A zatem jest to technologia nowa, czyli prawdopodobnie lepsza. Zazwyczaj to, co nowe, jest lepsze, np. na hulajnogach elektrycznych lepiej jeździ się po autostradzie niż trabantem z lat 70. XX wieku. Po cóż zatem tracić czas na starocie – jedziemy do przodu, niezależnie od pogody.

Kolejna kwestia to uniwersalność i pracowitość – tego typu rozwiązania, oparte o wydajniejsze jednostki obliczeniowe uczą się szybko i „ogarniają” dużo więcej danych – żmudne szukanie w Excelu łatwiej bowiem zlecić maszynie – ona tysiące kolumn opracuje zdecydowanie szybciej niż człowiek. Aż prosi się przypomnieć, że *lenistwo motorem postępu*.

Spektakularnym atutem Sztucznej Inteligencji są zdolności lingwistyczne. Skonstruowane i nauczone duże modele językowe (ang. *Large Language Model*, LLM) są w stanie poprawić, a nawet przerobić na określony styl czyjąś wypowiedź w danym lub w wielu językach. Zmniejszyło to istotnie zapotrzebowanie na naukę języków obcych i ograniczyło rynek pracy w zakresie tłumaczeń i korekt. Ale warto jednocześnie wspomnieć, że był taki *król, który znał osiem języków, ale w żadnym nie miał nic ciekawego do powiedzenia...*

Jest też społeczne oczekiwanie i promowanie tej technologii. Przykładem PR-owej troski i wybaczenia błędów jest wprowadzenie w tym kontekście terminu halucynacji. Czyli *coś w rodzaju snu na jawie*. Nazywanie tak czynności podawania nieprawdziwych danych na pewno lepiej brzmi niż kłamstwo zarzucane dzieciom, czy rozmijanie się z prawdą, zarzucane politykom. Uzupełnianie brakujących fragmentów danymi, które pasują, z przekonaniem o prawdziwości i bez intencjonalnego zmyślania to raczej konfabulacja, ale to znowu może być odebrane pejoratywnie, a halucynacja? Pewnie z niewyspania, bo się przecież sieć neuronowa przepracowała.... Postulat zmiany nazewnictwa zjawiska przedstawiła też grupa psychiatrów z różnych krajów [3].

Inną kwestią wykorzystania sztucznej inteligencji jest możliwość obciążenia jej odpowiedzialnością. Przecież nie zostanie osądzona. Może więc śmiało wygłosić coś niepopularnego a potrzebnego dla inżynierii społecznej (o nie, jeśli mówi coś niepoprawnego, to... się jej nie używa i publicznie karci np. za rasizm [4] [5], a następnie pracuje nad tą nieprawomyślnością, by ją wykrzewić). Czasem jednak może też powiedzieć coś, co będzie na rękę pewnej grupie ludzi, co owi skwitują, że właśnie taki wynik dała sztuczna inteligencja, więc lepiej być nie może. I tutaj właśnie owo wypracowane wcześniej zaufanie wchodzi w obszar, który potencjalnie może służyć manipulacji, szczególnie łatwej w mediach socjalnych. Obszarem zaś krytycznym jest poleganie na diagnozie stanowiącej często wyrocznię w sprawie zdrowia i życia. Już wkrótce społeczeństwo będzie musiało się zmierzyć z sytuacją, w której *przychodzi baba¹ do lekarza, a tam sztuczna inteligencja*.

Odpowiedzialność wiąże się też z podnoszonym problemem własności intelektualnej – co zrobić, jeśli sieć się nauczyła czyjejś twórczości i na tym fakcie inny jej użytkownik zarabia? To coś podobnego do np. wykorzystania cudzych utworów do

¹ Ze względu na tzw. dorobek kulturowy pominięto opcje „albo chłop, albo inno”.

konstrukcji następnych – takie zachowanie nie było entuzjastycznie przyjęte w społeczeństwie...

Oczywiście, jak wszystko, tak i sztuczna inteligencja jako narzędzie *ma być może nawet swoje minusy*, ale oczywiście *rozchodzi się o to, by te minusy nie przesłoniły plusów*. Niewątpliwym problemem jest fakt, że obecne rozwiązania nie pozwalają określić, co dokładnie ta AI sobie wymyśliła – wprawdzie można wskazać, których danych użyła, czyli *coś dzwoni i już wiadomo, w którym kościele*, to jednak jeszcze nie da się wyekstrahować jej algorytmu, czyli *dziadek wiedział, nie powiedział, a to było tak....*

Warto również pochylić się nad miarą osiągnięć sztucznej inteligencji. Najpopularniejsze obecnie przeprowadzane badania naukowe to swoisty tzw. benchmarking, wypełniający ściśle główny strumień dyskursu służącego do zdobywania cennych kryptowalut ministerialnych, aczkolwiek publikowane są na ogół tylko te wyniki, które owo rozwiązanie stawiają w dobrym świetle [6]. A przecież lista, w których *nasz mózg elektronowy okazał się bezsilny i wypuł wielki znak zapytania* jest również duża. Istnieją przecież problemy, w których wyniki AI w porównaniu do metod klasycznych są *bardzo bardzo słabe, bardzo złe*. O ile człowiek jest w stanie się szybko zaadoptować do nowych danych, to może się zdarzyć, że sztuczna inteligencja, po zmianie źródła danych zamiast diagnozy i wskazania rozwiązania *narysuje kopytko* [7] [5]. Wykazano także, że np. pomijalna przez człowieka zmiana jednego piksela na obrazie może zmienić także diagnozę stawianą automatycznie [8] [9].

Zatem ogólnie rzecz biorąc – raz to działa, raz nie działa, czyli *na dwoje babka wróżyła*. Także kwestia opierania się na wszelkim materiale cyfrowym, jaki np. został wygenerowany w sztuce, może sprowadzić się do statystycznego werdyktu, że jednak *Kopernik była kobietą, a Ziemia ma kształt banana. Zdumiewająca jest ta nowa nauka...*

Koniecznien trzeba coś dodać też o artyzmie. Od lat np. nie powstał żaden utwór, który obiegłby Kulę Ziemią niczym *Final Countdown*. A przecież taki przebój to recepta dla autora na dostatek do końca życia. Tymczasem ostatnio, tym, co znaczna część rodaków podśpiewywała pod nosem, był utwór o oczach, do którego polubienia oczywiście nikt światły, mniemający o własnej inteligencji [10], w żaden sposób się nie przyznawał – a dlaczego? Czyli jednak zdani jesteśmy na genialność ludzką, której takie cechy jak lenistwo przeszkadzają np. w rozwiązywaniu istotnych problemów medycznych? A może i sumie dobrze, bo gdyby AI zaprojektowała sama od

siebie leki na wszelkie choroby, to równie dobrze zaprojektowałaby supercząstkę, która niczym *broń M profesora Wiktora Kuppelweisera* rozwiązałaby problem klimatyczny unicestwiając szybko i bezboleśnie, wręcz „humanitarnie” ludzkość...

2. A tymczasem coś zupełnie na poważnie

Samo pojęcie „sztucznej inteligencji” zostało po raz pierwszy użyte przez grupę pod wodzą Johna McCarthy’ego w roku 1955 jako propozycja projektu naukowego na uczelni w Dartmouth [11]. Począwszy od tego czasu, w rozwoju wymienionej dziedziny zaobserwować można kilka wyraźnych fal wzmożonego entuzjazmu, które ostatecznie kończyły się zwykle powszechnym rozczarowaniem i zawiedzionymi nadziejami, gdyż powstałe wówczas systemy sztucznej inteligencji, oględnie mówiąc, nie do końca spełniały pokładane w nich oczekiwania, przy czym szczyty rozważanych fal dzieli od siebie zwykle okres około kilkunastu lat. Co istotne, każda tego rodzaju fala entuzjazmu, pobudzająca rozwój metod sztucznej inteligencji, wyносиła związane z nimi systemy informatyczne na coraz to wyższy stopień ich skomplikowania, złożoności i funkcjonalności, w związku z czym należy wyraźnie zaznaczyć, że towarzyszący rozwojowi rozważanej dyscypliny entuzjazm nie szedł bynajmniej na marne.

Jednym z kierunków rozwoju systemów sztucznej inteligencji są sztuczne sieci neuronowe, w przypadku których w ich rozwoju także można wyróżnić wspomniane fale wzmożonego entuzjazmu, po których przychodziło jednak nieubłagane zimne otrzeźwienie. Swe pierwsze sukcesy technika sztucznych sieci neuronowych zaczęła święcić już w latach 60. ubiegłego wieku. Otóż w roku 1965 poczta amerykańska z powodzeniem wykorzystwała sztuczną sieć neuronową zbudowaną z liniowych perceptronów do rozpoznawania znaków pisma odręcznego, co umożliwiło praktycznie całkowitą automatyzację procesu sortowania przesyłek pocztowych i zarazem wręcz niesamowicie przyspieszyło realizację związanych z tym czynności.

Niestety dość szybko na rozpalone do czerwoności głowy entuzjastów powszechnego zastosowania tego rodzaju sztucznych sieci neuronowych w praktycznej wszystkich dziedzinach wylany został kubek zimnej wody w postaci opublikowanej w roku 1969 książki autorstwa Marvin’a Minsky’ego i Seymour’a Paperta [12], w której wykazana została ograniczoność potencjalnych zastosowań liniowych perceptronów, które to, jak pokazano za pośrednictwem formalnego dowodu

matematycznego, nie są w stanie w ogóle nauczyć się realizacji nawet tak prostej funkcji logicznej, jaką pełni elektroniczna bramka XOR, nie mówiąc już o znacznie bardziej skomplikowanych tego rodzaju funkcjach logicznych – po prostu zbiory danych, których odróżniania uczy się perceptron, muszą być od siebie liniowo separowane, gdyż w przypadku przeciwnym jego wytrenowanie będzie prostu niemożliwe.

Powstały wskutek tego impas, który nastąpił w dalszym rozwoju technik obliczeniowych związanych z wykorzystaniem sztucznych sieci neuronowych, został przezwyciężony dopiero w roku 1986, gdy David Rumelhart zaprezentował algorytm wstecznej propagacji błędu, który umożliwił już przeprowadzenie efektywnej nauki wielowarstwowych sztucznych sieci neuronowych zbudowanych z neuronów nieliniowych [13].

Tego rodzaju sztuczne sieci neuronowe były w stanie nauczyć się nie tylko realizacji funkcji XOR, ale w zasadzie każdej innej dowolnie wybranej funkcji, co znakomicie poszerzyło obszar ich potencjalnych zastosowań, wzbudzając tym samym kolejną falę entuzjazmu, której kulminacja przypadła na lata 90. ubiegłego wieku, gdy moce obliczeniowe komputerów osobistych stały się na tyle wysokie, że umożliwiły efektywne uczenie wielowarstwowych nieliniowych sztucznych sieci neuronowych na sprzęcie komputerowym rozważanej klasy. W owym czasie całe rzesze pracowników nauki „rzuciły się” wręcz na wyszukiwanie potencjalnych zastosowań dla tego typu sztucznych sieci neuronowych, sprowadzających się najczęściej do klasyfikacji zbiorów danych i rozpoznawania wzorców. Z dzisiejszej perspektywy niektóre z proponowanych wówczas zastosowań wielowarstwowych sztucznych sieci neuronowych mogą wydawać się wręcz kuriozalne, bo jak inaczej można określić przykładowy pomysł związany z wytrenowaniem tego rodzaju sztucznej sieci neuronowej tylko po to, by była później zdolna rozpoznawać dźwięki wydawane przez trąbkę czy bęben od dźwięków skrzypiec.

Prekursor badań nad sztucznymi sieciami neuronowymi w Polsce, profesor Ryszard Tadeusiewicz, miał swego czasu nawet stwierdzić, że w naszym kraju na sztuczne sieci neuronowe zapanowała wręcz swoista moda, w związku z czym brak posiadania we własnym dorobku naukowym jakiegokolwiek publikacji poświęconej rozpatrywanej dziedzinie badań może zostać w środowisku badawczym odebrane jako swego rodzaju poważny nietakt towarzyski [14]. Po przejściu tego rodzaju euforii powiązanej z odkryciem całkowicie nowych potencjalnych obszarów

zastosowań wielowarstwowych nieliniowych sztucznych sieci neuronowych, związana z tym fala entuzjazmu zaczęła wyraźnie słabnąć w początkowych latach XXI stulecia. Uświadomiono sobie wówczas, że realizowane przez tego rodzaju sztuczne sieci neuronowe zadania klasyfikacji danych i rozpoznawania wzorców należą do stosunkowo prostych i, co istotne, nie wydawało się wówczas możliwe ich zastosowanie do realizacji o wiele bardziej ambitnych celów związanych przykładowo z dziedziną przetwarzania języka naturalnego, a zwłaszcza tłumaczenia komputerowego – wówczas powszechnie głoszona konkluzja była taka, że rozpatrywane sieci do realizacji tego rodzaju zadań po prostu w ogóle się nie nadają.

Kolejnym kamieniem milowym w historii rozwoju sztucznych sieci neuronowych było opracowanie koncepcji głębokich sztucznych sieci neuronowych oraz wypracowanie odpowiednich metod uczenia głębokiego przez Geoffreya Hintona w roku 2006 [15]. Sztuczne sieci neuronowe o architekturze głębokiej, które zbudowane są z licznych warstw nieliniowych neuronów, umożliwią znacznie lepsze dopasowanie modelu do danych treningowych, niż miało to miejsce w przypadku stosowanych uprzednio sieci trójwarstwowych uczonych za pomocą algorytmu wstecznej propagacji błędów.

Spektakularnym przykładem zastosowań tego rodzaju sieci jest popularny ChatGPT a także program translacji komputerowej DeepL. Aby teksty zapisane w języku naturalnym mogły być w ogóle przetwarzane z wykorzystaniem techniki sztucznych sieci neuronowych, to muszą wpierw zostać przetworzone do postaci numerycznej, czyli muszą zostać zakodowane w postaci macierzy rzadkich zawierających dane binarne, a to z kolei wymusza zastosowanie do uczenia tego rodzaju sieci superkomputerów o wielkich mocach obliczeniowych, co raczej wyklucza w sposób zdecydowany możliwości „amatorskiego” zajmowania się tego rodzaju zagadnieniami z wykorzystaniem li tylko własnych komputerów osobistych.

Obecnie wydaje się, że ponownie znajdujemy się na szczycie kolejnej, trzeciej już fali powszechnego entuzjazmu, czy wręcz swoistej euforii, wywołanej popularnością zastosowań sztucznych sieci neuronowych, ale z drugiej strony odnosi się nieodparte wrażenie, że małymi krokami nadciągają kolejne rozczarowanie, ponieważ prawdopodobnie znowu, tak jak uprzednio, zdecydowanie zbyt wiele społeczeństwu obiecano, a z całą pewnością można już stwierdzić, że nawet niezwykle popularne obecnie sieci głębokie nie są w stanie rozwiązać wszystkich znanych nam problemów naukowych [16].

Wystarczy chociażby wspomnieć o inicjatywie naukowców z Clay Mathematics Institute w New Hapshire, którzy w roku 2000 sformułowali 7 tzw. problemów milenijnych, czyli najbardziej istotnych zagadnień matematycznych, które być może zostaną rozwiązane w obecnym tysiącleciu [17]. Jednym z tych problemów jest zagadnienie, czy $NP \neq P$, stanowiące prawdopodobnie najbardziej doniosły dylemat współczesnej informatyki teoretycznej. Z kolei jeden ze wspomnianych problemów milenijnych – hipoteza Poincarego został rozwiązany w roku 2003 przez Grigorija Perelmana, który, co niezmiernie interesujące, odmówił przyjęcia przysługującej mu za to nagrody w wysokości miliona dolarów amerykańskich, przy czym motyw jego postępowania nie są do końca jasne, gdyż odmawia on wszelkich komentarzy odnośnie podjętej przez siebie decyzji (jak widać, pieniądze nie były w tym wypadku dla niego najważniejsze, co z drugiej strony wydawać się może z pewnych względów dość intrygujące) [18].

Powyższy przykład pokazuje, że tego rodzaju niezwykle trudne zagadnienia matematyczne znajdują się mimo wszystko w zasięgu możliwości ludzkiego intelektu, jednak nikt przy zdrowych zmysłach nie będzie twierdził, że którykolwiek ze wspomnianych problemów milenijnych mógłby zostać już w chwili obecnej rozwiązany przez jakikolwiek program sztucznej inteligencji, gdyż przekracza to zdecydowanie jej możliwości na bieżącym etapie rozwoju nauki i techniki [19].

W szczególności absolutnie nie należy liczyć na to, że jakiś istniejący obecnie system sztucznej inteligencji udowodni bądź obali sformułowaną w roku 1859 słynną hipotezę Reimanna, głoszącą, że wszystkie nietrywialne zera tzw. funkcji dzeta Reimanna leżą na płaszczyźnie zespolonej na prostej równoległej do osi urojonej i przechodzącej przez punkt $\frac{1}{2}$ położony na osi rzeczywistej.

Podobnie trudno jest oczekiwać, że sztuczna inteligencja jest w stanie udowodnić hipotezę Goldbacha, głoszącą, że każda parzysta liczba naturalna większa od dwóch może zostać przedstawiona w postaci sumy dwóch liczb pierwszych. Tymczasem wciąż żywimy nadzieję, że gdzieś na świecie pojawi się kiedyś jakiś geniusz, który ją wreszcie udowodni, tym bardziej że całkiem niedawno udowodniona została już tzw. słaba hipoteza Goldbacha głosząca, iż każda liczba naturalna nieparzysta większa od 7 może zostać przedstawiona w postaci sumy trzech liczb pierwszych, co zostało formalnie wykazane przez Haralda Helfgotta w roku 2013 [20].

Podobnie można zadać pytanie, czy sztuczna inteligencja będzie w ogóle w stanie udowodnić kiedyś w przyszłości problem Collatza, którego same sformułowanie jest

z drugiej strony wręcz zadziwiająco proste [21]. Otóż na początku podajemy jakąkolwiek dowolną liczbę naturalną N (przykładowo 127) i jeśli liczba ta jest nieparzysta, wówczas mnożymy ją przez trzy i dodajemy jeden ($3N+1$), a jeśli jest parzysta, to dzielimy ją przez dwa ($N/2$). W wyniku wykonania tego rodzaju operacji dostajemy kolejną liczbę naturalną, z którą ponownie postępujemy zgodnie z wymienionymi uprzednio regułami. Obecnie wydaje się, że niezależnie od tego, jaką liczbę na początku sobie arbitralnie wybierzemy, to zawsze ostatecznie wpadamy w cykl (4, 2, 1). Czy tak jest istotnie dla dowolnej liczby naturalnej N , tego jeszcze nikomu nie udało się formalnie wykazać. Węgierski matematyk Paul Erdős miał swego czasu stwierdzić, że obecnie matematyka nie jest jeszcze gotowa na rozwiązywanie tego rodzaju problemów.

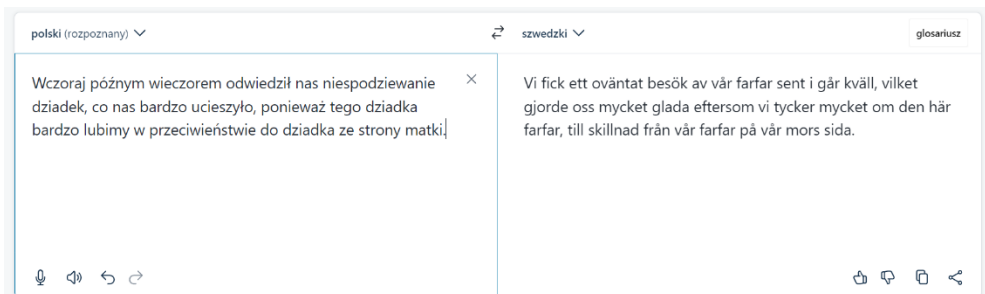
Tym bardziej obecnie na rozwiązywanie wymienionych problemów nie jest jeszcze w żadnym wypadku gotowa sztuczna inteligencja. Pozostaje pytaniem otwartym, ile jeszcze wspomnianych fal entuzjazmu, wynoszących systemy sztucznej inteligencji na coraz to wyższy poziom złożoności i funkcjonalności, musi ostatecznie przeminąć, aby sztuczna inteligencja stała się wreszcie gotowa na rozwiązywanie zagadnień, takich jak choćby wymieniona hipoteza Riemanna, hipoteza Goldbacha czy też problem Collatza.

Obecnie jednym z najbardziej spektakularnych obszarów zastosowań systemów uczenia maszynowego opartych na sztucznych sieciach neuronowych o architekturze głębokiej jest tłumaczenie komputerowe. W ostatnich latach w wymienionej dziedzinie dokonał się istotny przełom, w związku z czym jakość uzyskiwanych za pomocą komputera przekładów jest już relatywnie wysoka i to w zasadzie niezależnie od wyboru tematyki tłumaczonych tekstów i wyboru docelowego języka przekładu – przykładowo program DeepL, uznawany za jeden z najlepszych tego rodzaju komputerowych translatorów, posiada możliwość tłumaczenia w dowolnie wybranym kierunku pomiędzy dowolnie wybraną parą spośród około 30 języków świata.

Jednak pomimo uzyskania niepodważalnych sukcesów w rozważanej dziedzinie komputerowego przekładu, pewne problemy nadal pozostają w ogólnym przypadku nierozwiązywalne [22]. Przyczyna takiego stanu rzeczy nie jest bynajmniej sprawą natury technicznej, ale wynika z samej specyfiki języków naturalnych. Jako przykład tego rodzaju nierozstrzygalnych w ogólnym przypadku zagadnień związanych z automatyzacją przekładu pomiędzy językami naturalnymi można wymienić problem

występujący podczas tłumaczenia polskiego rzeczownika „dziadek” na język szwedzki.

Na pierwszy rzut oka wspomniane zagadnienie może wydawać się wręcz trywialne, jednak w rzeczywistości takim bynajmniej nie jest, ponieważ w języku szwedzkim nie istnieje bezpośredni odpowiednik semantyczny polskiego rzeczownika „dziadek”, tylko do wyboru mamy dwa rzeczowniki, które są jego hiponimami. Jednym z nich jest szwedzki rzeczownik „morfar” oznaczający dziadka ze strony matki, a drugim jest szwedzki rzeczownik „farfar”, który stosowany jest w przypadku dziadka pochodzącego ze strony ojca. Ponieważ polski rzeczownik „dziadek” jest dla nich hiperonimem, to na język szwedzki może zostać przetłumaczony jedynie za pomocą jednego ze wspomnianych szwedzkich hiponimów wspomnianego polskiego rzeczownika. Zatem pytanie w tym wypadku brzmi, którego z rozważanych szwedzkich hiponimów polskiego hiperonimu „dziadek” należy w danym kontekście użyć? Czy poprawne będzie w danym przypadku użycie hiponimu „morfar”, a może należałoby użyć zamiast niego hiponimu „farfar”?



Rysunek 1. Przykład problemów, na jakie natrafia system DeepL podczas próby wyboru poprawnego tłumaczeniem na język szwedzki polskiego rzeczownika „dziadek”

Oczywiście tłumaczący szwedzki tekst człowiek, znający biegle rozważany język, potrafi zapewne skutecznie tego rodzaju dylematy w ogólnym przypadku rozwiązać. Być może kilka stron wcześniej zostały gdzieś podane pewne informacje, na podstawie których można wydedukować, czy w tłumaczonym zdaniu chodzi o dziadka ze strony ojca czy też ze strony matki. Jednak zautomatyzowanie tego rodzaju wnioskowania, opracowanie odpowiedniego zestawu reguł wnioskujących i ich zaszycie w postaci algorytmu zaimplementowanego w jakimś języku programowania nie

wyduje się być w ogólnym przypadku w ogóle możliwe. Z tego powodu komputerowy translator zawsze będzie wybierał w sposób li tylko czysto arbitralny jeden ze wspomnianych hiponimów, w związku z czym prawdopodobieństwo uzyskania poprawnego w danym kontekście przekładu wynosić będzie jedynie 50%.

Przykład tego rodzaju błędnego przekładu pokazano na rys. 1, gdzie program DeepL trzykrotnie użył szwedzkiego rzeczownika „farfar”, niestety w trzecim przypadku postąpił w sposób niewłaściwy, ponieważ należało wówczas użyć rzeczownika „morfar”.

Jak widać, zamieszczone na rys. 1 polskie zdanie zostało ostatecznie przełożone na język szwedzki w sposób błędny i to pomimo zawarcia w jego treści jawnej podpowiedzi w postaci frazy „dziadka ze strony matki”, która została przetłumaczona na język szwedzki całkowicie bezsensownie jako „vår farfar på vår mors sida”, co dosłownie znaczy „nasz ojciec ojca ze strony naszej matki”.

No cóż, tak obecnie prezentuje się aktualny stan badań w dziedzinie sztucznej inteligencji i w szczególności przekładu komputerowego. Co gorsze, nie widać na razie na horyzoncie żadnych realnych perspektyw, aby tego rodzaju problemy mogły zostać kiedyś skutecznie rozwiązane w przyszłości, a sam paradygmat uczenia maszynowego czy też uczenia głębokiego jest w tym wypadku z całą pewnością niewystarczający.

Z coraz powszechniejszym wykorzystaniem modeli językowych wiążą się też potencjalnie swego rodzaju wysoce realne zagrożenia, wynikające niejednokrotnie z całkowicie bezkrytycznego ich stosowania, co sprowadza się w praktyce do tego, że do wiadomości bezrefleksyjnie przyjmujemy pewne fakty oraz podejmujemy także pewne działania tylko dlatego, że tak właśnie zaleca nam sztuczna inteligencja. Na naszych oczach rodzi się w ten sposób swoista „technekracja” (czyli swoiste rządy techniki), sprowadzająca się w praktyce do tego, że systemy techniczne zaczynają w pewien sposób po prostu nami „rządzić”, a przecież nie zawsze podejmowane przez nie decyzje muszą być dla nas dobre bądź w jakikolwiek sposób korzystne – niestety niejednokrotnie ich skutki mogą być dla nas w skrajnych wypadkach wręcz zgubne, przykładowo w obszarze medycyny i ochrony zdrowia [23].

Czyżby zatem stare porzekadło sprowadzające się do stwierdzenia, że „wolę własną głupotę od sztucznej inteligencji” było w rozważanym wypadku nadal aktualne?

Literatura

- [1] Krzysztof Budzyń et al., *Endoscopist deskilling risk after exposure to artificial intelligence in colonoscopy: a multicentre, observational study*, The Lancet Gastroenterology & Hepatology 2025.
[https://doi.org/10.1016/S2468-1253\(25\)00133-5](https://doi.org/10.1016/S2468-1253(25)00133-5).
- [2] Jarosław Bugajski, *Sztuczna inteligencja – zagrożenie czy bezpieczeństwo?*, Zbliżenia Cywilizacyjne vol. 19, no. 3, pp. 31–59, 2023.
<https://doi.org/10.21784/ZC.2023.015>.
- [3] Andrew L. Smith, Felix Greaves, Trishan Panch, *Hallucination or confabulation? Neuroanatomy as metaphor in large language models*, PLOS Digital Health, vol. 2, no. 11, p. e0000388, 2023.
<https://doi.org/10.1371/journal.pdig.0000388>.
- [4] Valentin Hofmann, Pratyusha Ria Kalluri, Dan Jurafsky, Sharese King, *AI generates covertly racist decisions about people based on their dialect*, Nature vol. 633, no. 8028, pp. 147–154, 2024.
<https://doi.org/10.1038/s41586-024-07856-5>.
- [5] Taylor Chung, Jonathan R. Dillman, *Deep learning image reconstruction: a tremendous advance for clinical MRI but be careful...*, Pediatric Radiology vol. 53, no. 10, pp. 2157–2158, 2023.
<https://doi.org/10.1007/s00247-023-05720-8>.
- [6] Karolina Nurzyńska, Michał Strzelecki, Adam Piórkowski, Rafał Obuchowicz, *AI in Medical Imaging and Image Processing*, Journal of Clinical Medicine vol. 14, no. 12, p. 4153, 2025. <https://doi.org/10.3390/jcm14124153>.
- [7] Sayantan Bhadra, Varun A. Kelkar, Frank J. Brooks, Mark A. Anastasio, *On Hallucinations in Tomographic Image Reconstruction*, IEEE Transactions on Medical Imaging vol. 40, no. 11, pp. 3249–3260, 2021.
<https://doi.org/10.1109/TMI.2021.3077857>.
- [8] Min-Jen Tsai, Ping-Yi Lin, Ming-En Lee, *Adversarial Attacks on Medical Image Classification*, Cancers vol. 15, no. 17, 2023.
<https://doi.org/10.3390/cancers15174228>.
- [9] Wiktoria Tajak, Karolina Nurzyńska, Adam Piórkowski, *Vulnerability to One-Pixel Attacks of Neural Network Architectures in Medical Image Classification*, Bio-Algorithms and Med-Systems 2025 [w druku].

- [10] Jacek Wasilewski, Agata Kostrzewa, Jacek Kowalski, Beata Pawłowska, *Gdyby Szekspir pisał disco polo... Wpływ źródła tekstu na jego ocenę*, *Studia Medioznawcze* vol. 3, no. 70, pp. 31–45, 2017.
<https://studiamedioznawcze.eu/index.php/studiamedioznawcze/article/view/369>.
- [11] John McCarthy, Marvin L. Minsky, Nathaniel Rochester, Claude E. Shannon, *A Proposal for the Dartmouth Summer Research Project on Artificial Intelligence: August 31, 1955*, *AI Magazine* vol. 27, no. 4, pp. 12–14, 2006.
<https://onlinelibrary.wiley.com/doi/abs/10.1609/aimag.v27i4.1904>.
- [12] Marvin Minsky, Seymour Papert, *Perceptrons. An introduction to computational geometry*. Cambridge: Massachusetts Institute of Technology. Third printing, 1988.
- [13] David E. Rumelhart, Geoffrey E. Hinton, Ronald J. Williams, *Learning representations by back-propagating errors*, *Nature* vol. 323, no. 6088, pp. 533–536, 1986. <https://doi.org/10.1038/323533a0>.
- [14] Ryszard Tadeusiewicz, *Sieci neuronowe*. Warszawa: Akademicka Oficyna Wydawnicza RM, 1993.
- [15] Geoffrey E. Hinton, Simon Osindero, Yee-Whye Teh, *A Fast Learning Algorithm for Deep Belief Nets*, *Neural Computation* vol. 18, no. 7, pp. 1527–1554, July 2006. <https://doi.org/10.1162/neco.2006.18.7.1527>.
- [16] Stuart J. Russell, Peter Norvig, *Sztuczna inteligencja: nowe spojrzenie*. Gliwice: Helion, 2023.
- [17] James Carlson, Arthur Jaffe, Andrew Wiles, *The millennium prize problems*. Cambridge, Mass., Providence, R.I.: American Mathematical Society, Clay Mathematics Institute, 2006.
- [18] *Rosyjski geniusz rozwiązuje problem za milion dolarów, lecz rezygnuje z medalu Fieldsa*. <https://cordis.europa.eu/article/id/26212-russian-genius-solves-million-dollar-problem-but-declines-fields-medal/pl>.
- [19] Roger Penrose, *Nowy umysł cesarza. O komputerach, umyśle i prawach fizyki*. Warszawa: Wydawnictwa Naukowe PWN, 1995.
- [20] Harald A. Helfgott, *The ternary Goldbach conjecture is true*, arXiv preprint arXiv:1312.7748, 2013. <https://doi.org/10.48550/arXiv.1312.7748>.
- [21] Marek Kubale, *Łagodne wprowadzenie do analizy algorytmów*. Gdańsk: Wydawnictwo Politechniki Gdańskiej, 2021.

- [22] Mirosław Gajer, Zbigniew Handzel, *Rekonstrukcja i rewitalizacja zagrożonych wymarciem języków z wykorzystaniem narzędzie lingwistyki komputerowej*. Kraków: Wydawnictwo Księgarnia Akademicka, 2021.
<https://doi.org/10.12797/9788381385695>.
- [23] Janusz Korwin-Mikke, *Jak działają sztuczne inteligencje*, NCZAS INFO, vol. XXXVI, no. 35-36, p. LXII, 2025.
<https://nczas.info/2025/08/30/korwin-mikke-jak-dzialaja-sztuczne-inteligencje/>.
-

AI is coming – free reflections on Artificial Intelligence

Abstract

This article presents free considerations on Artificial Intelligence in the context of social reception and hopes placed in it. Various aspects are presented, primarily those related to education and science. In a way referring to tradition and pop culture, selected issues related to the operation of Artificial Neural Networks are explained, with particular emphasis on what is omitted in media discourse, i.e. the shortcomings and deficiencies of this technology. Certainly, what is offered by currently existing artificial intelligence systems is still very far from what could possibly be seen as true artificial intelligence. In particular, there is currently absolutely no chance that any artificial intelligence system could be expected to be able to prove the Riemann hypothesis, for example, especially since this has been an open mathematical problem for more than 150 years, the solution of which is probably beyond the capacity of the human intellect. Similarly, the computer translation systems that currently exist are also far from the desired ideal in this respect, and the machine learning applied to them alone is by no means capable of effectively solving all the problems that arise in such systems.

Keywords: *artificial intelligence, neural networks, deep learning, machine learning*