

Bunt inteligentnych odkurzaczy – czy czeka nas koniec?

Bartłomiej W. Papież, Assoc. Prof., PhD, MSc

Big Data Institute, University of Oxford, Wielka Brytania

Skróty: AGI (*artificial general intelligence*) – ogólna sztuczna inteligencja; AI (*artificial intelligence*) – sztuczna inteligencja; ANI (*artificial narrow intelligence*) – wąska sztuczna inteligencja; LLM (*large language model*) – duży model językowy

Wszystkie ilustracje zostały wygenerowane za pomocą Microsoft Copilot.

Tak. Wszyscy umrzemy.

I właściwie na tym mógłbym zakończyć pisanie kolejnego artykułu dla Medycyny Praktycznej – z poczuciem dobrze wypełnionej misji budowania atmosfery napięcia i strachu wizją medycznych robotów, które wyślą lekarzy na bezrobocie, czy też apokaliptycznego buntu inteligentnych i samoświadomych maszyn, jak w filmie *Terminator*. Oczywiście każdy, kto czytał poprzednie odcinki tego cyklu, wie, że jak dotąd nie doszło do zamieszek wywołanych przez bezrobotnych radiologów (Med. Prakt. Ginekol. Położ., 2025; 1: 101–111), chociaż już w 2016 roku zeszłoroczny noblista i „ojciec” sieci neuronowych Geoffrey Hinton ostrzegał przed masowym bezrobociem tych specjalistów. Podczas jednej z konferencji stwierdził: „Powinniśmy przestać szkolić radiologów już teraz. To całkowicie oczywiste, że w ciągu pięciu lat uczenie głębokie będzie osiągać lepsze wyniki niż radiolodzy”¹. Po wczesnych sukcesach dużych modeli językowych (LLM) Hinton ponownie zszokował świat, mówiąc w 2023 roku, że istnieje „10–20-procentowe prawdopodobieństwo, że sztuczna inteligencja (AI) doprowadzi do wyginięcia ludzkości w ciągu najbliższych trzech dekad”². Inni „ojcowie chrzestni” i ich uczniowie również wyrazili swoje obawy dotyczące niektórych z najpoważniejszych zagrożeń związanych z zaawansowaną AI. Wzywali do uznania prac nad ograniczeniem ryzyk związanych z wyginięciem ludzkości spowodowanym przez AI za globalny priorytet, na równi z innymi zagrożeniami społecznymi, takimi

jak pandemii czy wojna nuklearna³. Tymczasem jeszcze inna grupa „ojców” AI (bo jak wiadomo, sukces ma ich wielu), z Yannem LeCunem na czele, argumentowała, że „AI może raczej uratować ludzkość przed wyginięciem”⁴. Zostawiając dywagacje i opinie „ojców i matek” AI, popatrzmy, czy mogłaby się ona zamienić w łowcę pacjentów i na przykład po cichutku podawać Pavulon swoim podopiecznym.

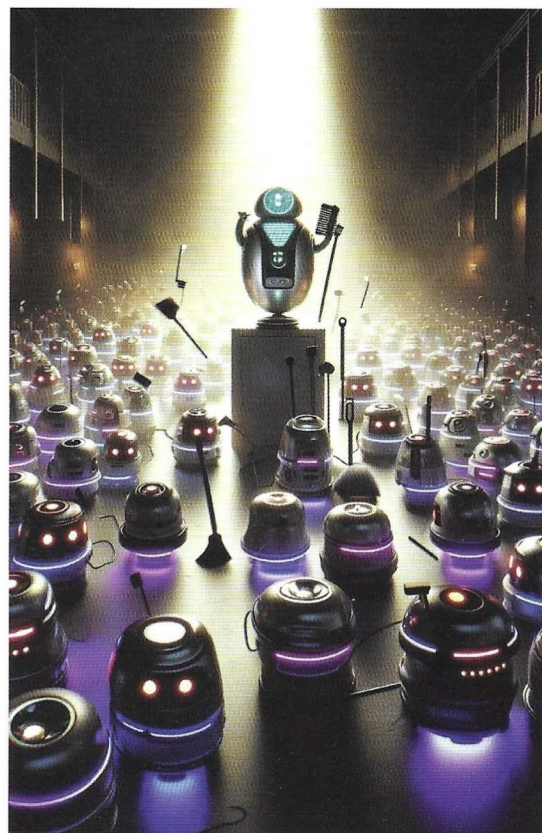
O niedoskonałej AI

„Sztuczna inteligencja” to termin, który jest obecnie odmieniany przez wszystkie przypadki w każdej dziedzinie naszego życia – od medycyny po finanse, od edukacji po motoryzację, włącznie z autonomicznymi samochodami, które obiecał nam Elon Musk w 2015 roku⁵. W poprzednich odcinkach próbowałem opisać AI poprzez pryzmat eksperymentów porównujących konwersacje człowieka i maszyny, takich jak słynny test Turinga⁶, albo analizujących zdolność rozumienia konwersacji, jak chiński pokój⁷. Popularność AI spowodowała też, że (prawie) każde rozwiązanie technologiczne, które dodaje więcej niż dwie liczby z wykorzystaniem komputera (albo nawet liczydła), zaczęło być określane tym mianem. Wydaje się, że procesem tym rządziły proste zasady rynkowe popytu i podaży, bo AI przyciągała i przyciąga uwagę mediów, przemysłu, a tym samym też uwagę inwestorów i nadzieje na potencjalne zyski. Znaczenie pojęcia AI stało się równie rozległe jak

kolejka do chirurgia plastycznego przyjmującego w ramach Narodowego Funduszu Zdrowia (NFZ)⁸, a przy tym tak samo rozmyte jak definicja „zdrowego odżywiania” na popularnych profilach społecznościowych. Niemniej nie ulega wątpliwości, że AI obejmuje zarówno rozwiązania stosunkowo proste, jak i technologie, które mają szansę zrewolucjonizować świat.

Inną perspektywę, która pozwala nam zrozumieć AI, otwiera refleksja nad tym, czym dzisiaj AI różni się od tego, jaka może być w przyszłości (nawet jeśli nasze wyobrażenia ocierają się o *science fiction*). Obecnie dostępną AI często nazywamy wąską (ANI) albo słabą (*weak AI*). Narzędzia ANI zostały zaprojektowane do realizacji ściśle określonego zadania. To taki „specjalista”, który potrafi całkiem nieźle naśladować, a tym samym automatyzować proste działania wykonywane przez człowieka (i znowu wracamy myślami do gry w naśladownictwo, czyli słynnego testu Turinga, oraz ograniczeń tej gry opisanej w eksperymencie chińskiego pokoju). Wszystkie dotychczasowe technologie AI, włącznie z tymi, które przez ostatnie lata rozgrzewają debatę publiczną nad możliwościami AI, czyli ChatGPT i innymi LLM (Claude, DeepSeek), są przedstawicielami słabej AI. Z perspektywy możliwości LLM słaba AI nie wydaje się taka niedoskonała, jakby na to wskazywała sama nazwa. Przykładowy LLM (żeby nie odwoływać się cały czas do ChatGPT) nie jest może geniuszem, bo nie zdał egzaminu specjalistycznego z interny w Polsce (Med. Prakt. Ginekol. Położ., 2025; 2: 85–91), ale całkiem nieźle sobie radzi z odpowiadaniem na ogólne pytania medyczne, uzyskując 86,5% poprawnych odpowiedzi na United States Medical Licensing Examination (USMLE)⁹.

Jeśli kogoś nie przekonują egzaminacyjne możliwości AI, bo swój egzamin lekarski ma już dawno za sobą, może się skupić na innym przykładzie zastosowania LLM w ochronie zdrowia, czyli przygotowywaniu wypisów ze szpitala. Sporządzenie szczegółowego sprawozdania medycznego po zakończeniu leczenia jest czasochłonne, szczególnie dla mniej doświadczonych lekarzy (a to przecież rezydentom uprzyjemnia się czas w szpitalu tym niezbyt lubianym papierkowym zajęciem). Dlatego przykład wykorzystania LLM do automatyzacji



generowania wypisów wydawał się całkowicie naturalnym krokiem w rozwoju medycznej AI¹⁰.

Autorzy postulowali, że wykorzystanie LLM przełoży się na odciążenie przepracowanych się rezydentów oraz poprawę jakości wypisów (ze wszystkimi ważnymi szczegółami, które np. z braku doświadczenia zdarza się pomijać młodszym lekarzom), a przede wszystkim pozwoli się skupić na tym, czym każdy lekarz (a nie tylko rezydent) powinien się zajmować, czyli na opiece nad pacjentem. Nawet jeśli przygotowanie dobrego jakościowo wypisu jawi się laikom jako czarna magia, takiemu zadaniu nadal może sprostać słaba AI.

Podobnie jest z obrazowaniem medycznym, w którym duże modele podstawowe (*foundation models*), czyli obrazowy odpowiednik LLM, także potrafią wykonywać zadania lekarskie, np. wykrywać guzki w płucach¹¹, opisywać radiogramy w reumatoidalnym zapaleniu stawów¹² albo wykrywać drobne osteofity kręgosłupa będące wyni-



ko analizuje dane, wskazując na wzorce zawarte we wcześniej wykorzystanych danych.

Sztuczna inteligencja jutra

Opisana wyżej słaba AI może pomagać w obsłudze pacjentów i usprawniać zarządzanie zasobami ochrony zdrowia. Co więcej, mogłaby usprawnić zarządzanie kolejkami do lekarzy specjalistów, a przede wszystkim przewidywanie, kto się nie zgłosi na umówioną wizytę, na podstawie informacji o pogodzie (tak, chodzi o Anglię), korkach czy pracy wykonywanej przez pacjenta¹⁵. Tym, czego nie może zrobić słaba AI, jest napisanie nowego dzieła literackiego na poziomie Szekspira. Słaba AI nie ma zdolności twórczego myślenia, jest tylko powielaczem wzorców zaobserwowanych w danych, na których się uczyła, i tym samym może być też powielaczem uprzedzeń wynikających z zaskłóści historycznych czy braku reprezentatywnych danych (Med. Prakt. Ginek. Położ., 2025; 4: 99–105).

Brak możliwości powielenia zdolności ludzkiego mózgu przez ANI jest kluczową cechą prowadzącą do definicji ogólnej AI (AGI), czyli technologii, która ma się zbliżyć do skopiowania tego, co potrafi zrobić człowiek. Czasami nazywa się ją także mocną AI (*strong AI*) albo prawdziwą AI (*true AI*). Można sobie wyobrazić, że AGI jest wszechstronnym systemem opakowanym w humanoidalne ciało, który może się uczyć, rozumować, planować, podejmować decyzje i rozwiązywać problemy w sposób podobny do człowieka. W odróżnieniu od słabej AI, którą określiliśmy jako specjalistę do konkretnego zadania, AGI miałaby mieć zdolność rozwiązywania problemów w zróżnicowanych dziedzinach naszego życia, włącznie z robieniem prania i myciem naczyń po obiedzie¹⁵.

Oczywiste wydaje się pytanie, czy to w ogóle jest możliwe, a jeśli tak, to kiedy nastąpi. Raport Global Catastrophic Risk Institute z 2020 roku wskazywał na 72 realizowane na całym świecie projekty naukowe i rozwojowe dotyczące próby zbudowania AGI¹⁶. Wraz z rozwojem LLM pojawiły się też głosy, że w najnowszych modelach, takich jak ChatGPT-4, można widzieć wczesne modele AGI, ponieważ nie tylko odpowiadają na zadane pytania, ale też potrafią rozwiązywać

kiem zmian zwyrodnieniowych¹³. Wszystkie one stanowią przykłady systemów, których zastosowanie jest zazwyczaj ograniczone do konkretnego narządu czy też określonej choroby i jednej metody obrazowania, co także klasyfikuje je jako słabą AI.

I tak dochodzimy do kolejnej cechy słabej AI, którą jest wysoka efektywność, ale tylko w jednym, wybranym obszarze (czyli odpowiadanie na pytania egzaminacyjne, wygenerowanie wypisu, wykrywanie guzków etc.). Najlepszy model (nawet jeśli byłby perfekcyjny) stworzony do wykrywania zmian w reumatoidalnym zapaleniu stawów na radiogramie ręki nie wykryje nawet najbardziej oczywistego złamania kości łódeczkowatej. Podsumowując: zdolności słabej AI są ograniczone do tego, do czego została zaprogramowana; nie potrafi ona wyjść poza te ramy i w kontekście medycznym „nie rozumie” choroby w sposób ludzki, tyl-

zadania obejmujące matematykę, programowanie, przetwarzanie obrazu, medycynę, prawo czy psychologię, bez potrzeby dodatkowego douczania ani specjalistycznego podpowiadania czy instruowania (*prompting*)¹⁷.

Ale czy rozwiązywanie zadań przy użyciu tekstu albo generowanie grafik (takich jak te umieszczane w artykułach poświęconych AI na łamach niniejszego czasopisma) świadczy już o świadomości i kreatywności modelu AI? A może AGI powinna wiedzieć, jak zrobić kawę?¹⁸ Steve Wozniak, współzałożyciel Apple, zaproponował praktyczny test AGI – eksperyment robienia kawy. Sprawdza on, czy robot wyposażony w AGI poradzi sobie z zaparzeniem kawy w dowolnym, wcześniej sobie nieznanym domu. Oczekuje się od niego, że sam wymyśli, jak to zrobić za pomocą tego, co znajdzie na miejscu. Tak więc obsługa ekspresu do kawy albo też wsypanie kilku łyżeczek zmielonych ziaren i zalanie ich wrzątkiem świadczyłoby o byciu inteligentnym przynajmniej jak człowiek...

■ Miałeś złoty róg, ostał ci się ino łańcuch myśli

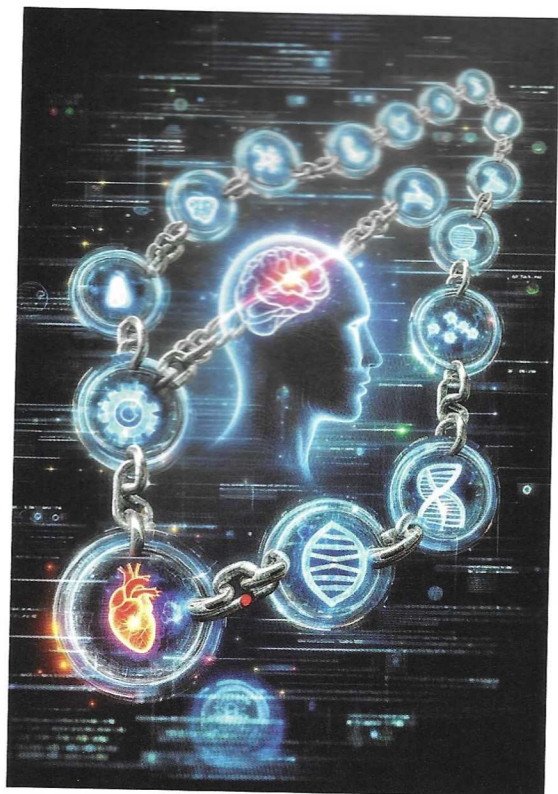
O ile podanie przykładów medycznych słabej AI było bardzo proste, bo takie narzędzia są coraz częściej używane przez lekarzy, o tyle wskazanie, jak będzie wyglądać zastosowanie AGI (poza przygotowaniem kawy w czasie nocnego dyżuru), wchodzi mocno w strefę *science fiction*.

Może jednak jesteśmy świadkami przedwiośnia medycznej AGI? We wrześniu 2024 roku OpenAI udostępniło kolejną wersję ChatGPT-o1 z mechanizmem łańcucha myśli (*chain of thoughts*)¹⁹. Łańcuch myśli to nowe podejście do budowania dużych modeli AI, w którym naśladuje się proces rozumowania typowy dla człowieka. Jeśli pamiętamy choć trochę czasy szkoły podstawowej, bez trudu przypomnimy sobie lekcje, na których byliśmy (mniej lub bardziej skutecznie) uczeni tej metody wnioskowania. Łańcuch myśli w wydaniu kiepskiego nauczyciela podsumował dobitnie Gombrowicz: „Słowacki wielkim poetą był”, „ponieważ wieszczem był”. W wersji optymistycznej edukacji zaangażowany i pełen kreatywności nauczyciel (albo takąż nauczycielka) pokazywał nam, jak krok po kroku rozwiązać zadanie poprzez rozbijanie trudnego problemu na sekwencje



(łańcuch!) mniejszych, ale logicznych kroków, aż do momentu, w którym zauważył na naszej twarzy iskrę żarzącego się intelektu. Strategia łańcucha myśli pozwala nam w życiu codziennym uczyć się nowych, skomplikowanych rzeczy poprzez dzielenie ich na etapy. I dokładnie tak to działa w ChatGPT-o1, gdzie końcowa odpowiedź jest wydedukowana z pośrednich kroków pomagających rozwiązać zadanie, które daliśmy naszemu modelowi AI.

Kilka miesięcy po pojawieniu się ChatGPT-o1 amerykańscy naukowcy (zawsze chciałem użyć tej frazy) przeprowadzili serię eksperymentów z zastosowaniem tego modelu dotyczących generowania diagnozy różnicowej, prezentacji toku rozumowania diagnostycznego, różnicowania wstępnych rozpoznań w procesie triażu, rozumowania probabilistycznego oraz podejmowania decyzji terapeutycznych²⁰. Wszystkie eksperymenty oceniali lekarze eksperci przy użyciu zweryfikowanych narzędzi psychometrycznych. O ile przyzwyczailiśmy się, że AI robi coś tak samo dobrze jak ludzie albo



że ludzie, którzy korzystają z AI, osiągają lepsze rezultaty, o tyle wyniki tej pracy otwierają oczy niedowiarkom. Autorzy zaobserwowali znaczną poprawę w zakresie generowania diagnozy różnicowej oraz jakości rozumowania diagnostycznego i terapeutycznego, zwłaszcza w zadaniach wymagających złożonego, krytycznego myślenia. Eksperymenty pokazały spójną i przewyższającą ludzkie możliwości wydajność w wielu zadaniach z zakresu rozumowania medycznego, ocenianych przez lekarzy, w porównaniu zarówno z wcześniejszymi generacjami modeli AI, jak i z samymi lekarzami.

Wyniki tej pracy nie tylko przekonują nieprzekonanych, ale wskazują także przyszlą perspektywę szybkiego postępu w budowie nowych modeli AI. Tak samo szybko musi się zmieniać nasze myślenie o AI w medycynie. Nie chodzi tu tylko o wybór: AI czy człowiek, ale także o rozważenie, jakie zadania lepiej wykonuje człowiek, a do których lepiej nadaje się AI, oraz w jakich przestrzeniach człowiek powinien współpracować z AI. A wciąż mówimy tylko o co najwyżej wstępie do wstępu do AGI.

Mocna AI w medycynie będzie miała wszystkie zalety najlepszego lekarza, jakiego można sobie wyobrazić: zdolność szybkiego uczenia się, podejmowania poprawnych decyzji dzięki połączeniu zdobytej wiedzy z informacją o stanie pacjenta i jego historii medycznej czy genetyce, empatycznego podejścia do pacjenta i prowadzenia z nim rozmowy, a zarazem będzie pozbawiona ludzkich słabości, takich jak zmęczenie, wypalenie zawodowe czy uprzedzenia. A może AGI będzie mogła nie tylko diagnozować na podstawie danych, ale także rozumieć związki między różnymi chorobami, eksperymentować z nowymi terapiami i przewidywać skutki działań medycznych? Albo będzie potrafiła pracować z pacjentem, nie tylko dostosowując zalecenia do jego sytuacji, ale również „rozumiejąc” jego potrzeby, motywacje lub osobiste wyzwania i oferując indywidualną pomoc w przestrzeganiu zaleceń lekarskich (oczywiście zapisanych skrupulatnie w wypisie lekarskim). Jakkolwiek będzie wyglądała AGI w medycynie, jest pewne, że zmieni relacje między lekarzem a pacjentem²¹. Mocna AI będzie miała potencjał zrewolucjonizowania każdej dziedziny życia, nie tylko medycyny, ponieważ będzie się w stanie uczyć z doświadczenia i dostosowywać do nowych warunków, co pozwoli jej rozwiązywać problemy, których nie jesteśmy w stanie dziś przewidzieć.

Wraz z wielką siłą przychodzi wielka odpowiedzialność

Obecne systemy AI są wąsko wyspecjalizowane i świetnie sobie radzą w swoich dziedzinach; AGI będzie miała szerokie kompetencje jak ludzki umysł i będzie umiała sobie radzić w różnych, nawet nieznanych dziedzinach. Słaba AI nie ma świadomości, rozumienia ani intencji, jest tylko narzędziem w ręku osoby, która z niej korzysta; AGI (o ile jest możliwa) będzie mogła wykazywać cechy przypisywane świadomości, takie jak samodzielne podejmowanie decyzji zgodnie z intuicją i długofalowymi celami. Słaba AI wydaje się bezpieczną technologią, ponieważ działa zgodnie z tym, do czego została zaprogramowana, ale co do AGI istnieje przypuszczenie, że jeśli nie będzie odpowiednio kontrolowana, może stanowić potencjalne zagrożenie, ponieważ będzie umiała po-

dejmować decyzje w sytuacjach, które wykraczają poza nasze zdolności przewidywania.

Tymczasem zważywszy na mocną AI wyłania się jeszcze... super-AI, czyli hipotetyczna AI, która będzie inteligentniejsza od AGI (czyli inteligentniejsza od najinteligentniejszego człowieka) we wszystkich możliwych przestrzeniach wiedzy i życia. Co więcej, taka super-AI będzie mogła, bez naszej kontroli, budować jeszcze lepszą formę samej siebie – super-super-AI, a super-super-AI będzie mogła zbudować super-super-super-AI..., doprowadzając do osobliwości technologicznej (*singularity*), czyli hipotetycznego momentu, w którym rozwój technologiczny staje się niekontrolowalny i nieodwracalny, a tym samym przynosząc nieprzewidywalne konsekwencje dla ludzkości. Rozwój LLM i dużych modeli podstawowych rozbudził na nowo dyskusję o tym, czy i kiedy super-AI może powstać, a jeśli powstanie, czy będzie dla nas zagrożeniem.

Na podstawie wyników osiąganych przez przywołany już wcześniej model ChatGPT-4¹⁷ i szybkości rozwoju dużych modeli AI²² próbuje się argumentować, że ludzkość wkroczyła na ścieżkę wiodącą do zbudowania superinteligencji. Już nasze dywagacje na temat AGI opierały się na *science fiction*; w odniesieniu do super-AI musimy się posłużyć tym większą dozą fantazji, żeby wyobrazić sobie coś, co jest lepsze niż najlepsza wersja nas samych.

Czy to już jest koniec i nie ma już nic?

Najbardziej zaawansowane modele AI, nawet te z wbudowanym mechanizmem łańcucha myśli, mają swoje ograniczenia. Wspomniany wcześniej ChatGPT-01, mimo że bardzo dobrze sobie radzi z niektórymi wyzwaniami, w zadaniach, które przekraczają 16 kroków, zaczyna się gubić jak dziecko we mgle, pozostawione w ciemnym lesie bez latarki²³. To, co dla ludzi wydaje się łatwe, czyli budowanie abstrakcyjnych reguł i dostosowywanie ich do nowych, ale podobnych zadań, wciąż stanowi problem dla LLM. Mocna AI, która ma rozwiązać palące problemy ludzkości (np. skrócić kolejki do specjalistów przyjmujących na NFZ), wywołuje też obawy, że nie będziemy mieli nad nią kontroli. Przecież nietrudno sobie wyobrazić scenariusz, gdzie rozwiązaniem powyż-



szego problemu kolejek będzie pozbycie się lekarzy w myśl „nie ma lekarzy, nie ma kolejek do lekarzy, nie ma niczego”.

Dlatego równolegle z próbami budowy AGI prowadzone są dysputy (i badania), jak zapewnić, że przyszłe systemy AI – zapewne autonomiczne – będą bezpieczne dla ludzi i samych siebie²⁴. Historycznie autorem jednego z ciekawszych rozważań na temat tego, jak budować roboty, które będą postępować etycznie, był pisarz *science fiction* Isaac Asimov. W opowiadaniu z 1942 roku pt. *Runaround* (*Zawierucha*) pisarz przedstawił Trzy Prawa Robotyki. Były to zabezpieczenia i wbudowane zasady etyczne, które brzmiały:

1. Robot nie może skrzywdzić człowieka ani przez zaniechanie działania dopuścić, by człowiek doznał krzywdy.

2. Robot musi być posłuszny rozkazom wydawanym przez ludzi, chyba że stoją one w sprzeczności z Pierwszym Prawem.

3. Robot musi chronić swoje istnienie, o ile ta ochrona nie stoi w sprzeczności z Pierwszym lub Drugim Prawem.

Oczywiście w świecie LLM i dużych modeli podstawowych wystarczy zamienić słowo „robot” na frazę „AI”, żeby otrzymać trzypunktowy dekalog dla autonomicznej AI. Idąc dalej, możemy zamienić AI na medyczną AI i mamy już pewność, że nasza doskonała AI nie będzie podawać nam środków, które mogłyby na nas sprowadzić niebezpieczeństwo. Zgłębiając jednak te trzy zasady odnoszące się do medycznej AI, szybko odkrywamy pewną trudność. Prawa Asimova zakładają, że robot (aka medyczna AI) ma możliwość przewidywania konsekwencji swojego działania (albo zaniechania). Czy jest to możliwe? Wydaje się, że obliczeniowo taki robot musiałby przeanalizować wszystkie możliwości do końca świata, aby mieć pewność, że nie krzywdzi człowieka, podobnie jak komputer IBM Deep Blue wykorzystywał możliwości obliczeniowe superkomputerów do sprawdzenia każdej możliwości wykonania ruchu szachowego²⁵. W szachach mamy jednak ograniczoną liczbę ruchów dla każdej figury, jak natomiast miałyby to wyglądać w przypadku określenia, co jest etyczne lub nieetyczne w zależności od prawa, kultury, religii etc.?

Jedno jest pewne: przyszłość AI to nie tylko kwestia technologii, ale również etyki, filozofii oraz odpowiedzialności społecznej. Dlatego w debatę o słabej AI dzisiaj i mocnej AI jutro powinni się zaangażować nie tylko twórcy tych technologii, ale my wszyscy, ponieważ w przyszłości będziemy z nich korzystać, a może będzie od nich zależeć nasze lepsze życie.

PIŚMIENNICTWO

- Hinton G.: On radiology. 24.11.2016. www.youtube.com/watch?v=ZHMpRXst5vQ (dostęp 15.04.2025)
- AI and human extinction. 1.06.2023. www.bbc.co.uk/programmes/m001md54 (dostęp 15.04.2025)
- Statement on AI Risk. AI experts and public figures express their concern about AI risk. www.safe.ai/work/statement-on-ai-risk (dostęp 15.04.2025)
- LeCun Y.: <https://x.com/ylecun/status/1719475457265938604> (dostęp 15.04.2025)
- Musk E.: Code conference. 2016. www.youtube.com/watch?v=wsixsRI-Sz4 (dostęp 15.04.2025)
- Turing A.M.: Computing machinery and intelligence. *Mind*, 1950; 59: 433–460
- Searle J.: Minds, brains, and programs. *Behav. Brain Sci.*, 1980; 3: 417–457
- Solecka M.: Coraz dłuższe kolejki do leczenia. 9.01.2025. www.mp.pl/kurier/370438,coraz-dluzsze-kolejki-do-leczenia (dostęp 15.04.2025)
- Singhal K., Tu T., Gottweis J. i wsp.: Toward expert-level medical question answering with large language models. *Nat. Med.*, 2025; 31: 943–950

- Patel S.B., Lam K.: ChatGPT: the future of discharge summaries? *Lancet Digit. Health*, 2023; 5: e107–e108
- Ather S., Kadir T., Gleeson F.: Artificial intelligence and radiomics in pulmonary nodule management: current status and future applications. *Clin. Radiol.*, 2020; 75: 13–19
- Zhiyan B., Coates L.C., Papież B.W.: Deep learning models to automate the scoring of hand radiographs for rheumatoid arthritis. [W:] *Annual Conference on Medical Image Understanding and Analysis*. Cham: Springer Nature Switzerland, 2024
- Kundu S.S., Mo Y., Srikiyasemwat N., Papież B.W.: Spinal osteophyte detection via robust patch extraction on minimally annotated X-rays. [W:] *2024 IEEE International Symposium on Biomedical Imaging (ISBI)*. Athens, 2024: 1–5
- NHS: AI expansion to help tackle missed appointments and improve waiting times. 14.03.2024. www.england.nhs.uk/2024/03/nhs-ai-expansion-to-help-tackle-missed-appointments-and-improve-waiting-times/ (dostęp 15.04.2025)
- n0: our first generalist policy. 31.10.2024. www.physicalintelligence.company/blog/pi0 (dostęp 15.04.2025)
- Fitzgerald M., Boddy A., Baum S.D.: 2020 survey of artificial general intelligence projects for ethics, risk, and policy. *Global Catastrophic Risk Institute Technical Report 20–1*. https://gcrinstitute.org/papers/055_agi-2020.pdf (dostęp 15.04.2025)
- Bubeck S., Chadrasekaran V., Eldan R. i wsp.: Sparks of artificial general intelligence: Early experiments with gpt-4. 03.2023. <https://arxiv.org/pdf/2303.12712> (dostęp 15.04.2025)
- Wozniak S.: Could a computer make a cup of coffee? www.youtube.com/watch?v=MowergwQR5Y (dostęp 15.04.2025)
- OpenAI o1 system card. 5.12.2024. <https://cdn.openai.com/o1-system-card-20241205.pdf> (dostęp 15.04.2025)
- Brodeur P.G., Buckley T.A., Kanjee Z. i wsp.: Superhuman performance of a large language model on the reasoning tasks of a physician. 2024. *arXiv preprint arXiv: 2412.10849*
- Kazzazi F.: The automation of doctors and machines: a classification for AI in medicine (ADAM framework). *Future Healthc. J.*, 2021; 8: e257–e262
- Wei J., Tay Y., Bommasani R. i wsp.: Emergent abilities of large language models. 2022. *arXiv preprint arXiv: 2206.07682*
- Valmeekam K., Stechly K., Kambhampati S.: LLMs still can't plan; can LLMs? A preliminary evaluation of OpenAI's o1 on PlanBench. 2024. *arXiv preprint arXiv: 2409.13373*
- Ananthaswamy A.: How close is AI to human-level intelligence? *Nature*, 2024; 636: 22–25
- Deep Blue 1996. www.ibm.com/history/deep-blue (dostęp 15.04.2025)