



Roboty

A co z Panem Bogiem?

A co z Panem Bogiem?

Elon Musk, założyciel Tesla Motors, ostrzega – w ciągu pięciu lat roboty mogą zagrażać człowiekowi.

Elon Musk włączył się do dyskusji na [futurolologicznym forum Egde.org](http://futurolologicznymforum.Egde.org), ostrzegając czytelników, że rozwój sztucznej inteligencji może doprowadzić do sytuacji, w której roboty będą decydować o zabijaniu ludzi – [relacjonuje serwis BusinessInsider](http://relacjonuje.serwis.BusinessInsider).

Założyciel **Tesla Motors** ocenia, że coś naprawdę groźnego może się wydarzyć w perspektywie 5 lat. W odpowiedzi na zarzuty o „dziwaczne” prognozy, Musk tłumaczył, że jego przewidywania nie są bajkami o żelaznym wilku, a on sam nie pisze o czymś, czego nie rozumie. Po kilku minutach jednak post Elona Muska został usunięty.

[Robot, praca, monitoring](#)

Powstanie sztucznej inteligencji to kwestia czasu. Szacuje się, że ludzkość stworzy ją pomiędzy 2040, a 2050 rokiem. A kilka dekad później sztuczna inteligencja wyprzedzi w swoich możliwościach ludzką i stanie się... autonomiczna.

Od filozofów **starożytnej Grecji** po prekursorkę powieści **science-fiction Mary Shelley** – ludzie wszystkich epok mieli podobną obawę: że wynalazki zaczną żyć własnym życiem i staną się potężniejsze niż jego twórcy.

Obawa ta nie jest obca także współczesnym. Wielu naukowców boi się, że **sztuczna inteligencja** z czasem przewyższy człowieka. Obawy te nie są pozbawione swoich racji i w perspektywie długoterminowej ryzyka z tym związane powinny być ostrożnie rozważone. Stworzenie sztucznej inteligencji opierałoby się na zaprojektowaniu [oprogramowania](#) lub maszyny, która w sposób inteligentny mogłaby zrealizować dany cel. Od lat dziedzina ta jest jednak nasycona przesadnym optymizmem – naukowcom bowiem jest niezwykle trudno stworzyć ogólną inteligencję, czy też maszynę, która potrafiłaby niezależnie rozwiązywać problemy i niczym człowiek dostosowywać się do nowych warunków.

Postęp technologiczny jednak przyspiesza. Wąskie aplikacje z pewnymi rodzajami sztucznej inteligencji już dziś są dookoła nas. Za przykład może posłużyć wyszukiwarka **Google**, pokazująca pożądane rezultaty, **iPhone** reagujący na głos właściciela, czy **Amazon** sugerujący, jaką książkę można jeszcze kupić. Już dziś sztuczna inteligencja wygrywa z

człowiekiem w szachy. Z pewnością rozwinie się również w innych obszarach, takich jak [energetyka](#), nauka czy medycyna.

Różni eksperci szacują w ankiecie, że gatunek ludzki wymyśli sztuczną inteligencję pomiędzy 2040 a 2050 rokiem, zaś kilka dekad później sztuczna inteligencja wyprzedzi możliwości człowieka. Eksperci sugerują także, że istnieje 33 proc. prawdopodobieństwo, że taki rozwój wydarzeń będzie dla ludzkości „zły” lub „bardzo zły”.

Ankietowani eksperci nie są jedynymi, którzy wyrażają takie obawy. **Elon Musk**, twórca **PayPal** i **TeslaMotors**, uważa, że sztuczna inteligencja może być nawet bardziej niebezpieczna niż **bomba atomowa**. Z kolei słynny fizyk **Stephen Hawking** nazwał ją „potencjalnie najlepszą lub najgorszą rzeczą w historii, jaka może się przytrafić ludzkości”. Ludzie ci nie są wariatami, a zatem o co cały szum?

Po pierwsze, postęp sztucznej inteligencji nie jest łatwy do bezpośredniego zaobserwowania i zmierzenia. Wraz ze wzrostem siły i autonomii **sztucznej inteligencji** może wzrastać trudność w zrozumieniu i przewidzeniu jej działań. Sztuczna inteligencja może bowiem ewoluować bardzo szybko i na wiele niespodziewanych sposobów.

Głębsza obawa dotyczy sytuacji, w której sztuczna inteligencja nie tylko przewyższyłaby ludzką, ale stale się ulepszała, tworząc za każdym razem coraz doskonalsze wersje. W tym wypadku motywy sztucznej inteligencji mogłyby stać się coraz bardziej niejasne, jej logika coraz bardziej nieprzenikniona, a jej zachowanie nieprzewidywalne. Stałaby się bardzo trudna lub wręcz niemożliwa do kontroli.

Użyjmy pewnego uproszczonego przykładu. Samoucząca się sztuczna inteligencja byłaby zaprogramowana do obliczania dziesiętnych liczby pi. W miarę, jak stawałaby się coraz bardziej inteligentna, zdecydowałaby, że w celu zwiększenia swoich mocy obliczeniowych zaczęnie wykorzystywać całą moc wszystkich komputerów na Ziemi. Z czasem mogłoby się okazać, że to wciąż za mało i potrzebuje większych zasobów. I tak dalej. Rola ludzi w takim scenariuszu nie jest oczywista. A może jednak jest?

Choć naukowcy i badacze zaczynają brać te ryzyka coraz bardziej poważnie, przygotowanie wielu z nich na taką dyskusję może nie będzie łatwe.

Jednym ze sposobów mogłoby być przeprowadzenie eksperymentu w wirtualnym świecie, gdzie można by podać obserwacji ewolucję **sztucznej inteligencji**. W wyniku tego badania łatwiejsze byłoby zidentyfikowanie potencjalnych niebezpieczeństw.

Drugim sposobem byłoby wkalkulowanie symulacji moralności w system sztucznej inteligencji. Naukowcy poczynili pewne postępy w stosowaniu tzw. logiki deontycznej (czyli logiki obowiązków i pozwoleń) wobec robotów, aby ulepszyć ich możliwości rozumowania i przewidywania zachowań. To samo można zastosować do sztucznej inteligencji w ogóle. Mówiąc bardziej prozaicznie, naukowcy z tego obszaru powinni opracować ogólne wytyczne dla eksperymentowania ze sztuczną inteligencją, aby móc przewidzieć różne ryzyka z tym związane. Monitorowanie postępu technologicznego może również wymagać pewnego dnia

globalnej współpracy politycznej, być może na podobnej zasadzie, jak działa dziś

Międzynarodowa Agencja Energii Atomowej.

Na razie nie ma powodów do paniki. Sztuczna inteligencja znajduje się w początkowym stadium. Dzięki niej pewnie uda się stworzyć użyteczne [technologie](#) oraz polepszyć los wielu ludzi. Ale bez odpowiedniego przygotowania, sztuczna inteligencja może zamienić się w potwora **Frankensteina** – dużego, nieprzewidywalnego i zupełnie odmiennego od wizji, którą mieli jego twórcy.

Czytaj też: [Roboty zagrażają rynkowi pracy. Czy połowa zawodów przestanie istnieć?](#)

autor tekstu: [KN](#)

źródło: Bloomberg

Autor: KN