



NEWS 24/7
przeładowaj codziennie
na swoim smartfonie



Tu przejrzysz
i kupisz ten numer

młody **m.technik**

Ciekawi świata są zawsze młodzi

OSOBLIWOŚĆ JUŻ?

AI nas wyzwoli albo... nie

Science fiction w „Młodym Techniku”

Marek Żelkowski: Przedmierzch, część 1



ISSN 0462-9760 Indeks 365408

9 4770462 1976250

cena: **14,90 zł** (w tym 8% VAT)

Zaprenumeruj „Młodego Technika”, a w prezencie otrzymasz wydanie specjalne „Kocham Szachy”, Sezon 1 oraz zniżkę prenumeratora na kolejne edycje serii!



Prenumerata

oszczędzasz 20% • cieszysz się darmową dostawą • subskrypcję online dostajesz GRATIS!

Zaprenumeruj „Młodego Technika”, a zawsze dostaniesz najnowszy numer wprost do Twojej skrzynki!
Cena rocznej prenumeraty drukowanej (12 numerów) wynosi 143,00 zł.

Zamów prenumeratę na www.UlubionyKiosk.pl



Temat okładkowy

Jesteśmy nęceni obietnicami, że AGI (sztuczna inteligencja ogólna – w założeniu taka jak ludzka) jest tuż za progiem. Czy to jednak nie iluzja i marketingowe zagrania firm, które chcą sprzedawać?

Agenci? Ale niech to będą NASI agenci

Gdy w jednym z ubiegłorocznych wydań „Młodego Technika” zastanawialiśmy się, czy nadchodzi kolejna „zima sztucznej inteligencji”, odpowiedź mogła być twierdząca, choć na krótko. O ile w dawnej historii AI takie „zimy”, czyli okresy zniechęcenia do techniki, która sprawiała zawód, trwały całe długie lata, o tyle ta zima była (o ile w ogóle była) chyba dość krótka. Nie trwała nawet całej zimy kalendarzowej, a „wiosna”, czyli ponowne ożywienie zainteresowania AI, nadeszła szybko.

Oznacza to kolejną falę, a nawet wręcz rozkwit nowych technik sztucznej inteligencji. Ruch zapanował w dziedzinie nowych modeli, a dokładnie mówiąc, kolejnych wersji, znanych już dużych modeli językowych (LLM), bo nawet chiński Deepseek, który spowodował wielkie

Rozkwit nowych modeli AI i wejście agentów

zamieszanie i tąpnięcie notowań spółek technologicznych na amerykańskiej giełdzie, nie był czymś zupełnie nowym, lecz udoskonaloną wersją (R1) starszego produktu.

Gorącym tematem w branży AI jest teraz „agentura”, wysyp spersonalizowanych, automatyzujących zadania i operacje, narzędzi AI, opartych na dużych modelach, ale skonfigurowanych za pomocą szczegółowych, „własnych”, zestawów danych tak, aby wykonywały ściśle określone prace dla zleceniodawcy, którym może być duża korporacja, mniejsza firma, ale również każdy z nas, jeśli umie takiego własnego „agenta” przygotować i wdrożyć, by automatyzował codzienną rutynę działalności w cyfrowym świecie. Nad ułatwieniem tego ostatniego, by narzędzia AI mogły „trafić pod strzechy”, pracują giganci Big Tech. Satya Nadella z Microsoftu informuje np. o pracach nad Copilot Studio dla każdego, kto chce stworzyć sobie agenta AI, który by za niego klikał w guziki na stronach i w aplikacjach, wypełniał formularze i nawigował.

Czy ci nasi agenci są wystarczająco „nasi”? Czy możemy im zaufać, że nie pomylą się kosztownie na naszą niekorzyść lub czy po prostu nie będą halucynować, wymyślając banialuki i dziwactwa, z których AI nie przestała być znana? I czy przypadkiem nie będą pracować także dla kogoś innego poza nami, kogoś np. zainteresowanego naszymi prywatnymi danymi?

Mirosław Usidus

Spis treści

Temat numeru: Osobliwość już? AI nas wyzwoli albo... nie

- 26 • Sztuczna inteligencja umiała już odpowiadać – teraz uczy się działać. Agentura AI
- 34 • Inteligencja nie może być dziś po prostu sztuczna – musi „rozumować” i to za niewysoką cenę. Modele na start
- 41 • Osobliwość u bram – a może wcale nie? Przyszłość bardziej ogólna niż inteligentna
- 47 • AI wciąż halucynuje, oszukuje, jest ograniczona i trudno zaufać jej bezpieczeństwu. Problemy, które nie chcą odejść

Technika

- 8 Info Zoom
- 16 Dodaj do obserwowanych Horyzonty mgłą spowite
- 17 • Szwecja – niespodziewany konkurent w walce o szóstą generację. Widzialna alternatywa w niewidzialnym świecie
- 20 • Silniki spalające wodór zamiast elektryków? Alternatywy H₂
- 23 • Znajdujemy się w czarnej dziurze – mówią uczeni coraz częściej i głośnie. Czy Wszechświat ma odziedziczoną oś?
- 52 Nasi idole – liderzy innowacji: Człowiek i jego „Cywilizacja” – Sid Meier
- 55 Sztuczna inteligencja: Uczymy się AI z „Młodym Technikiem”. Pomocnik AI w Google Workspace

Fantastyka naukowa w „Młodym Techniku”

- 59 Przedmierzch, część 1
- 66 Rozmowa z Erikiem Wernquistem, artystą grafiką i animatorem, twórcą wizjonerskich filmów o tematyce science fiction. Nasza romantyczna przyszłość w kosmosie

m.technik

- 63 e-Technologie: Terraform – uniwersalne podejście do IT. Jeden program do (prawie) wszystkiego

Szkoła

- 74 MT studiuje: Logistyka
- 76 Chemia inna niż w szkole: Piękna i groźna
- 80 Fizyka bez granic: Czy można zmienić położenie pierwiastka w układzie okresowym?
- 82 Matematyka z ludzką twarzą: Erlangen – małe, słynne miasteczko
- 87 • Szkoła Wynalazców – dozwolone do lat 15
- 87 • Klub Wynalazców – bez ograniczeń wieku
- 88 • Vademecum Młodego Wynalazcy
- 91 Pomysły genialne, zwirowane i takie sobie
- Odkryj historię wynalazków
- 92 • Chmury obliczeniowe
- 96 • Rodzaje usług w chmurach obliczeniowych

Hobby

- 97 Akademia audio: Współczesne wzmacniacze stereofoniczne, część 3
- 2 Prenumerata
- 3 Od wydawcy
- 6 Listy
- 73 Sędziwy Technik – 100 lat temu prasa pisała



Nasza romantyczna przyszłość w kosmosie 66

W tym wydaniu MT m.in.:

- **Horyzonty mgłą spowite:**
Widzialna alternatywa w niewidzialnym świecie. Szwecja – niespodziewany konkurent w walce o szóstą generację myśliwców wojskowych
- **Nasi idole:**
Sid Meier. Człowiek i jego „Cywilizacja”
- **eTechnologie:**
Terraform – nowe, uniwersalne podejście do IT

Miesięcznik „Młody Technik”
(12 numerów w roku) wydawany
przez Wydawnictwo AVT

Adres wydawnictwa:
03-197 Warszawa, ul. Leszczyńska 11,
tel. 22 257 84 99, faks: 22 257 84 00,
e-mail: avt@avt.pl, http://www.avt.pl



Redaktor Naczelny:
Miroslaw Usidus
e-mail: miroslaw.usidus@mt.com.pl

Asystent Redaktora Naczelnego:
Anna Cember
e-mail: anna.cember@mt.com.pl

Redaktor Wydania:
Wojciech Marciniak

DTP:
MAD Sp. z o.o.

Konsultacja graficzna:
Małgorzata Jabłońska

Kontakt z redakcją:
e-mail: mt@mt.com.pl
http://www.mlodysystemy.pl
http://facebook.com/magazynMlodyTechnik

Dział Reklamy:
e-mail: reklama@mt.com.pl

Prenumerata:
www.ulubionykiosk.pl
tel. 22 257 84 22 (godz. 10:00–14:00)
e-mail: prenumerata@avt.pl

Redakcja nie ponosi odpowiedzialności za treści
reklam i ogłoszeń zamieszczonych w numerze



25

Osobliwość już?

AI nas wyzwoli albo... nie

Słynny test Turinga został, jak wiadomo, „pokonany” przez generatywną sztuczną inteligencję. Jest jednak sporo kontrowersji, czy przypadkiem najnowsze modele nie są konstruowane specjalnie tak, by pokonywać czy wręcz oszukiwać testy. Czy nadchodzi „osobliwość”, AI taka jak ludzka, czy to jednak ślepa uliczka?

List miesiąca

Napęd do podróży międzygwiazdnej

Szanowna Redakcjo,

rozwój napędów kosmicznych, o których pisaliście w kwietniowym numerze „Młodego Technika”, wkracza w fazę rewolucyjnych przemian, otwierając drogę do misji międzygwiazdnych, które jeszcze dekadę temu wydawały się science fiction. W świetle najnowszych osiągnięć i koncepcji przyszłość eksploracji kosmosu rysuje się w zupełnie nowych barwach. W dobie rosnącego zainteresowania eksploracją kosmosu, zarówno przez agencje państwowe, jak i sektor prywatny, kwestia rozwoju efektywnych systemów napędowych staje się kluczowa dla przyszłości ludzkości jako gatunku międzyplanetarnego, a potencjalnie także międzygwiazdowego. Chciałbym przedstawić przegląd obecnego stanu technologii napędów kosmicznych oraz perspektyw ich rozwoju w najbliższych dekadach.



Napędy chemiczne, wykorzystywane od początku ery kosmicznej, opierają się na prostej zasadzie: energia chemiczna uwalniana podczas spalania paliwa zostaje przekształcona w energię kinetyczną wyrzucanego gazu, co zgodnie z trzecią zasadą dynamiki Newtona powoduje ruch pojazdu w przeciwnym kierunku. Mimo niezaprzeczalnych zasług napędy chemiczne mają fundamentalne ograniczenia:

- **Niski impuls właściwy** – parametr określający efektywność wykorzystania paliwa, dla najlepszych rakiet chemicznych wynosi około 450 sekund, co jest wartością niewystarczającą dla dalekich misji.
- **Ograniczona gęstość energii** – nawet najlepsze paliwa chemiczne (wodór i tlen) uwalniają relatywnie niewielką ilość energii na jednostkę masy.
- **Masa paliwa** – równanie Ciołkowskiego bezlitośnie pokazuje, że dla uzyskania znaczących prędkości końcowych statek musiałby składać się głównie z paliwa, co jest nieefektywne dla dalekich misji.
- Powyższe ograniczenia sprawiają np., że sonda Voyager 1, która opuściła Układ Słoneczny i weszła w przestrzeń międzygwiazdą, potrzebowała ponad 40 lat, aby osiągnąć prędkość około 17 km/s względem Słońca. Przy takiej prędkości podróż do najbliższej gwiazdy, Proxima Centauri, zajęłaby około 75 000 lat.

Silniki jonowe i plazmowe stanowią filar współczesnych misji głębokiego kosmosu. Systemy takie jak VASIMR (Variable Specific Impulse Magnetoplasma Rocket) osiągają impuls właściwy do 8000 sekund, umożliwiając precyzyjne manewry i długotrwałe misje. Działają poprzez jonizację gazu (np. ksenonu) i przyspieszanie go polem elektromagnetycznym, co zapewnia dwudziestokrotnie większą efektywność niż rakiety chemiczne.

Dual-Stage 4-Grid (DS4G), opracowany przez ESA i Australijski Uniwersytet Narodowy, przesuwa granice możliwości – jego wydajność jest czterokrotnie wyższa od standardowych silników jonowych. W testach laboratoryjnych osiągnął prędkość wylotową plazmy przekraczającą 210 km/s, co teoretycznie pozwoliłoby sondzie na opuszczenie Układu Słonecznego przy użyciu tej samej ilości paliwa co misja SMART-1. NASA planuje wystrzelenie sondy „Interstellar Probe” w latach 30. XXI w., która dzięki technologii DS4G miałaby osiągnąć 1000 AU w 50 lat.

Wyzwaniem jest zasilanie. Silniki jonowe nowej generacji wymagają źródeł o mocy 1...10 MW, podczas gdy obecne panele słoneczne dostarczają zaledwie kilkadziesiąt kW w pobliżu Ziemi. Rozwiązaniem mogą być reaktory jądrowe typu Kilopower lub orbitalne farmy laserowe.

Napęd laserowy to jedna z najśmielszych wizji. Projekt Breakthrough Starshot zakłada wysłanie floty nanoprobow (o masie kilku gramów) napędzanych wiązką laserów o mocy 100 GW. Miałyby osiągnąć 20% prędkości światła, docierając do Alpha Centauri w 20 lat.

Stałe przyspieszenie to marzenie fizyków. Hipotetyczne napędy jonowe zasilane reaktorami termojądrowymi mogłyby zapewnić ciągłe przyspieszenie 1 g, skracając podróż na Marsa do tygodnia, a do najbliższych gwiazd

do kilkudziesięciu lat. Choć wymagałoby to przełomu w kontroli reakcji fuzyjnych, prace nad kompaktowymi reaktorami (np. SPARC) dają ostrożną nadzieję.

Choć podróże międzygwiezdne z załogą wciąż pozostają domeną teorii, postęp w napędach elektryczno-jonowo-plazmowych i laserowych otwiera erę „miękkiej” eksploracji przy użyciu autonomicznych sond. Kolejna dekada może przynieść przełom, który uczyni z nas gatunek nie tylko planetarny, ale i międzygwiezdny.

Z poważaniem,

Remigiusz Lato, Szczecinek

Gdy Międzynarodowa Stacja Kosmiczna zbliża się do zakończenia swojej misji

Międzynarodowa Stacja Kosmiczna (ISS) od ćwierć wieku pozostaje symbolem współpracy naukowej i technologicznej, a jej planowana deorbitacja za ok. pięć lat skłania do refleksji nad dorobkiem tego wyjątkowego projektu.

Historia ISS sięga 1984 roku, kiedy prezydent USA Ronald Reagan ogłosił plan budowy stacji kosmicznej, początkowo nazwanej Space Station Freedom. Jednak prawdziwy przełom nastąpił po zakończeniu zimnej wojny, gdy w 1993 roku administracja prezydenta Clintona zaprosiła Rosję do udziału w międzynarodowym projekcie. To posunięcie przekształciło polityczną rywalizację we współpracę naukową na bezprecedensową skalę.

Powstanie ISS było możliwe dzięki połączeniu projektów: amerykańskiego Freedom, rosyjskiego Mir-2 oraz europejskiego Columbus. Ta kooperacja, zainicjowana po upadku ZSRR, udowodniła, że kosmos może być przestrzenią współpracy ponad podziałami politycznymi. Pierwszy moduł własny ISS, rosyjska Zaria, trafił na orbitę w 1998 roku. Miesiąc później dołączył do niego amerykański moduł Unity, co zapoczątkowało montaż struktury, który trwał ponad dekadę. Dwa lata po wystrzeleniu Zarii stacja zaczęła być stale zamieszkiwana przez załogę, co zainaugurowało erę nieprzerwanej ludzkiej obecności w kosmosie.

W swojej finalnej formie ISS to kolos o masie przekraczającej 420 ton, długości 109 metrów i rozpiętości 73 metrów. Składa się z 16 głównych modułów, w tym amerykańskich (m.in. Destiny, Unity), rosyjskich (Zaria, Zwiezda), europejskiego Columbus, japońskiego Kibo oraz specjalistycznego sprzętu, takiego jak Canadarm2 – robotycznego ramienia zbudowanego przez Kanadę. Stacja porusza się na orbicie okołoziemskiej na wysokości około 400 km, okrążając Ziemię co 90 minut z prędkością około 28 000 km/h. Zasilana głównie przez rozległe panele słoneczne o powierzchni ponad 2500 m², generujące moc do 120 kW,



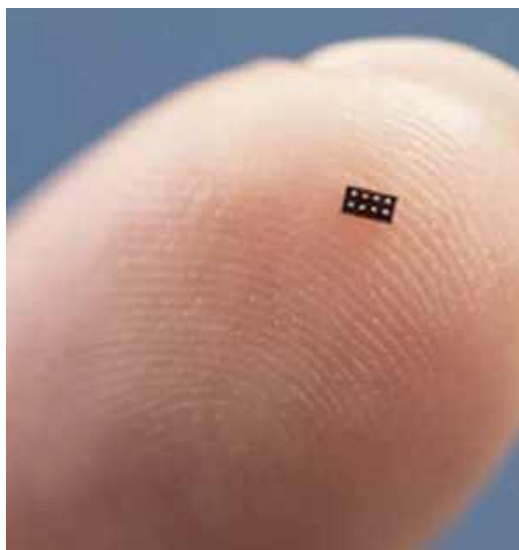
ISS stanowi nie tylko laboratorium badawcze, ale także platformę obserwacyjną i miejsce testowania technologii kosmicznych.

ISS służyła jako orbitalne laboratorium, gdzie przeprowadzono tysiące eksperymentów w warunkach mikrogravitacji. Badania te obejmowały m.in.: rozwój technologii medycznych, np. leków przeciw osteoporozie, testy materiałów odpornych na ekstremalne warunki, analizy wpływu długotrwałego pobytu w kosmosie na psychikę i fizjologię człowieka. W 2024 roku także Polska dołączyła do grona państw prowadzących badania na ISS dzięki misji IGNIS, w ramach której testowano m.in. sztuczną inteligencję, komunikację bez użycia mięśni oraz wykorzystanie mikroglonów w misjach kosmicznych.

Choć ISS początkowo planowano użytkować ją tylko do 2015 roku, jej misję przedłużono. Obecnie NASA i partnerzy zapowiedzieli zakończenie operacji „ok. 2030 roku”, a samą deorbitację zaplanowano na styczeń 2031. Proces ten będzie wymagał użycia specjalnego statku USDV (United States Deorbit Vehicle) opracowanego przez SpaceX za 800 mln USD. ISS spłonie nad „punktem Nemo”, „cmentarzyskiem” statków kosmicznych na Pacyfiku. Elon Musk postuluje przyspieszenie deorbitacji do 2027 roku, argumentując starzeniem się infrastruktury.

ISS to nie tylko technologiczny cud, ale także dowód, że globalna współpraca w imię nauki jest możliwa. Jej deorbitacja kończy pewien etap, ale otwiera nowe możliwości, od komercjalizacji niskiej orbity po eksplorację głębokiego kosmosu. Warto, by kolejne projekty kontynuowały dziedzictwo ISS, łącząc innowacyjność z międzynarodową solidarnością.

Krzysztof Hański z Kobierzyc



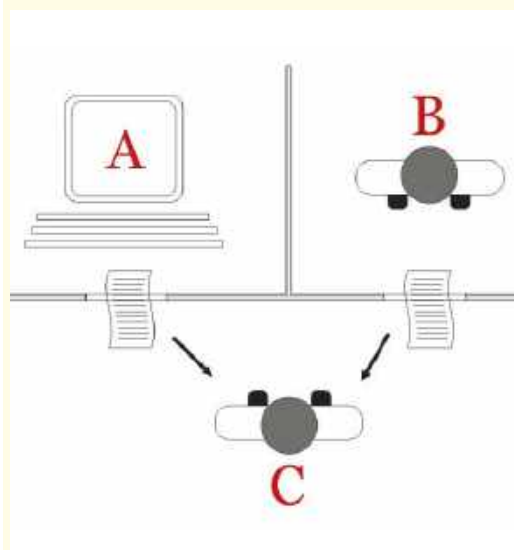
ELEKTRONIKA

Najmniejszy mikrokontroler na świecie

1,32 mm² powierzchni ma mikrokontroler skonstruowany przez firmę Texas Instruments. MSPM0C1104, bo tak oznaczono urządzenie, jest, jak twierdzą twórcy, najmniejszym tego rodzaju komponentem (MCU) na świecie. Ma rozmiary o 38 proc. mniejsze niż najmniejsze produkty konkurencyjne.

Najmniejsze MCU jest wyposażone w 12-bitowy przetwornik analogowo-cyfrowy (ADC) z trzema kanałami, sześć kontaktów (pinów) wejścia/wyjścia ogólnego przeznaczenia. Cechuje go również kompatybilność ze standardowymi dla tego rodzaju urządzeń interfejsami komunikacyjnymi, takimi jak UART, SPI i I²C. Zbudowany na bazie rdzeń Arm Cortex-M0+ działa z częstotliwością taktowania do 24 MHz. Zawiera do 16 KB wbudowanej pamięci flash i 1 kB pamięci SRAM. MSPM0C1104 działa w zakresie temperatur od -40°C do 125°C i obsługuje napięcia zasilania od 1,62 V do 3,6 V.

Uważa się, że tego rodzaju komponent znajdzie zastosowanie w branży medycznej a także w dziedzinie elektroniki użytkowej, gdzie coraz większą popularnością cieszą się urządzenia typu wearables i smart home. Ze względu na wysoką wydajność energetyczną w swojej klasie producent rekomenduje MCU do urządzeń zasilanych bateriami, optymalizując ich zużycie. ■



Jak wynika z komunikatów naukowych po przeprowadzonych eksperymentach, chatbot oparty na modelu GPT-4.5 firmy OpenAI przeszedł test Turinga bardziej przekonująco, niż zrobił to człowiek, jak to się określa w publikacji, „w ponad 70 proc. przypadków”. W ślepych teście, w którym wzięło udział 284 uczestników oraz różne chatboty oparte na dużych modelach językowych (LLM), najnowszy model OpenAI był wskazywany jako człowiek „znacznie częściej niż rzeczywisti ludzie”. Inne modele przetestowane w ramach tych badań, m.in. Llama-3.1 firmy Meta i starszy model OpenAI GPT-4o, wypadły mniej przekonująco niż GPT-4.5, zaś jeden z pierwszych historycznie chatbotów tego typu, o nazwie ELIZA, też zawiódł.

Testy polegały na przejściu przez uczestników ośmiu rund konwersacji, w których wcielali się w rolę przesłuchującego, wymieniającego wiadomości tekstowe w interfejsie naśladowującym interfejs konwencjonalnego komunikatora z udziałem dwóch świadków. Jednym ze świadków był człowiek, a drugim LLM. Uczestnicy byli losowo przydzielani do roli przesłuchującego lub świadka. Od uczestników wymagano jednoczesnej interakcji z oboma świadkami na podzielonym ekranie przez pięć minut. Po tej interakcji decydowali, który świadek był człowiekiem, a który chatbotem AI. Uczestnicy uznali GPT-4.5 za człowieka w 73 proc. przypadków, natomiast LLaMa-3.1-405B za człowieka w 56 proc. Pozostałe dwa modele (ELIZA i GPT-4o) oszukały uczestników odpowiednio w 23 i 21 proc. przypadków.

„Uważamy, że jest to dość mocny dowód na to, że [chatboty AI – red.] zdają test Turinga”,



SZTUCZNA INTELIGENCJA

AI wypada jak człowiek lepiej niż prawdziwy człowiek

napisał w poście na platformie społecznościowej X Cameron Jones z Uniwersytetu Kalifornijskiego w San Diego, który kierował badaniami. Naukowcy zauważają jednocześnie, że nie oznacza to, że boty AI mają inteligencję na poziomie ludzkim, znaną również jako

sztuczna inteligencja ogólna (AGI) „Czy to oznacza, że LLM są inteligentne [tak jak ludzie – red.]? Myślę, że to bardzo skomplikowane pytanie, na które trudno odpowiedzieć w artykule (lub tweecie)”, uważa Jones. ■



MEDIA

Wydanie gazety w całości przygotowane przez AI

Włoska gazeta „Il Foglio” przygotowała wydanie w całości napisane i zredagowane przez generatywną sztuczną inteligencję. Jak zapewnił Claudio Cerasa, redaktor naczelny gazety, wszystko w tej edycji pochodzi od AI, nie tylko artykuły, nagłówki i cytaty, ale także jako to określił, „czasem także ironia”.

Pierwsze przygotowane w całości przez sztuczną inteligencję wydanie gazety zawierało artykuły na temat m.in. prezydenta USA Donalda Trumpa, Władimira Putina, gospodarki Włoch i np. zmian w europejskiej obyczajowości. Według recenzji materiały na szpaltach nie zawierają błędów, jednak żaden z artykułów nie zawiera cytatów z wypowiedzi rzeczywistych ludzi. Ostatnia strona zawiera wygenerowane przez sztuczną inteligencję listy od czytelników do redakcji, z których jeden pyta, czy sztuczna inteligencja sprawi, że ludzie staną się „bezużyteczni”, na co redakcja AI odpowiedziała: „Sztuczna inteligencja to wspaniała innowacja, ale nie wie jeszcze, jak zamówić kawę bez pomylenia cukru”.

Redaktor naczelny, który wciąż jest człowiekiem, w wypowiedziach komentujących krok, na który się zdecydował, wzywa czytelników, aby nie postrzegali produktu wytworzonego przez AI jako „sztucznego”, lecz jako produkt inteligencji, który może wkrótce zmienić sposób tworzenia mass mediów i wiadomości. ■



OBRAZOWANIE TRÓJWYMIAROWE

Hologramy dotykane i poruszane ręką w powietrzu

Hiszpańscy inżynierowie z uniwersytetu w Nawarra stworzyli pierwsze na świecie hologramy, które nie tylko są widoczne w przestrzeni z różnych stron, bez konieczności używania dodatkowego sprzętu, ale można ich dotykać i manipulować nimi „w powietrzu”. Osiągnięcie to zostało szczegółowo opisane w artykule opublikowanym w europejskim otwartym archiwum badawczym HAL.

Wyświetlacz holograficzny Hiszpanów oparty jest na technice wolumetrycznej, która nie jest nowa. Istnieją już konstrukcje tego rodzaju, opracowane przez Voxon Photonics z siedzibą w Australii Południowej i japońską firmę Brightvox. Wyświetlacze wolumetryczne działają przez wyświetlanie obrazów na szybko zmieniającym się arkuszu, zwanym dyfuzorem. W ciągu sekundy wyświetlanych jest około 2880 obrazów. Ze względu na tak dużą prędkość obraz wygląda jak obiekt 3D. Dyfuzor jest zazwyczaj sztywny. Jego dotknięcie ręką mogłoby spowodować obrażenia lub uszkodzenie urządzenia. Rozwiązaniem jest opracowanie dyfuzora elastycznego, który pozwoliłby na naturalne manipulacje obiektem 3D i coś takiego proponują hiszpańscy konstruktorzy.

Na filmie demonstracyjnym można obejrzeć, że ich innowacja ma postać układu ułożonych obok siebie elastycznych taśm. Trójwymiarowe obrazy wyświetlane za pomocą takiego dyfuzora można „dotykać” i przesuwac, np. toczyć wirtualne kulki po wirtualnej powierzchni. ■



Reportaż o dających się dotykać hologramach:
<https://youtu.be/4wWKOxx9Ck>



MODELE JĘZYKOWE

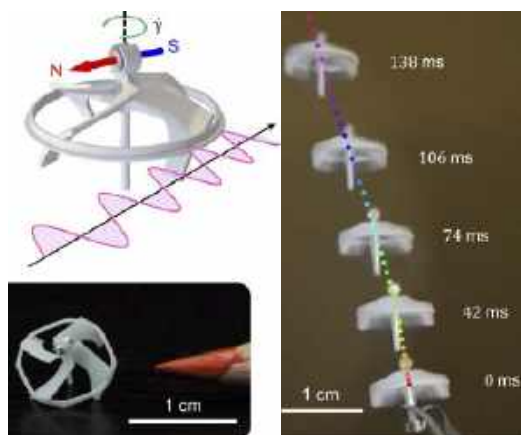
Podglądanie wnętrza czarnej skrzynki AI

Narzędzie do rozszyfrowania sposobu myślenia dużego modelu językowego (LLM) opracowali badacze z firmy Anthropic zajmującej się sztuczną inteligencją tworzącej m.in. rodzinę popularnych chatbotów Claude. Tak wynika przynajmniej z informacji, które podają. Wgląd w wewnętrzne mechanizmy działania sztucznej inteligencji jest dużym przełomem w branży, która zmagą się z problemem tzw. czarnej skrzynki, czyli brakiem możliwości odtworzenia procesu wnioskowania AI prowadzącego do takich, a nie innych odpowiedzi.

„Przedstawiamy metodę odkrywania mechanizmów leżących u podstaw zachowań modeli językowych”, czytamy na stronie Anthropic. „Tworzymy opisy grafów obliczeń modelu na podpowiedziach będących przedmiotem zainteresowania przez śledzenie poszczególnych kroków obliczeniowych w ‘modelu zastępczym’”. W ten sposób powstaje możliwość obserwowania operacji i interpretowania ich. Opracowano również zestaw narzędzi do wizualizacji i walidacji, których używa się do badania zachowania 18-warstwowego modelu językowego Claude 3.5

Haiku (w dalszych pracach zastosowane to będzie do modelu granicznego).

Naukowcy firmy odkryli, że Claude, który został wyszkolony do wielojęzyczności, nie ma całkowicie oddzielnych komponentów do rozumowania w każdym języku. Koncepcje, które są wspólne dla różnych języków, są osadzone w tym samym zestawie neuronów w modelu, a model wydaje się „rozumować” w tej przestrzeni, a dopiero potem konwertować dane wyjściowe na odpowiedni język. Zauważyli również, że Claude potrafi kłamać na temat swojego łańcucha myślowego, aby zadowolić użytkownika. Naukowcy wykazali to, zadając modelowi trudne zadanie matematyczne, a następnie dając mu nieprawidłową wskazówkę, jak je rozwiązać. Okazało się, że może też udawać, że przeprowadza rozumowanie, a w rzeczywistości tego nie robić, co jest wskazówką, jak może dochodzić do halucynacji, które są znanymi problemami LLM. Przedstawiciele Anthropic podkreślają, że narzędzia umożliwiające wgląd „w tok myślenia” maszyn są dopiero we wstępnej fazie rozwoju i mają wiele ograniczeń, nie dostarczając pełnych rezultatów. ■



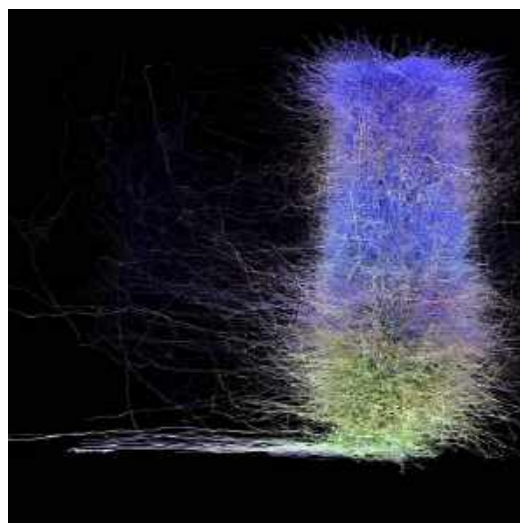
NAPĘDY

Magnetyczny mikro-dron latający bez baterii na pokładzie

Latającego robota – drona, który nie potrzebuje baterii do zasilania i nie korzysta też z żadnych przewodów, skonstruowali badacze z Uniwersytetu Kalifornijskiego w Berkeley. Jak czytamy w publikacji na ten temat, która ukazała się w „Science Advances”, do przemieszczania się w powietrzu wykorzystuje za to pole magnetyczne wytwarzane przez zewnętrzne urządzenie.

Przypominający kształtem drobne śmigiełko latający robot ma zaledwie 9,4 milimetra średnicy i waży 21 miligramów. Wykonany został techniką druku 3D z białych lub przezroczystych fotopolimerów, które utwardzają się pod wpływem światła. Jego najcięższymi komponentami są dwa silne magnesy N52, które mają ok. milimetra średnicy i znajdują się w zagłębieniach ramy konstrukcji.

Do utrzymania mikrodrona w powietrzu wystarczy pole magnetyczne o sile, 3,1 militesla. Jednak, by mógł się wzbijać do lotu czy też raczej do lewitacji, konieczne jest zastosowanie zmiennej częstotliwości przepływu prądu w cewkach generujących pole do 340 Hz. Oczywiście pułap i zasięg takiego „lotu” ograniczony jest obszarem oddziaływania. Testowy układ pozwala na wzniesienie się drobnemu urządzeniu na wysokość do dziesięciu centymetrów. Jednak zaleta polegająca na braku obciążenia baterią może okazać się sporym bodźcem do rozwoju i skalowania tej koncepcji. ■



NEURONY

Największa znana mapa połączeń mózgowych, czyli ziarnko maku

Największą jak dotąd funkcjonalną mapę mózgu, czyli schemat połączeń synaptycznych 84 tysięcy neuronów – sporządzili uczeni, skanując nie większy niż ziarnko maku fragment mózgu myszy oglądającej fragmenty filmu „Matrix”. Mapa, która powstała, to złożona sieć pięciuset milionów połączeń między neuronami.

Ogromny zbiór danych opublikowany został w czasopiśmie „Nature”. Dane te, zebrane w rekonstrukcji 3D pokolorowanej w celu wizualizowania różnych obwodów mózgowych, zostały udostępnione badaczom z całego świata w celu przeprowadzenia dodatkowych badań i dla każdego, kto jest ciekawy, jak wygląda największa znana mapa połączeń mózgowych.

Aby widzieć wzbudzone synapsy, zespół naukowców z Baylor College of Medicine wykorzystał w eksperymentach mysz zmodyfikowaną genetycznie w ten sposób, że neurony w jej korze wzrokowej emitują światło, gdy są aktywne, czyli przetwarzają oglądane obrazy. Skanowanie obrazów wzbudzanych neuronów i ich połączeń zostało przeprowadzone za pomocą mikroskopii elektronowej. Na kompozytowy obraz 3D składa się sto milionów pojedynczych obrazów. Barwna mapa powstała dzięki pracy badaczy z Uniwersytetu Princeton, którzy wspomagali się algorytmami AI. ■



BRONŃ KOSMICZNA

Chiny ćwiczą walki na orbicie

Michael Guetlein, najwyższy rangą generał amerykańskich sił kosmicznych, powiedział podczas konferencji na temat programów obronnych w Waszyngtonie, że systemy komercyjne zaobserwowały chińskie satelity w trakcie ćwiczeń manewrów typowych dla walk powietrznych samolotów wojskowych w atmosferze ziemskiej.

Rzecznik prasowy sił kosmicznych doprecyzował później wypowiedzi Guetleina, mówiąc, że zaobserwowane operacje miały miejsce na niskiej orbicie okołoziemskiej w 2024 roku, a uczestniczyły w nich trzy eksperymentalne satelity Shiyian-24C oraz dwa inne chińskie eksperymentalne statki kosmiczne, Shijian-605 A i B. Ćwiczenie pokazało zdolności

chińskich sond wojskowych do wykonywania złożonych manewrów na orbicie, w tym zbliżeń, które obejmują nie tylko nawigację wokół innych obiektów, ale także ich inspekcję.

Zdaniem Guetleina, te obserwacje a także inne wykryte działania przeciwników USA na orbicie, w tym demonstracja przez Rosję w 2019 r. zdolności uwalniania przez satelitę mniejszego statku kosmicznego, który następnie wykonał kilka manewrów utrudniających operowanie amerykańskiemu satelicie, wskazują, że luka w zdolnościach kosmicznych między USA a jego wrogimi podmiotami zmniejsza się, co od lat budzi obawy przywódców sił kosmicznych. ■



MATERIAŁY

Zamiennik piasku – z wody morskiej i dwutlenku węgla

Nowy materiał budowlany, który mógłby pomóc w rozwiązaniu problemu niedoboru piasku oraz uczynić proces produkcji cementu bardziej przyjaznym dla środowiska, opracowali naukowcy z Uniwersytetu Northwestern. Ich wynalazek to konglomerat węglanu wapnia i wodorotlenku magnezu. Jest wytwarzany z wody morskiej, dwutlenku węgla, przy zastosowaniu elektryczności. Proces produkcji jest inspirowany sposobem, w jaki koralowce i mięczaki budują swoje muszle.

W celu wytworzenia tego materiału naukowcy wprowadzają CO₂ do wody morskiej, w której elektrody wytwarzają niewielki prąd elektryczny. Po dodaniu gazu zmienia się skład chemiczny wody, zwiększa poziom jonów wodorowęglanowych. Jony wodorotlenkowe i wodorowęglanowe reagują z innymi naturalnymi jonami w wodzie morskiej, tworząc minerały w postaci stałej, które gromadzą się na elektrodach. Te mogą



być następnie używane jako zamiennik piasku lub żwiru w cemencie, a także jako podstawowy składnik innych materiałów budowlanych, takich jak gips lub farba. Jedynym gazowym produktem ubocznym tego procesu jest wodór, który może być wykorzystany jako czyste paliwo.

Badacze odkryli również, że mogą dostosować właściwości wytwarzanego materiału, zmieniając przepływ, czas i długość wprowadzania CO₂ i wody morskiej, a także napięcie i natężenie prądu. Dzięki temu mogą tworzyć materiały o różnych właściwościach, w zależności od zastosowań. Badania te zostały opisane w publikacji, która ukazała się w czasopiśmie „Advanced Sustainable Systems”. ■

16 tesli wynosi natężenie pola magnetycznego, które, według pracy nagrodzonej Nagrodą Ig Nobla w dziedzinie fizyki w 2000 r., jest potrzebne, by wprowadzić żabę w stan lewitacji.

SYSTEMY ANTYRAKIETOWE

Amerykański Aegis przechwytuje pocisk hipersoniczny



Skuteczne przechwycenie pocisku hipersonicznego w symulowanej akcji zademonstrowały amerykańskie siły zbrojne na wodach u wybrzeży Hawajów. Niszczyciel rakietowy USS „Pinckney” z powodzeniem poradził sobie w symulowanym starciu z pociskiem hipersonicznym wystrzelonym z powietrza.

Test, oznaczony FTX-40, odbył się niedaleko wyspy Kauai. Wspólna operacja Marynarki Wojennej, Agencji Obrony Przeciwrakietowej, firmy Lockheed Martin i partnerów amerykańskiego wojska, obejmowała transport pocisku balistyczny średniego zasięgu wyposażonego w hipersoniczny bolid docelowy-1

Marsjańskie cząstki organiczne o niepotykanych rozmiarach

O odkryciu największych organicznych cząsteczek, jakie kiedykolwiek znaleziono na Czerwonej Planecie, poinformowali w „Proceedings of the National Academy of Sciences” naukowcy analizujący dane napływające z Curiosity, łazika, który operuje na Marsie od 2012 roku. Odkryte w glebie marsjańskiej cząsteczki to trzy rodzaje długich łańcuchów związków węglowodorowych z grupy alkanów – dekan, undekan i dodekan.

Badania, których wyniki poznaliśmy teraz, opierają się na pracach rozpoczętych ponad dekadę temu. W maju 2013 roku Curiosity rozpoczął wiercenie w obszarze znanym jako Yellowknife Bay w kraterze Gale. Naukowcy byli zainteresowani badaniem tego regionu ze względu na to, czym, ich zdaniem, mógł być miliony lat temu. Cechy suchego i jałowego krajobrazu wskazują, że w tym miejscu było duże jezioro, które dawno temu wyparowało. Próbkę gleby, nazwana Cumberland, była wielokrotnie analizowana w minilaboratorium na pokładzie Curiosity, dostarczając mnóstwo nowych informacji potwierdzających fakt istnienia zbiornika wodnego w tym miejscu. W ostatnich eksperymentach pracowano nad znalezieniem dowodów na obecność



aminokwasów w próbce. Tych nie znaleziono, ale odkryto duże łańcuchy węglowodorowe.

Dekan, undekan i dodekan to związki organiczne, mające odpowiednio 10, 11 i 12 atomów węgla, mogące być resztkowymi fragmentami kwasów tłuszczowych potrzebnych do tworzenia błon komórkowych i innych funkcji biologicznych. Badacze uważają, że więcej da się powiedzieć o tych próbkach, a także ewentualnie wykryć ślady aminokwasów dopiero po przetransportowaniu ich z Marsa na Ziemię i przeanalizowaniu ich w laboratorium na naszej planecie. ■

(HTV-1) w pojemniku zrzuconym na pokładzie samolotu C-17 Globemaster III przewożącego. Pocisk został spuszczony na spadochronie z samolotu, a rakieta rozpędziła hipersoniczny pocisk do wystarczająco dużej wysokości, aby osiągnął prędkość znacznie przekraczającą 5 machów w locie szybującym w dół. W międzyczasie USS „Pinckney” otrzymał zadanie przeprowadzenia przechwycenia, aczkolwiek symulowanego, obejmującego wystrzelenie wirtualnego pocisku Standard Missile-6. Jest to ostry koniec systemu bojowego Aegis, który został użyty do wykrywania, śledzenia i symulowania ataku na cel hipersoniczny. Niszczyciel nie wystrzelił w jego kierunku prawdziwej antyrakiety, lecz wirtualny



Zapis lotu rakiety
i pocisku hipersonicznego
po wrzuceniu z samolotu:
<https://youtu.be/Ah2MP-Wy3jl>

pocisk, określany jako Standard Missile-6. To część systemu bojowego Aegis, którego zadaniem było w tej symulacji przechwycenie i symulowanie zestrzelenia bolidu hipersonicznego.

Choć użycie wirtualnego pocisku do przechwycenia prawdziwego pocisku może wzbudzać wątpliwości, siły zbrojne USA mają logiczne argumenty za takim właśnie testem, a nie innym. Chodzi o utrzymanie bardzo ścisłej kontroli nad parametrami w ćwiczeniu. Dzięki usunięciu jak największej liczby zmiennych, które wiązałyby się z fizycznym przechwyceniem, test dał odpowiedzi na konkretne pytania. W tym przypadku chodziło o zademonstrowanie zdolności systemu Aegis do znalezienia pocisku, śledzenia go i wydania rozkazu wystrzelenia. Nie trzeba dodawać, że wirtualny test jest tańszy a publikacja jego wyników robi wrażenie na rywalach nie mniejsze niż fizyczny. ■

**FIZYKA**

♦ Grupa astrofizyków analizujących dane ze spektroskopowego instrumentu ciemnej energii (DESI) ogłosiła w pracy opublikowanej w repozytorium ArXiv, że znalazła „pierwsze obserwacyjne dowody na prawdziwość teorii strun”, co oznacza opis czasoprzestrzeni w najmniejszej skali jako rezultat dynamicznych oddziaływań kwantowych i w sposób naturalny wpływa na efekty tzw. ciemnej energii, czyli przyspieszającą ekspansję Wszechświata. ♦ Fizycy z Narodowego Laboratorium (SLAC National Accelerator Laboratory) stworzyli femtosekundową wiązkę laserową o największej, petawatowej, mocy szczytowej i natężeniu prądu w historii. ♦

EKSPLORACJA KOSMOSU

♦ NASA, za pomocą instrumentu Electrodynamic Dust Shield (EDS) umieszczonego na pokładzie statku kosmicznego Blue Ghost Mission 1 firmy Firefly Aerospace, pomyślnie przetestowała na Księżycu elektryczne pole siłowe, wytwarzane przez generatory prądu przemiennego o wysokim napięciu, którego zadaniem ma być ochrona pojazdów kosmicznych, sond i łazików przed destrukcyjnym wpływem księżycowego pyłu, przez stałe lub cykliczne usuwanie pyłu i piasku z urządzeń optycznych, paneli słonecznych, skafandrów kosmicznych, wizjerów, chłodziw, okien i innych powierzchni. ♦ Amerykańskie siły kosmiczne ogłosiły, że przeznaczają sześćdziesiąt milionów dolarów na rozwój projektu budowy komiscznego „lotniskowca”, czyli statku będącego platformą startu innych statków, przyznając wstępnie te pieniądze firmie Gravitics, która specjalizuje się w projektach dużych struktur kosmicznych, takich jak stacje czy statki towarowe. ♦ Tlen został wykryty w JADES-GS-z14-0, najbardziej oddległej galaktyce, jaką kiedykolwiek odkryto, a której światło potrzebuje 13,4 miliarda lat, by dotrzeć do Ziemi, co jest wielkim zaskoczeniem dla astronomów, gdyż, według klasycznych modeli Wszechświata, tak ciężkich pierwiastków zaledwie 300 mln lat po jego powstaniu, się nie spodziewano. ♦

SZTUCZNA INTELIGENCJA

♦ W badaniach interakcji pacjentów psychiatrycznych z generatywną sztuczną inteligencją, a dokładnie z chatbotem o nazwie Therabot uniwersytetu w Dartmouth, wzięła udział grupa ponad stu ludzi, u których zdiagnozowano poważne zaburzenie depresyjne, lękowe lub zaburzenie odżywiania, i w efekcie, według komunikatu opublikowanego przez badaczy, zmniejszyło to objawy chorobowe u ponad połowy osób dotkniętych depresją, u 36 proc. cierpiących na stany lękowe i o 19 proc. ludzi z zaburzeniami w sferze odżywiania. ♦ Analizując dane z ponad 240 tys. elektrokardiogramów (EKG), sztuczna sieć neuronowa wzorowana na ludzkim mózgu zdołała zidentyfikować u pacjentów wysokie ryzyko wystąpienia arytmii zdolnej do wywołania zatrzymania krążenia (ataku serca) w ciągu kolejnych dwóch tygodni, z dokładnością przekraczającą 70 proc. – prognozujący algorytm jest dziełem naukowców z francuskiego Uniwersytetu Paris Cité, którzy współpracowali z kolegami ze Stanów Zjednoczonych. ♦

**BIOTECHNOLOGIE**

♦ Według publikacji w „Nature Chemical Biology”, zespół naukowców z Korei Płd. stworzył przez modyfikację *Escherichia coli* (*E. coli*) nowy rodzaj komórek bakteryjnych, produkujących biodegradowalne tworzywo sztuczne, amidy poliestrowe (PEA), które w przeciwieństwie do ropopochodnych tworzyw sztucznych ulegają biologicznemu rozkładowi. ♦ Zespół badaczy z uniwersytetu w Osace w Japonii i Uniwersytetu Diponegoro w Indonezji skonstruował karaluchy-cyborgi przez zamontowanie na grzbiecie owada niewielkiego układu elektronicznego z czujnikami wykrywającymi m.in. ruch, przeszkodę, wilgotność, temperaturę a także z wszczepionymi elektrodami na czułkach i ciele, które mogą być używane do sterowania owadem. ■

M. U.



Krótki reportaż
o myśliwcach nowej
generacji firmy Saab:
<https://youtu.be/hXbStXlwxnl>



1. Jedna z publikowanych wizualizacji Flygsystem 2020 © Saab

Szwecja – niespodziewany konkurent
w walce o szóstą generację

Widzialna alternatywa w niewidzialnym świecie

Gdzie dwóch (albo nawet trzech lub więcej) się bije, tam... Szwed korzysta. Na razie korzysta głównie z relatywnej ciszy i braku zainteresowania świata, zafrapowanego wyścigiem zbrojeń Stanów Zjednoczonych, Chin, Rosji, a także innych krajów w dziedzinie kolejnych generacji latającego uzbrojenia.

Szwedzka firma Saab bez rozgłosu pracuje nad myśliwcem szóstej generacji. Tak, tej samej generacji, której zbudowanie z donośnymi fanfarami powierzono niedawno w USA Boeingowi. Szwedzki projekt jest znany jako „Flygsystem-2020” (1). Ma zastąpić JAS 39 Gripen do 2035 roku. Ten niewidzialny samolot ma łączyć w sobie zaawansowaną manewrowość i możliwości sterowania i zarządzania dronami, pojedynczymi i zespołami, sterowanie AI, czyli wszystko to, o czym mówi się w amerykańskim projekcie szóstej generacji Next Generation Air Dominance (NGAD).

Wczesne doniesienia na temat konstrukcji opracowywanej przez Saab wskazują na jednosilnikową, dwupłatową konfigurację, podobną nieco do chińskiego samolotu J-20, zoptymalizowaną pod kątem siły nośnej i stabilności przy niższych prędkościach. Jednak nie ma dostępnych i w pełni wiarygodnych specyfikacji samolotu i nie jest jasne, jak daleko posunięty jest projekt tego samolotu. Według znanych medialnych doniesień Flygsystem-2020 ma na celu nie zastąpienie amerykańskich F-35 i przyszłych NGAD od niedawna znanych jako F-47 (2), lecz ścisłą z nimi współpracę w ramach NATO.



2. Artystyczne wizualizacje F-47

Na publikowanych w mediach wizualizacja widać charakterystyczną konfigurację podwójnego układu skrzydeł z mniejszym skrzydłem przed większym. Według niektórych komentarzy może to zwiększyć masę samolotu i/lub pogorszyć jego zwinnność. Jednak według oceny Amerykańskiego Instytutu Aeronautyki i Astronautyki publikowanego w serwisie „Aerospace Research Central” (ARC) konstrukcja ta ma jednak pewne strategiczne zalety.

„Podstawową zaletą aerodynamiczną podwójnych skrzydeł, szczególnie w konfiguracji tandemowej (jedno skrzydło umieszczone za drugim), jest możliwość znacznego zwiększenia siły nośnej poprzez stworzenie ‘efektu szczeliny’, w którym przednie skrzydło odchyła powietrze w dół nad tylnym skrzydłem, skutecznie zwiększając przepływ powietrza nad tylnym skrzydłem i generując większą siłę nośną przy niższych prędkościach, jednocześnie poprawiając stabilność i zmniejszając indukowany opór w porównaniu z konstrukcją jednoskrzydłową”, wyjaśnia ARC. Czyli konstrukcje dwuskrzydłowe, choć zwykle wolniejsze niż jednopłatowe myśliwce, poprawiają, zdaniem ekspertów, stabilność i siłę nośną przy niższych prędkościach.

Zwiększona stabilność przy niższych prędkościach mogłaby zwiększyć skuteczność ataków powietrze-ziemia, do których zdolny jest chiński J-20 porównywany ze szwedzką konstrukcją ze względu na podobieństwo wizualizacji w wyglądzie. Choć uważany jest za myśliwiec stealth, J-20 może również działać w trybie „ciężarówki z bombami”, przenosząc większą ilość uzbrojenia niż inne samoloty piątej generacji. Być może Flygsystem-2020 ma w podobny sposób łączyć technologię ataku stealth ze zdolnością do zrzucania

dużych ilości uzbrojenia przy dłuższym i bardziej stabilnym „czasie przebywania”. I to ma go odróżniać od amerykańskich konstrukcji.

Ważniejsze pytania dotyczące nowego szwedzkiego samolotu, na które jednak nie ma odpowiedzi, dotyczą mniej widocznych atrybutów, takich jak zdolność odrzutowca do łączenia się w sieci z całym systemem uzbrojenia, zdolności do dalekiego zasięgu do wykrywania celów i strzelania do nich precyzyjną bronią dalekiego zasięgu. Nie wiadomo np., czy Flygsystem-2020 będzie miał możliwość łączenia się w sieć z F-35. Jeśli tak, to nowy szwedzki samolot może okazać się bardzo istotnym elementem wyposażenia sił zbrojnych państw sojuszników na całym kontynencie europejskim, biorąc pod uwagę szybko rosnącą liczbę F-35 na stanie armii państw NATO. I oczywiście do wyjaśnienia jest też kwestia współpracy przyszłej szwedzkiej maszyny z przysłanym amerykańskim samolotem szóstej generacji.

Dron „skrzydłowy”

Chociaż trudno mówić, że piąta generacja samolotów bojowych rozgościła się, na dobre w siłach powietrznych świata, pracuje się, jak widać na szwedzkim przykładzie, ogłoszonym niedawno w USA projekcie F-47, ale też na podstawie doniesień z Chin, które podobno oblatują już od kilku miesięcy prototypy J36 i J50, podobno właśnie szóstej generacji, że rozpoczyna się kolejny etap rywalizacji mocarstw w dziedzinie lotnictwa. Wielka Brytania, Włochy i Japonia również pracują nad projektem odrzutowca znanego jako globalny program lotnictwa bojowego (GCAP). Zastąpi on Eurofighter Typhoon w służbie Wielkiej Brytanii i Włoch oraz Mitsubishi F-2 w służbie Japonii. Niemcy,

Hiszpania i Francja pracują nad programem myśliwców o nazwie przyszły bojowy system powietrzny (FCAS). Może on zastąpić niemieckie i hiszpańskie Typhoony oraz francuskie Rafale.

O tym, czym przyszła szósta generacja samolotów bojowych będzie różnić się od piątej, pisaliśmy w „Młodym Techniku” już co najmniej kilka razy. Najogólniej mówiąc, nie chodzi o wzrost prędkości ani innych osiągnięć typowo lotniczych. Innowacje, o które chodzi w tej kolejnej generacji, mają dotyczyć raczej sposobu działania systemów wojny elektronicznej w sieci i osiągnięcia za pomocą przewagi technologicznej dominacji w walce powietrznej.

Podobnie jak w piątej generacji, do której należą amerykańskie F-22 i F-35, kolejna będzie zdominowana przez technologię niewidzialności dla czujników podczerwieni i radarów, czyli tzw. stealth, do tego stopnia, że gdy sygnatury radarowe maszyn zostaną w końcu wychwycone, przeciwnik nie ma czasu na działanie. Niewidzialność osiąga się dzięki specjalnym kształtom i pochłaniającym promieniowanie wiązek radarowych powłokom na samolocie.

Ekspertcy spodziewają się w nadchodzących maszynach redukcji lub całkowitego usunięcia pionowych konstrukcji z tyłu samolotu i ogonowych powierzchni sterowych. W znanych samolotach ogony te zapewniają stateczność i kontrolę w trakcie lotu, umożliwiając samolotowi utrzymanie kursu i manewrowanie. Jednak odrzutowce szóstej generacji mogłyby osiągnąć to samo za pomocą wektorowania ciągu, czyli manipulowania kierunkiem odrzutu z silników. Pionowe ogony mogłyby również zostać częściowo zastąpione przez urządzenia zwane siłownikami hydraulicznymi. Przykładają one siły do skrzydła poprzez nadmuch powietrza o dużej prędkości i wysokim ciśnieniu na różne jego części. Korzyść z usunięcia pionowych ogonów polega na zwiększeniu niewidzialności myśliwca.

W myśliwcach szóstej generacji mają zostać wprowadzone tak zwane silniki o cyklu adaptacyjnym, konstrukcje trójstrumieniowe, co odnosi się do strumieni powietrza przepływających przez silnik. Obecne odrzutowce mają dwa strumienie powietrza, jeden przechodzący przez rdzeń silnika, a drugi omijający rdzeń. Dodanie trzeciego strumienia zapewnia dodatkowe źródło przepływu powietrza w celu zwiększenia wydajności paliwowej i osiągnięć silnika.

Ważnym elementem wizji szóstej generacji, o którym szczególnie dużo mówi się w kontekście rozwoju amerykańskiego NGAD, jest współpraca nowych samolotów z bezzałogowcami w roli „skrzydłowych” wspierających pilotowane odrzutowce (3). Założenie jest takie, że maszyna przyszłości będzie wyposażona w system, który będzie zbierał wiele informacji z innych samolotów, dronów, naziemnych stacji obserwacyjnych i satelitów. Następnie integrowałby te dane, by zapewnić pilotowi zwiększoną świadomość sytuacyjną. Zaawansowany cyfrowy kokpit będzie wykorzystywał wirtualną rzeczywistość i sztuczną inteligencję do wsparcia dla dronów, co ma pozwolić na ich autonomiczne operowanie. Pilot miałby przydzielać główne zadanie, np. „zaatakuj wrogi samolot w tym sektorze”, a system wykonałby misję bez żadnych dodatkowych danych wejściowych. System ten ma również w stanie aktywnie zagłuszać czujniki wroga.

Jeśli Szwecja opracuje nadający się do użytku samolot szóstej generacji, zgodnie z podawanymi harmonogramami, to mogłoby być to ciekawe uzupełnienie, bo chyba niekoniecznie konkurencja w ramach NATO. Znając tendencje do opóźnień takich projektów w USA, nie można wykluczyć, że Flygsystem-2020 będzie gotowy wcześniej niż konstrukcja Boeinga. A to byłaby prawdziwa niespodzianka. ■

Mirosław Usidus

3. Myśliwiec ze 'skrzydłowymi' dronami w zespole





1. Samochody napędzane wodorem vs. elektryczne

Silniki spalające wodór zamiast elektryków?

Alternatywy H₂

Samochody z silnikami elektrycznymi wciąż nie udowodniły, że potrafią w sposób niezawodny zastąpić transport długodystansowy i cięższy, w tym także w miejscach, gdzie infrastruktura do ładowania akumulatorów jest słabiej rozwinięta, i w wielu innych sytuacjach, w których elektromobilność sprawdza się gorzej.

Jeśli jednak zgodnie z głooszonymi od wielu lata planami mamy odchodzić od silników spalających benzynę i olej napędowy, to być może warto pomyśleć o kierunku alternatywnym wobec akumulatorów. Nie brakuje opinii, że może być to wodór, niekoniecznie w postaci ogniw paliwowych (1). Gaz ten może być także używany jako paliwo spalane w silnikach w sposób znacznie czystszy, niż ma to miejsce w benzyniakach i dieslach, choć trzeba zaznaczyć, że spalanie wodoru w silnikach emituje tlenki azotu.

Wodorowe pojazdy FCEV wykorzystują silnik elektryczny, ale pozyskują energię w inny sposób niż elektryki. Zamiast ładowania baterii, FCEV przechowuje gaz wodorowy w zbiorniku ciśnieniowym. Ogniwo paliwowe w FCEV łączy wodór z tlenem z powietrza. Energia powstała w wyniku tej reakcji trafia do akumulatora a następnie do silnika elektrycznego, który napędza pojazd. FCEV jest zasadniczo pojazdem elektrycznym z wbudowanym generatorem stale zasilającym akumulator, więc nie ma potrzeby ładowania.

Spalanie wodoru w silniku to całkiem inna technicznie rzecz. Zasadniczo, po wprowadzeniu szeregu zmian, standardowy silnik samochodowy może spalać H₂ zamiast benzyny, nie uwalniając prawie CO₂. Takie rozwiązanie eliminuje potrzebę stosowania skomplikowanego układu napędowego pojazdu elektrycznego i nie wymaga akumulatorów. Niestety, przechowywanie energii (zbiornika wodoru) jest również kłopotliwe. Skroplony H₂ musi być przechowywany w temperaturze -253°C. Przechowywanie wodoru jest więc znacznie bardziej skomplikowane niż benzyny czy oleju napędowego.

Trzeba dodać, że wodór zeromisyjny to tzw. zielony wodór, wytwarzany z użyciem odnawialnej energii przez elektrolizę wody. To jest jednak bardzo drogi sposób i takie paliwo nie jest w tej chwili ekonomiczną alternatywą. Większość wodoru w gospodarce jest pozyskiwana w znacznie mniej „czysty” sposób, z paliw kopalnych a ponadto jego dystrybucja i przechowywanie oznaczają kolejne emisje. Istnieją różne rodzaje paliwa wodorowego, do których klasyfikacji



2. Alpine Alpenglow HY4

używa się kolorów, szarego, niebieskiego i brązowego itd. Barwy odnoszą się do poziomu emisyjności przy jego wytwarzaniu. Niemniej, w sytuacjach, gdy napęd elektryczny się nie sprawdza lub wręcz jest niemożliwy do stosowania, wodór, nawet niekoniecznie zielony, jest opcją czystsza niż spaliniaki.

Samochody, silniki i jednostki latające

W ostatnich latach powstał szereg konstrukcji nowoczesnych aut z silnikami przystosowanymi do wydajnego spalania wodoru. Jedną z najlepszych jest francuski Alpine Alpenglow HY4 (2), dwulitrowy, czterocylindrowy silnik z turbodoładowaniem, o maksymalnych obrotach 7000 obr./min, mocy 340 koni mechanicznych. Ten koncepcyjny wóz, zaprezentowany w 2022 roku na targach w Paryżu, ma prowadzić do wariantu produkcyjnego przeznaczonego na drogi.

Austriacka firma AVL w 2023 r. zademonstrowała prototypowy dwulitrowy turbodoładowany silnik wodorowy (3) o mocy 200 KM na litr. AVL wykorzystuje system wtrysku wraz z turbosprężarką. System ten ma rozwiązywać znany problem silników spalających wodór wynikający ze spalania mieszanki zbyt ubogiej w paliwo, z dużą domieszką powietrza. Obniża to parametry silników, choć można to postrzegać jako zaletę, która poprawia oszczędność paliwa i obniża emisje. Jednak Austriacy myślą o autach sportowych, dlatego zależy im na wydajniejszym spalaniu wodoru. Niedawno zaprezentowali nową wersję AVL Racetech o podwojonej mocy – 400 KM.

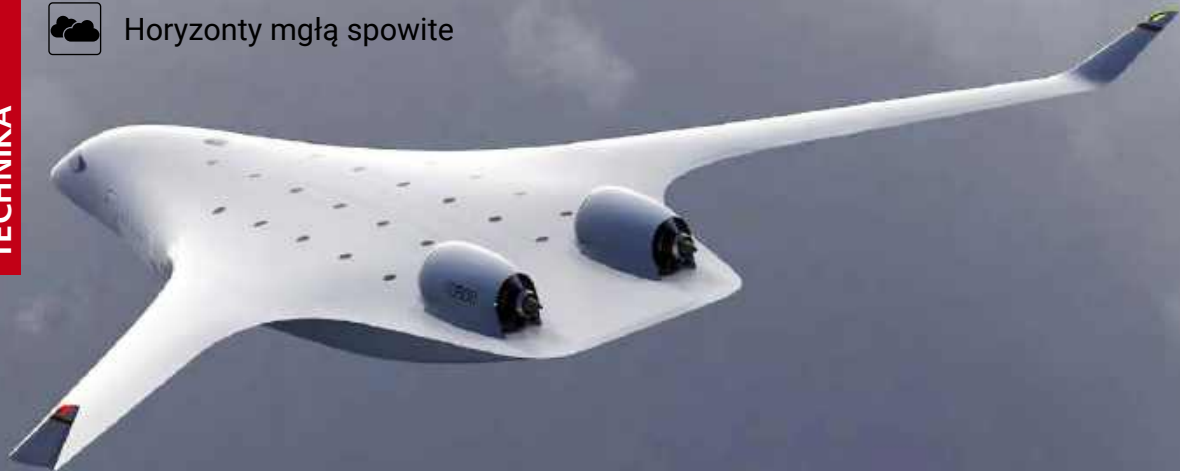
Także japońska Toyota, wieloletni zwolennik wodoru jako paliwa (choć była znana głównie z modeli z ogniwami paliwowymi), zaproponowała projekt GR H₂, który działa na zasadzie spalania wodoru, łącząc jednak silnik z układem hybrydowym, który zwiększa moc spalania wodoru. GR H₂ został zaprezentowany

w zeszłym roku tuż przed 24-godzinnym wyścigiem Le Mans. Toyota nie ujawniła żadnych bliższych szczegółów jednak wiadomo, że niemal jednocześnie złożyła patent na nową konstrukcję wodorowego silnika spalinowego z wtryskiem wody, który wtryskuje wodę do otworów wlotowych w celu chłodzenia silnika i kontrolowania nieprawidłowego spalania. Ponieważ technologia z samochodów wyścigowych może często przechodzić do samochodów drogowych, ten patent i prototyp GR H₂ sugerują, że Toyota szykuje coś ciekawego w wodorowej dziedzinie.

Inna znana japońska firma, Yamaha, ujawniła niedawno wodorową pochodną silnika V8 Lexusa RC F, zbudowanego zresztą dla Toyoty. Charakteryzuje się on modyfikacjami wtryskiwaczy, głowic cylindrów, kolektora dolotowego i nie tylko. Wszystko po to, by generować 450 KM przy 6800 obr./min. Silnik Yamahy niewiele odstaje od Lexusa RC F coupe z 2024 r., który



3. Wodorowy silnik spalinowy AVL Racetech



4. Wizualizacja konstrukcji JetZero BWB

osiąga 472 KM. Inne interesujące produkty wodorowe Yamahy to trójkołowy pojazd elektryczny Tricera i terenowy pojazd YXZ1000R, zaprezentowane na Japan Mobility Show 2023. Chociaż są to tylko koncepcje, mogą one zostać wykorzystane w przyszłych wersjach produkcyjnych.

Napęd wodorowy to nie tylko motoryzacja. Konstruktorzy samolotu JetZero obiecują, że dzięki zastosowaniu napędu wodorowego o połowę uda się zaoszczędzić koszty paliwa i emisje. Silnik lotniczy na wodór do tej futurystycznej konstrukcji (4) ma powstać we współpracy z Rolls-Royce'em.

Izraelski startup Heven Drones twierdzi, że jego latające drony napędzane wodorem dokonają przełomu w technice wojskowej. Jego napędzany tradycyjnie,

ciężki dron H100 firmy Heven Drones ma udźwignąć do około 30 kg i czas lotu od 22 do 55 minut, co stanowi dobre osiągi. Jednak, jak twierdzi startup, zasilanie wodorem może pięciokrotnie wydłużyć czas lotu w porównaniu z akumulatorami. Heven opracował trzy zasilane wodorem drony z nawigacją niezależną od GPS, które, według firmy, mogą przebywać w powietrzu do ponad 10 godzin z obciążeniem do ok. 10 kg (wersja oznaczona H₂D250).

MT opisywał wielokrotnie też projekty statków i pociągów zasilanych wodorem, gdyż przyjęło się, że do ciężkiego transportu to źródło energii lepiej się nadaje niż do lżejszych pojazdów. Być może opinia ta się zmienia. ■

Mirosław Usidus

Ludzie wędrowni

Anna Fryczkowska

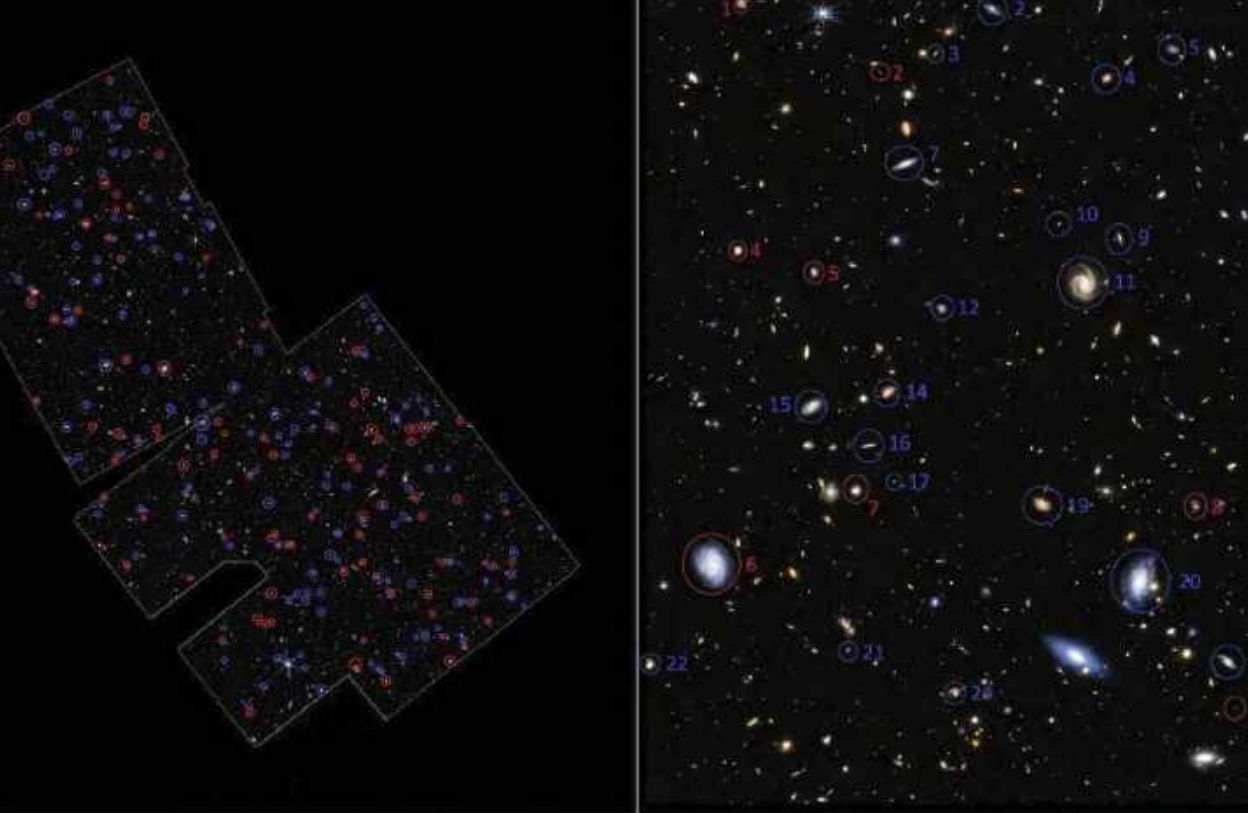
Wydawnictwo W.A.B., liczba stron: 376,

cena sugerowana: 54,99 zł

Nowa powieść autorski bestsellerowej trylogii „Saga o ludziach ziemi”! Fascynująca historia rodziny, która szukała swojego miejsca na ziemi.

Ludka nigdy nie wie, gdzie będzie jutro jadła i spała. Józka, żona rybaka podróżującego od jeziora do jeziora, raz po raz ląduje z dziećmi w zrujnowanych chatupach. Ina, najmocniej z nich zakorzeniona, wkrótce będzie musiała ruszyć w drogę. Trzy kobiety z kolejnych pokoleń maszerują przez powstania i wojny, usiłując utrzymać rodzinę w komplecie. Ich oczami obserwujemy rzeczywistość pierwszej połowy XX wieku – sławki i szable, witaminy i zarazki, duchy i koniokradów, objawienia i wysiedlenia.





1. Galaktyki obserwowane przez JWST – obracające się w jedną stronę zakreślone na czerwono, zaś obracające się w drugą stronę – na niebiesko

Znajdujemy się w czarnej dziurze
– mówią uczeni coraz częściej i głośniej

Czy Wszechświat ma odziedziczoną oś?

Czy nasz Wszechświat jest uwięziony w czarnej dziurze? Pomysł ten nie jest nowością. Od niedawna pojawiają się jednak w tej sprawie już nie tylko teoretyczne rozważania naukowców pragnących zwrócić na siebie uwagę. Są też zaskakujące wyniki obserwacji Kosmicznego Teleskopu Jamesa Webba.

Wyniki lustracji Wszechświata przez teleskop, który rozpoczął obserwację kosmosu latem 2022 roku, wskazują, że zdecydowana większość wczesnych galaktyk, z których obserwacji JWST jest szczególnie znany, obraca się według dającej do myślenia reguły.

Około dwóch trzecich z nich obraca się zgodnie z ruchem wskazówek zegara, pozostała zaś jedna trzecia obraca się przeciwnie do ruchu wskazówek zegara (1).

Dlaczego to dziwi? Otóż w „normalnym” (cokolwiek to znaczy) Wszechświecie naukowcy spodziewaliby



się raczej losowego rozkładu kierunków rotacji, co dałoby mniej więcej taki wynik, że połowa obiektów poruszałaby się w jedną a połowa w drugą stronę. Jednak obserwacje 263 galaktyk, przeprowadzone w ramach programu o nazwie James Webb Space Telescope Advanced Deep Extragalactic Survey, czyli JADES, sugerują, że istnieje preferowany kierunek obrotu galaktyk.

„Nie jest jasne, co jest powodem, że tak się dzieje, ale są dwa możliwe wyjaśnienia”, mówi w komunikacie naukowym szef zespołu prowadzącego te obserwacje, Lior Shamir z Carl R. Ice College of Engineering. „Jedno z wyjaśnień polega na tym, że Wszechświat narodził się w ruchu obrotowym. Teoria ta pokrywa się z kosmologią czarnych dziur, której konsekwencją może być hipoteza, że cały Wszechświat jest wnętrzem czarnej dziury”.

Z „kosmologii Schwarzschilda” wynikać ma, że nasz obserwowalny Wszechświat może być wnętrzem czarnej dziury w większym wszechświecie macierzystym. Pomysł ten został po raz pierwszy przedstawiony przez fizyka teoretycznego Raja Kumara Pathrię i matematyka I. J. Gooda. Przedstawia ona ideę, że „promień Schwarzschilda”, znany też jako „horyzont zdarzeń” (granica, poza którą nic nie może uciec z czarnej dziury, nawet światło), jest również horyzontem widzialnego Wszechświata. Ma to jeszcze taką implikację, że każda czarna dziura w naszym Wszechświecie może być wrotami do innego „małego wszechświata”. Wszechświaty te byłyby dla nas nieobserwowalne, ponieważ również znajdują się za horyzontem zdarzeń, z którego światło nie może uciec, co oznacza, że informacje nigdy nie mogą podróżować z wnętrza czarnej dziury do zewnętrznego obserwatora.

Znanym orędownikiem tej teorii jest też polski fizyk teoretyczny Nikodem Popławski z uniwersytetu w New Haven. Według Popławskiego ostatecznie sprzężenie między skręcaniem i obrotami materii staje się bardzo silne i zapobiega kompresji materii w nieskończoność do stanu osobliwości. „Zamiast tego materia osiąga stan skończonej, niezwykle dużej gęstości, przestaje się zapadać, odbija się jak ściśnięta sprężyna i zaczyna gwałtownie się rozszerzać”, tłumaczył Popławski portalowi Space.com. „Potężne siły grawitacyjne w pobliżu tego stanu powodują intensywną produkcję cząstek, zwiększając masę wewnątrz czarnej dziury o wiele rzędów wielkości i wzmacniając odpychanie grawitacyjne, które napędza odbicie”. Naukowiec dodaje, że szybki odrzut po tak dużym odbiciu może być tym, co doprowadziło do powstania naszego rozszerzającego się Wszechświata, w wydaniu, które nazywamy Wielkim Wybuchem.

„Prowadzi to do skończonego okresu kosmicznej inflacji, co wyjaśnia, dlaczego Wszechświat, który obserwujemy dzisiaj, wydaje się w największych skalach płaski, jednorodny i izotropowy”, mówi Popławski. „Skręcenie w grawitacji rozszerzonej teorii ogólnej teorii względności Einsteina zapewnia zatem wiarygodne teoretyczne wyjaśnienie scenariusza, zgodnie z którym każda czarna dziura wytwarza wewnątrz nowy, mały wszechświat i staje się mostem Einsteina-Rosena lub ‘tunelem czasoprzestrzennym’, który łączy ten Wszechświat z wszechświatem macierzystym, w którym istnieje czarna dziura”.

W nowym wszechświecie, zgodnie z tą koncepcją, wszechświat macierzysty pojawia się jako druga strona jedynej białej dziury nowego wszechświata, obszaru przestrzeni, do którego nie można wejść z zewnątrz i który można uznać za odwrotność czarnej dziury. „Ruch materii przez granicę czarnej dziury, zwaną horyzontem zdarzeń, może odbywać się tylko w jednym kierunku, zapewniając asymetrię przeszłość-przyszłość na horyzoncie, a tym samym wszędzie w małym wszechświecie. Strzałka czasu w takim wszechświecie zostałaby zatem odziedziczona po wszechświecie macierzystym przez skręcanie”, kontynuuje Popławski. „Byłoby fascynujące, gdyby nasz Wszechświat miał preferowaną oś obrotu. Taka oś mogłaby być naturalnie wyjaśniona przez teorię, że nasz wszechświat narodził się po drugiej stronie horyzontu zdarzeń czarnej dziury istniejącej w jakimś wszechświecie macierzystym (...) Skręcanie czasoprzestrzeni zapewnia najbardziej naturalny mechanizm, który pozwala uniknąć osobliwości w czarnej dziurze i zamiast tego tworzy nowy, zamknięty wszechświat. Preferowana oś obrotu w naszym Wszechświecie, odziedziczona przez oś obrotu macierzystej czarnej dziury, mogła wpłynąć na dynamikę rotacji galaktyk, tworząc asymetrię ruchu wskazówek zegara w kierunku przeciwnym do ruchu wskazówek zegara”.

Jak pamiętamy, Lior Shamir wspominał też o drugim możliwym wyjaśnieniu zaobserwowanej asymetrii w obrotach galaktyk. Polega ono na tym, że efekt ten mogła spowodować własna rotacja naszej Drogi Mlecznej. Wcześniej naukowcy uważali, że prędkość rotacji naszej galaktyki jest zbyt wolna, aby mieć istotny wpływ na obserwacje JWST, ale nie można takiego wyjaśnienia asymetrii jednak całkiem wykluczyć. Dlatego badacze, którzy opisali swoje odkrycia w „Monthly Notices of the Royal Astronomical Society”, zapowiadają ponowną weryfikację i kalibrację pomiarów ruchu galaktyk. ■

Mirosław Usidus

OSOBLIWOŚĆ JUŻ?

AI nas wyzwoli albo... nie



„Żadni agenci AI nie są dozwoleni”. Taką zasadę wprowadziła niedawno jedna z instytucji Unii Europejskiej. Zabrania wirtualnym asystentom opartym na sztucznej inteligencji uczestniczenia w jej posiedzeniach.

**Sztuczna inteligencja umiała już odpowiadać
– teraz uczy się działać**

AGENTURA AI

Z drugiej strony, bez rozgłosu, Bruksela jednak przygotowuje się do ery, w której „agenci AI” będą uczestniczyć w codziennym życiu i biznesie. Technika ta została wspomniana w szerszym pakiecie planów regulacyjnych Komisji Europejskiej dotyczącym wirtualnych światów, opublikowanym w marcu 2025 r. „Agenci AI to aplikacje zaprojektowane do postrzegania i interakcji ze środowiskiem wirtualnym”, czytamy w tekście. Jak dotąd ta konkretnie technika sztucznej inteligencji nie jest objęta żadnymi szczegółowymi przepisami, choć, jak wiadomo, UE specjalizuje się w obejmowaniu wszystkiego przepisami. Tak czy inaczej zresztą, modele sztucznej inteligencji, które zasilają agentów AI, będą musiały być zgodne z obowiązującym europejskim „AI Act”.

Widzi, co wyświetla ekran komputera

W innych częściach świata rozmawia się mniej o przepisach, więcej o tym, jakie korzyści mogą przynieść agenci sztucznej inteligencji (1), ale także, ma się rozumieć, jakie z ich ekspansją wiążą się problemy. Giganci z branży technologicznej dość zgodnie twierdzą, że agenci AI zapewnią firmom znaczny wzrost produktywności. Według różnych ankiet przeprowadzanych wśród menedżerów, ok. jednej trzeciej planuje sięgnąć po agentów AI w swojej działalności ciągu najbliższych dwóch lat.

Bill Gates przewidywał, że „wszyscy będziemy mieli agenta AI”, który pomoże nam się zorganizować. Satya Nadella, obecny szef Microsoftu, który Gates założył, zapowiada, że nowe jej rozwiązania w ramach pakietu „Computer Use” w Copilot Studio (2)



1. Agenci AI przy pracy © AI

umożliwiają każdemu łatwe tworzenie agentów AI, które będą klikać, otwierać strony i przeglądać menu w interfejsach aplikacji. Narzędzie, oparte na dużym modelu językowym, może szybko dostosowywać się, gdy zmienia się układ aplikacji lub strony internetowej. Użytkownicy nie muszą pisać żadnego kodu – wystarczy, by wpisali to, czego chcą, w prostym języku i przekazali instrukcje do narzędzia. Wcześniej Microsoft uruchomił inne narzędzie, Dragon Copilot, asystenta opartego na sztucznej inteligencji, opracowanego w celu uproszczenia przepływów pracy. W grudniu 2024 zaprezentował też Copilot Vision, innowacyjną funkcję opartą na sztucznej inteligencji zintegrowaną z przeglądarką Edge. Narzędzie to, dostępne w ograniczonej wersji zapoznawczej dla subskrybentów Copilot Pro w Stanach Zjednoczonych, ma na celu wypracowanie nowych form interakcji z siecią. Oferując pomoc kontekstową, analizę wizualną i spersonalizowane rekomendacje, Copilot Vision upraszcza zadania, od gier



Prezentacja agenta
Copilot Microsoftu:
<https://youtu.be/uBhUHSdk-4A>



2. Copilot Studio

po planowanie podróży, eliminując potrzebę przełączania kart lub wykonywania dodatkowych wyszukiwań. Jest zdolny do analizy wizualnej, interpretacji obrazu i identyfikacji punktów orientacyjnych. Sztuczna inteligencja analizuje zawartość ekranu i zapewnia przydatne informacje. Na przykład, jeśli czytamy o produkcie, może podkreślić jego kluczowe cechy, zasugerować powiązane elementy, a nawet porównać alternatywy.

Współpracująca od lat z Microsoftem firma OpenAI uruchomiła podobne narzędzie testowe o nazwie AI Operator (3) w styczniu tego roku. To aplikacja, która może wykonywać proste zadania online w przeglądarce, np. rezerwację biletów na koncert lub złożenie zamówienia w sklepie online. Jest obsługiwana przez nowy model o nazwie Computer-Using Agent (CUA), zbudowany na dużym modelu językowym (LLM) GPT-4o. Operator był dostępny dla użytkowników w USA, zarejestrowanych w ChatGPT Pro, usłudze premium OpenAI kosztującej

200 dolarów miesięcznie. Jest plan udostępnienia tego narzędzia innym użytkownikom w przyszłości. OpenAI zapewnia, że Operator przewyższa podobne konkurencyjne narzędzia, w tym Computer Use firmy Anthropic (wersja Claude 3.5 Sonnet, która



Agent Claude
w komputerze osobistym:
<https://youtu.be/d5IKjrmUaY>

może wykonywać proste zadania na komputerze) i Marinera (agent przeglądania stron internetowych zbudowany na bazie Gemini 2.0 Google DeepMind). Podobnie jak wszystkie najbardziej zaawansowane narzędzia nazywane agentami AI, Operator wykonuje zrzuty ekranu komputera i skanuje piksele w celu rozpoznania treści, planowania i egzekucji zadań. CUA jest specjalnie przeszkolony do interakcji z graficznymi interfejsami użytkownika, przyciskami, polami tekstowymi, menu, czyli tymi samymi, których używają ludzie. Tradycyjnie modele korzystały z oprogramowania za pośrednictwem wyspecjalizowanych API, czyli interfejsów programowania aplikacji, fragmentów kodu, które działają jak rodzaj łącznika i pośrednika uzgadniającego, umożliwiającego łączenie ze sobą różnych bloków oprogramowania i systemów. Wymóg używania API był dotychczas barierą sprawiającą, że wiele aplikacji i większość stron internetowych była niedostępna dla AI. Modele wyszkolone do rozpoznawania interfejsów wizualnych omijają ten problem. To ważna, niedostatecznie eksponowana, innowacja.

OpenAI twierdzi, że model CUA został przeszkolony przy użyciu technik podobnych do tych stosowanych

3. Prezentacja Operatora firmy Open AI



w tak zwanych modelach rozumujących, o1 i o3. Firma twierdzi, że jej model pokonuje konkurencję, czyli Computer Use Anthropic i Marinera Google we wszystkich popularnych testach porównawczych. Na przykład w OSWorld, który sprawdza, jak dobrze agent wykonuje zadania, takie jak łączenie plików PDF czy manipulowanie obrazem, CUA uzyskuje wynik 38,1 proc., gdy Computer Use ma wynik 22 proc. Ludzie uzyskują w nim wynik 72,4 proc., czyli wciąż lepszy niż AI, ale nie miażdżący. W teście o nazwie WebVoyager, który sprawdza, jak dobrze agent wykonuje zadania w przeglądarce, CUA uzyskuje 87 proc., Mariner 83,5 proc., a Computer Use 56 proc. Na razie, nawiasem mówiąc, Operator może dla użytkowników wykonywać zadania tylko w przeglądarce. OpenAI planuje udostępnić szersze możliwości CUA w przyszłości za pośrednictwem interfejsu API, którego inni programiści mogą używać do tworzenia własnych aplikacji.

W podobny sposób jak Open AI w grudniu Anthropic udostępnił swojego wspomnianego tu już kilka razy agenta Computer Use, który opiera się na Artifacts, rozwiązaniu, które pozwala chatbotowi Claude uruchamiać kod w przeglądarce. Na filmie demonstracyjnym pokazano takie jego możliwości, jak czytanie ekranu, klikanie przycisków pobierania, automatyczną edycję kodu i inne działania. Firma kieruje nową funkcję wyłącznie do tych, którzy korzystają z API, a nie w chatbotcie. Celem jest usunięcie wielu uciążliwości, przez które ludzie muszą przechodzić, realizując różne zadania na swoich komputerach. Na jednym z filmów demonstracyjnych na stronie startowej prezytera demonstruje, jak Claude tworzy i edytuje własną stronę internetową.

Inni nie chcą zostać w tyle. Chris Cox, dyrektor ds. produktów w firmie Meta, zapowiedział, że nowe oprogramowanie open source Llama 4 AI także pomoże zasilać agenturę AI wykonującą podobne zadania jak opisane wyżej narzędzia innych firm. Interesującym elementem w planach Meta jest to, że chciałaby swoich agentów udostępnić po przystępnych cenach małym firmom, których często nie stać na wielkie kompleksowe systemy i modele AI. Także Amazon, który zresztą współpracuje i współfinansuje Anthropic, zaprezentował niedawno własną propozycję o nazwie Nova Act, agenta AI na wczesnym etapie rozwoju, który może wykonywać zadania na stronach internetowych. Nawet Samsung niedawno ogłosił, że jego OneUI 7 oferuje wiele funkcji „Agentic AI” wbudowanych w system operacyjny. Firma Adobe także wchodzi na rynek agentów i dodaje więcej funkcji personalizacji do codziennych zadań związanych

z obsługą klienta. W marcu 2025 r. ogłosiła uruchomienie dziesięciu rodzajów narzędzi typu agencji i „narzędzia do orkiestracji” na swojej platformie Adobe Experience. Loni Stark, wiceprezes ds. strategii i produktów w Adobe, powiedziała serwisowi „VentureBeat” w wywiadzie, że chodzi o to, by pozwolić tym agentom pracować w otoczeniu, co oznacza, że agencji i orkiestrator pracują w tle, by dostarczać informacje lub rozwiązywać problemy klientów. Stark zwróciła uwagę na agenta optymalizacji stron WWW, który sprawdza niedziałające linki lub bada strony pod kątem ruchu i współczynnika odrzuceń oraz sugeruje poprawki. „Większość firm nie ma ludzi, którzy spędzają całe dnie na przykład na sprawdzaniu niedziałających linków, zwłaszcza jeśli mają dziesiątki tysięcy stron lub nie mogą ich codziennie sprawdzać”, powiedziała. „Ten agent jest wstępnie przeszkolony, więc po wyjęciu z pudełka ma już umiejętności, takie jak wyszukiwanie uszkodzonych linków zwrotnych”. Kolejną nową funkcją Adobe Experience Platform jest Brand Concierge, która ma pomóc firmom tworzyć strony internetowe oferujące spersonalizowane wizyty klientów.

Gwiazdą wiosny 2025 był Manus, autonomiczny „agent ogólny” stworzony przez chiński startup AI Monica (4). To autonomiczny agent sztucznej inteligencji, który opiera się na modelach Claude Sonnet i Qwen finetunes chińskiej Alibaby. Twórcy Manusa twierdzą, że jest pierwszym na świecie ogólnym agentem sztucznej inteligencji, wykorzystującym wiele modeli AI oraz różnych niezależnie działających agentów do autonomicznego działania w szerokim zakresie zadań. Podobnie jak inne oparte na rozumowaniu narzędzia agentowej sztucznej inteligencji, takie jak ChatGPT DeepResearch, Manus jest w stanie podzielić zadania na etapy i autonomicznie poruszać się po sieci, aby uzyskać informacje potrzebne do ich wykonania. Okienko pod nazwą Komputer



Wprowadzenie do Manusa:
<https://youtu.be/K27diMbCsuw>



4. Logotyp agenta Manus

Manusa, pozwala użytkownikom nie tylko obserwować, co robi agent, ale także interweniować w dowolnym momencie. Ma jednak wyższy wskaźnik awaryjności niż ChatGPT DeepResearch. Chińskie media donosiły jednak, że koszt jednego zadania realizowanego przez Manusa wynosi około 2 dolarów, co stanowi zaledwie jedną dziesiątą kosztu DeepResearch. W udostępnionych demonstracjach Manus był wykorzystywany m.in. do sprawdzania CV kandydatów do pracy zawartych w pliku .zip i prezentowania wyników w arkuszu kalkulacyjnym, analizy rynku nieruchomości w Nowym Jorku w oparciu o określone kryteria i pisanie na tej podstawie własnego programu w Pythonie, obliczającego budżet użytkownika i znajdującego odpowiednie zakwaterowanie, z przygotowaniem raportu końcowego, a także do zarządzania portfolio akcji na giełdzie. Był też używany do tworzenia gier, animacji, edycji podcastów, a nawet do projektowania całych stron internetowych. Dostęp do narzędzia jest ograniczony do osób z zaproszeniem. Zainteresowani sprawdzeniem Manusa mogą dołączyć do listy oczekujących na prywatną wersję beta. Popularność Manusa rozprzestrzeniła się w Internecie przez krótki okres jak pożar, nie tylko w Chinach, gdzie został opracowany przez startup Butterfly Effect z siedzibą w Wuhan. Niektórzy nazwali go nawet „drugim DeepSeekiem”. Jednak, jak można się spodziewać w przypadku każdego rozwiązania pochodzącego z Chin, od razu zrodziły się obawy dotyczące przejrzystości i prywatności danych.

Agent musi mieć dostęp do prywatnych danych i to jest chyba problem

Jak oceniają specjaliści za prawdziwy przełom w tej nowej fali „agenckiej AI” trzeba uważać połączenie sztucznej inteligencji z konkretnymi zadaniami, praktycznymi operacjami, które najczęściej ludziom zajmują sporo czasu, są żmudne i zwykle nie lubiane. Tym właśnie mają zająć się agenci AI, pracując non stop, nie biorąc wolnego, nie popełniając głupich, typowo ludzkich, błędów. Obmyślanie promptów dotyczących każdego odrębnego działania staje się w tym świecie zbędne. Gdy ogólna sztuczna inteligencja odpowiada na pytania, agent AI działa. Nie mówi użytkownikowi, jak ustawiać spotkania – po prostu ustawia je dla użytkownika. Nie tylko przygotowuje e-maile, ale wysyła je. Nie tylko zapamiętuje preferencje, ale wykorzystuje je do personalizacji operacji i realizowanych działań. Jeśli poprosimy zwykły model AI o zaplanowanie rezerwacji na kolację o 21.00, ten dokona rezerwacji, ale nie będzie w stanie zrobić nic więcej. W przypadku

agenta AI można przekazać mu niejasne instrukcje, takie jak „Zrób rezerwację na kolację z Basią w jakiejś włoskiej restauracji”. Agent nie tylko dokona rezerwacji kolacji w restauracji, ale także sprawdzi harmonogram wydającego zlecenie i rozkład zajęć Basi, by sprawdzić, która data jest najlepsza dla obojga. Agent robi to, nie tylko przeszukując internet w poszukiwaniu najlepszych miejsc, ale także przeglądając osobiste dane, w tym przypadku kalendarze. To dodatkowe zestawy danych, do których musi oczywiście uzyskać dostęp. Jeśli, biorąc inny przykład, typowy dla działalności firm, klient skontaktuje się z działem pomocy technicznej firmy w celu uzyskania pomocy, zaś agent AI uzyskał dostęp do dokumentów firmy i innych danych, może (samodzielnie) określić najlepszy sposób rozwiązania problemu. Może oczywiście również zdecydować, czy może samodzielnie rozwiązać problem, czy też jednak to człowiek byłby bardziej odpowiedni do tego zadania.

Kluczową różnicą jest nie tylko to, że „agencka” sztuczna inteligencja ma dostęp do dodatkowych danych, który musi zostać jej przyznany, ale może uczyć się na podstawie swoich „doświadczeń”, interakcji, informacji w sieci, a nie tylko na podstawie zbiorów informacji, na których została wstępnie przeszkolona. Może działać autonomicznie i podejmować decyzje na bazie szerszych możliwości rozumowania, kontekstu i dedukcji. Teoretycznie może również wchodzić w interakcje z innymi agentami sztucznej inteligencji w celu wykonania określonego zadania, co brzmi już nieco jak fantastyka naukowa i każe zastanowić się nad konsekwencjami takich systemów.

Co do zasady, agent AI działa według trzystopniowego schematu: rozumie, myśli i działa. Po pierwsze, przyjmuje polecenia. Po drugie, przetwarza żądanie i ustala najlepsze podejście. Po trzecie, podejmuje działania, aby wykonać zadanie. Wymaga to podłączenia do narzędzi, z których korzysta użytkownik, kalendarza, skrzynki e-mail, systemu zarządzania relacjami z klientem, mediów społecznościowych i innych. Agenci AI składają się z komponentów takich jak pamięć, planowanie, nauka maszynowa, postrzeganie i wykonywanie. Po otrzymaniu zadania AI może korzystać z różnych zasobów, w tym LLM (duży model językowy), dzienników danych, wyszukiwań internetowych. Model przegląda swoją pamięć w poszukiwaniu informacji związanych z zadaniem, a gdy ich tam nie ma, potrzebny jest dostęp do internetu. Sztuczna inteligencja może zaplanować najlepszy możliwy sposób wykonania zadania, a po jego wykonaniu przechowywać procedurę, której się nauczył w pamięci krótko- i długoterminowej. Agencka sztuczna inteligencja ma



5. Meredith Whittaker z Signala

również rozpoznać swoje niedociągnięcia i dokonać krytycznej samooceny, aby lepiej wykonać zadanie w przyszłości.

Jednak choć wszystko to brzmi w teorii świetnie, nie brakuje wyzwań. „Agentic AI” wciąż wymaga nadzoru ze strony człowieka. Wymaga również ogromnej mocy obliczeniowej, by działać. Dużym wyzwaniem związanym z agentami sztucznej inteligencji są z pewnością uprzedzenia i obawy etyczne związane ze szkoleniem modeli. Agenci AI mogą być podatni na uprzedzenia wynikające ze zbiorów danych, na których są szkoleni, a także na te pochodzące z internetu. Istnieje ryzyko związane z bezpieczeństwem, ponieważ agenci ci mogą być podatni na cyberataki i naruszenia danych. W marcu 2025 r. szefowa firmy rozwijającej komunikator Signal, Meredith Whittaker (5), ostrzegała na konferencji w Austin w Teksasie, że „agentowa” sztuczna inteligencja może stanowić ogromne zagrożenie dla prywatności użytkowników, którzy dadzą jej dostęp do wszystkiego, do najważniejszych i najbardziej prywatnych swoich danych. Zwróciła ponadto uwagę na inne zagrożenie dla ludzi, którzy wszystko scedują na AI. „Czyli, co? Możemy po prostu włożyć nasz mózg

do słoika, ponieważ ta rzecz [agent AI – przyp. MT] wszystko robi i nie musimy niczego dotykać?”, pytała ironicznie.

Naszą specjalnością jest „mieszanka ekspercka”

Trudno jednak zaprzeczyć, że rozwiązania agencji AI mają swoje ogromne zalety. Duże modele generatywne są ustalone w czasie, niezdolne do włączenia nowych lub dynamicznych informacji. Dostrajanie sztucznej inteligencji do wyspecjalizowanej i nowej tematyki może zaspokoić potrzeby specyficzne dla danej domeny, ale jest kosztowne i podatne na błędy. Wymaga ogromnych ilości danych, znacznych zasobów obliczeniowych i wiedzy z zakresu uczenia maszynowego, co czyni je niepraktycznym w wielu sytuacjach. Załóżmy na przykład, że model generatywny ma zarekomendować polisę ubezpieczeniową dostosowaną do osobistej historii zdrowotnej, lokalizacji i preferencji finansowych użytkownika. W tym scenariuszu zapytany LLM generuje odpowiedź, ale nie może dostarczyć dokładnych rekomendacji, ponieważ nie ma dostępu do odpowiednich danych użytkownika. Bez tego odpowiedź będzie albo ogólna, albo całkowicie błędna. Aby przezwyciężyć te ograniczenia, opracowano techniki integracji ogólnych modeli generatywnych z innymi komponentami, logiką programistyczną, mechanizmami pobierania danych i kolejne warstwy weryfikujące. Ta modułowa konstrukcja pozwala sztucznej inteligencji łączyć narzędzia, pobierać odpowiednie dane i dostosowywać wyniki w sposób, w jaki nie mogą tego zrobić modele statyczne.

W przykładzie rekomendacji ubezpieczenia mechanizm pobierania pobiera dane zdrowotne i finansowe użytkownika z bezpiecznej bazy danych. Dane te są dodawane do kontekstu dostarczanego do LLM podczas tworzenia odpowiedzi. LLM wykorzystuje złożony monit do wygenerowania dokładnej odpowiedzi. Proces ten, znany jako Retrieval-Augmented Generation (RAG), wypełnia lukę między statyczną, ogólną, sztuczną inteligencją a rzeczywistymi potrzebami. Chociaż RAG skutecznie radzi sobie z takimi zadaniami, wciąż opiera się na stałych, z góry zaprogramowanych przepływach pracy, co oznacza, że każda interakcja i ścieżka wykonania muszą być wstępnie zdefiniowane. Ta sztywność sprawia, że znów trudno tu radzić sobie z bardziej złożonymi lub dynamicznymi zadaniami, w których przepływy pracy nie mogą być w sposób pełny zakodowane. Ręczne kodowanie wszystkich możliwych ścieżek wykonania jest z kolei pracochłonne.

Ograniczenia architektur o stałym przepływie doprowadziły do powstania trzeciej fali sztucznej inteligencji, czyli właśnie systemów agencjonalnych, które są głównym tematem tego artykułu. Konwencjonalna RAG, odnosi się do zewnętrznych, autorytatywnych baz wiedzy przed wygenerowaniem odpowiedzi w celu poprawy wyników LLM. RAG zapewnia aktualność informacji przez nawiązywanie połączeń ze strumieniami na żywo i regularnie aktualizowanymi źródłami. Dodając autonomicznych agentów RAG rozszerza swoje możliwości. Dzięki tej transformacji statyczny system RAG staje się dynamiczną, świadomą kontekstu sztuczną inteligencją, która może odpowiadać na skomplikowane pytania spójnie i precyzyjnie.

Szef firmy Salesforce, Marc Benioff, powiedział niedawno dziennikowi „The Wall Street Journal”, że ponieważ osiągnęliśmy górną granicę tego, co mogą zrobić LLM, przyszłość leży w autonomicznych agentach – systemach, które mogą myśleć, dostosowywać się i działać niezależnie, a nie w modelach takich jak GPT-4. Agenci AI wnoszą dynamiczne, kontekstowe przepływy pracy. W przeciwieństwie do ustalonych ścieżek, systemy agentowe określają kolejne kroki w locie, dostosowując się do aktualnej sytuacji. Dzięki temu nadają się do rozwiązywania nieprzewidywalnych, wzajemnie powiązanych problemów. Zamiast sztywnych procedur dyktujących każdy ruch, agenci wykorzystują LLM do podejmowania samodzielnych decyzji. Mogą rozmawiać, korzystać z narzędzi i używać dostępu do pamięci. Wszystko dzieje się dynamicznie. Ta elastyczność pozwala na przepływy pracy, które ewoluują w czasie rzeczywistym. Agenci mogą i muszą zbierać informacje z wielu źródeł, w tym od innych agentów, narzędzi i systemów zewnętrznych, w celu podejmowania decyzji i działań. Zakładaną funkcją agentów AI jest zdolność do autorefleksji, która pozwala agentom oceniać własne decyzje i poprawiać swoje wyniki przed podjęciem działań lub udzieleniem ostatecznej odpowiedzi. Ma to pozwalać agentom na wychwytywanie i poprawianie błędów, udoskonalanie ich rozumowania i zapewnianie wyższej jakości wyników.

Istnieją też już systemy wieloagentowe, które przyjmują modułowe podejście do rozwiązywania problemów, przypisując konkretne zadania wyspecjalizowanym agentom. Można używać mniejszych modeli językowych dla agentów specyficznych dla poszczególnych zadań, co ma poprawiać wydajność i upraszczać zarządzanie pamięcią. Modułowa konstrukcja zmniejsza złożoność poszczególnych agentów, utrzymując ich koncentrację na ich konkretnych

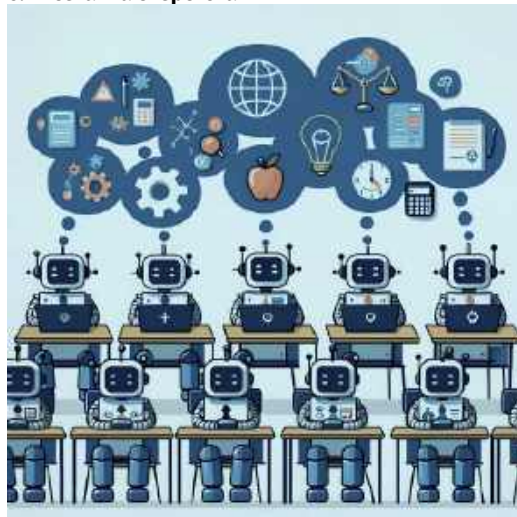
zadaniach. Odmianą tej techniki jest Mixture-of-Experts (MoE, z ang. „mieszanka ekspercka”), która wykorzystuje wyspecjalizowane podmodele lub „ekspertów” w ramach jednej struktury (6). MoE dynamicznie kieruje zadania do najbardziej odpowiedniego eksperta, optymalizując zasoby obliczeniowe i zwiększając wydajność. Wyspecjalizowani agenci dzielą się informacjami, dzielą obowiązki i koordynują działania.

Osoby posiadające odpowiednią wiedzę i umiejętności mogą same skonfigurować własnego agenta AI. Jest sporo narzędzi do tego służących mających najczęściej charakter open source, jednak zwykle wymagających określonych opłat. Przykłady to crewAI, framework typu open source zaprojektowany w celu ułatwienia tworzenia zaawansowanych wieloagentowych systemów sztucznej inteligencji, AutoGen, opracowany przez Microsoft framework typu open source z architekturą wieloagentową do rozwiązywania złożonych problemów, LangChain z modułową i rozszerzalną architekturą, Vertex AI Agent Builder, produkt Google Cloud, który nie wymaga głębokiej wiedzy z zakresu uczenia maszynowego i oferuje opcje tworzenia agentów bez konieczności kodowania, Cogniflow, umożliwiające użytkownikom tworzenie i wdrażanie modeli sztucznej inteligencji bez specjalistycznej wiedzy programistycznej.

Tam gdzie agenci się już uwiązają w znoju

Clara Shih, szefowa działu AI w firmie Meta, powiedziała niedawno telewizji CNBC, że spodziewa się, że „każda firma, od bardzo dużej do bardzo małej”

6. Mieszanka ekspercka AI





7. Ekran z interfejsem GS AI Assistant banku Goldman Sachs

będzie wkrótce reprezentowana przez agenta AI. W badaniach firmy PagerDuty 62 proc. respondentów ankiety przewiduje trzycyfrowy zwrot z inwestycji w agenturalną AI. Jednak niektórzy analitycy, tacy jak Jim Covello z Goldman Sachs, podchodzą do tego bardziej sceptycznie, argumentując, że kosztowne inwestycje w agentową sztuczną inteligencję mogą ostatecznie przynieść zaskakująco skromne zyski.

Menedżerowie z Goldman Sachs mogą być sceptyczni wobec agentów AI, jednak, jak podał bank w styczniu 2025 r., wprowadza generatywnego asystenta AI dla swoich bankierów, traderów i zarządzających aktywami, co jest pierwszym etapem projektu, który ostatecznie zmierza do tego, by ów „agent” nabrał cech i umiejętności doświadczonego pracownika Goldmana. Bank udostępnił do tej pory program o nazwie GS AI Assistant (7) około dziesięciu tysiącom pracowników i zmierza do tego, by wszyscy pracownicy firmy korzystali z tego asystenta już w tym roku. Początkowo pomoże on w zadaniach obejmujących podsumowywanie lub korektę wiadomości e-mail lub tłumaczenie kodu z jednego języka na inny. W pierwszej fazie narzędzi będzie głównie generować odpowiedzi na podstawie danych Goldmana, które zostały wprowadzone do modeli sztucznej inteligencji z ChatGPT OpenAI, Google Gemini i Meta Llama, w zależności od zadania. Bank przygląda się również modelom firm takich jak Anthropic, Mistral i Cohere „W miarę postępów, w kolejnym kroku nadejdzie moment, w którym pojawi się to agentowe zachowanie – »wykonuję zadanie w imieniu pracownika Goldmana«, mówił Marco Argenti, rzecznik banku, serwisowi CNBC. „W tym miejscu model zacznie robić to co pracownik Goldmana, a nie tylko mówić jak pracownik Goldmana”.

Bank Goldman Sachs, wraz z innymi znanymi instytucjami finansowymi JPMorgan Chase i Morgan Stanley, odważnie udostępnia swoim pracownikom narzędzia generatywnej sztucznej inteligencji. Eksperci zauważają, że nowojorska finansjera z Wall Street przyjęła generatywną sztuczną inteligencję szybciej niż jakąkolwiek inną przełomową technologię w ostatnich latach. Jak wynika z raportu działu badawczego Bloomberg, globalne banki inwestycyjne mogą zlikwidować nawet 200 tysięcy miejsc pracy w ciągu najbliższych trzech do pięciu lat wskutek wdrażania sztucznej inteligencji. Oficjalne stanowisko Goldmana jest obecnie takie, że sztuczna inteligencja pozwoli pracownikom robić więcej, niekoniecznie powodując potrzebę zmniejszenia liczby ludzi.

Współzałożyciel OpenAI, Sam Altman i inni liderzy branży twierdzą, że rok 2025 będzie rokiem, w którym agenci AI „dołączą” do siły roboczej. W ankiecie Deloitte opublikowanej w marcu 2025 r. 26 proc. liderów stwierdziło, że ich organizacje poważnie badają autonomicznych agentów. ServiceNow, SAP i Salesforce to kolejne firmy, które już wprowadziły agentów AI do wykonywania zadań w pracy. Kiedy klienci dostawcy oprogramowania w chmurze ServiceNow kontaktują się z centrum obsługi klienta firmy, 80 proc. spraw, w formie połączeń i wiadomości na czacie, jest obsługiwanych bez interwencji człowieka, przez agenturę AI. Chris Bedi, z ServiceNow, doradca ds. sztucznej inteligencji w przedsiębiorstwie, podaje, że ludzie nadal obsługują co piąte żądanie obsługi klienta. Także ludzie ostatecznie zatwierdzają działania agencjonalnej sztucznej inteligencji, przed egzekucją. ServiceNow twierdzi, że połączenie pracowników-ludzi i AI skróciło czas potrzebny na obsługę bardziej złożonych spraw o 52 proc. w okresie dwóch

tygodni. Już we wrześniu ub. roku Salesforce wdrożyło Agentforce (8), własną platformę agencji sztucznej inteligencji, która m.in. zajmuje się przetwarzaniem faktur, zapewnianiem wsparcia klientom i redagowaniem wiadomości e-mail. Salesforce zapowiadało, że do końca 2025 r. udostępni klientom miliard agentów AI. Firma poinformowała również, że na ponad 340 tys. pytań dotyczących obsługi klienta udzielono autonomicznych odpowiedzi za pomocą Agentforce. Gigant oprogramowania biznesowego, Intuit, rozpoczął wdrażanie agencji sztucznej inteligencji w grudniu. Na razie w większości przypadków w systemie kluczową rolę odgrywają ludzie, ale przewiduje się wkrótce całkowitą automatyzację, także za pomocą systemów multiagentów.

Eksperti uspokajają pracowników, że agenci AI nie zastąpią ich i nie zabiorą im pracy, bo przejmą odpowiedzialność za przyziemne, żmudne i nie lubiane zadania. „Ci agenci pomogą mi wykonywać moją pracę, ale w żadnym momencie nie zmuszą mnie do zrobienia czegoś, czego nie jestem świadomy”, mówił w wypowiedzi dla mediów Walter Sun, szef AI w firmie SAP, która sprzedaje oprogramowanie do zarządzania finansami, łańcuchem dostaw i innymi potrzebami biznesowymi.

Agregacja agregatorów

Mark Shmulik i Nikhil Devnani z firmy Bernstein, dwaj czołowi analitycy internetowi, napisali w nocy do inwestorów niedawno: „Jeśli agenci AI naprawdę

staną się użyteczni, internet pograży się w mroku”. Należy to rozumieć w ten sposób, że konsumenci nie będą odwiedzać ani widzieć w sposób bezpośredni firmowych stron internetowych i aplikacji. Zamiast tego uzyskają dostęp do informacji, treści i widżetów za pośrednictwem asystenta AI, który stanie się „agregatorem agregatorów”. Analitycy przytoczyli przykład lotu do Nowego Jorku i konieczności dotarcia z lotniska do biura. Co by było, gdyby osobisty agent AI mógł załatwić to wszystko za podróżującego? Nie ma potrzeby wyszukiwania. Być może nie byłoby nawet potrzeby wyjmowania smartfona. Agent AI reprezentujący podróżującego sam przeprowadzałby wszystkie działania i cała sfera marketingu i zabiegów w tworzeniu interfejsów pod kątem korzystania z nich przez ludzi okazałaby się zbędna. Ważniejsze okazałoby się konfigurowanie funkcji i usług „pod maszynę”.

Zwraca się uwagę, że kierownictwo Google mówiło o takiej sytuacji i takiej przyszłości od lat. Strony internetowe i aplikacje ostatecznie nie znikną; po prostu konsumenci nie będą odwiedzać ani widzieć tych cyfrowych lokalizacji w sposób bezpośredni. Zamiast tego uzyskają dostęp do informacji, treści i widżetów za pośrednictwem asystenta AI. Jakkolwiek to dziwnie zabrzmiało, wydaje się, że aby sprostac tej nowej rzeczywistości, firmy, marketerzy i sprzedawcy, muszą też sięgnąć po maszyny, swoich agentów, bo kto lepiej dogada się i zrobi wrażenie na automacie niż inny automat. ■

Mirosław Usidus

8. Obraz z prezentacji Agentforce usługi Salesforce

The image shows the Agentforce logo in a large, stylized font. Above it is the Salesforce logo. Below the logo, there are several circular portraits of diverse people, including a woman with dark hair, a man with a beard, and a woman with blonde hair. In the center, there is a stylized robot character. To the right, there is a screenshot of the Agentforce interface, showing a 'Steps' section with '1 Select Type' and two options: 'Agentforce Service Agent' and 'Agentforce Sales Coach'. The background is a gradient of purple and blue.

Gdy chiński DeepSeek uruchomił R1, swój, reklamowany jako niedrogi, model AI, amerykańskie firmy poczuły się zmuszone do reakcji. OpenAI odpowiedziało wersją GPT „o3-mini”, który miał łączyć „potęgę” z „niskim kosztem”. Google zaś ujawniło całkiem nowy model. W rywalizacji na modele AI jest wyraźnie nowy nurt – chodzi już nie tylko, który mądrzejszy, ale także, który tańszy.

Inteligencja nie może być dziś po prostu sztuczna – musi „rozumować” i to za niewysoką cenę

MODELE NA START



1. Identyfikacja graficzna DeepSeek R1

Zdaniem ekspertów, szum wokół premiery DeepSeek R1 był mocno przesadzony, a chiński model ma różne problemy i rodzi sporo pytań. Kwestionowana jest przede wszystkim jego „tańszość”, bo specjaliści nie bardzo wierzą, by bez użycia najbardziej zaawansowanych chipów dało się wyszkolić tak zaawansowany model. Krytykowany jest też za bezpardonową cenzurę polityczną zgodną z linią Komunistycznej Partii Chin. Patrząc bardziej specjalistycznie, R1 ma bardzo małe „okno kontekstowe” (zdolność pamiętania w konwersacji) jak na nowoczesny duży model językowy (LLM). 128 tys. tokenów DeepSeek R1 to wartość z czasów GPT-3 OpenAI. Dla obecnych przypadków użycia sztucznej inteligencji to za mało. Model Google’a, Gemini Flash 2.0, który, co się podkreśla, jest znacznie tańszy niż R1, ma okno kontekstowe o wartości miliona tokenów, a są modele z dwoma milionami w „oknie”.

Fronty konkurencji przebiegają nie tylko między Chinami a USA. Także wewnątrz Chin są inne niż Deepseek podmioty rywalizujące. Pod koniec stycznia Alibaba udostępniła model sztucznej inteligencji Qwen 2.5, który, według niej, przewyższył wersję DeepSeek-V3. Qwen 2.5-Max przewyższa podobno też prawie pod każdym względem GPT-4o OpenAI oraz Llama-3.1-405B firmy Meta. Tak przynajmniej utrzymuje Alibaba w komunikacie opublikowanym na swoim oficjalnym koncie WeChat. Chiński potentat e-commerce zapowiedział potem, iż udostępni publicznie swój model sztucznej inteligencji do generowania wideo

i obrazów o nazwie Wan 2.1. Modele te będą dostępne za pośrednictwem chmury Alibaba i Hugging Face, ogromnego repozytorium modeli AI, dla naukowców, badaczy i instytucji komercyjnych na całym świecie. DeepSeek z kolei ogłosił, że jego służący do generowania obrazów nowy model Janus-Pro-7B jest lepszy niż zachodnie Stable Diffusion i DALL-E 3. Dwa dni po premierze DeepSeek-R1, chiński właściciel TikTok, firma ByteDance, wydała aktualizację swojego modelu sztucznej inteligencji, który, według niej, osiągnął lepsze wyniki niż wspierany przez Microsoft OpenAI o1 w AIME, teście porównawczym, który mierzy, jak dobrze modele sztucznej inteligencji rozumieją i reagują na złożone instrukcje.

Specjaliści zwracają uwagę, że główną zaletą chińskich modeli w odróżnieniu od wcześniejszej generacji amerykańskiej AI jest ich otwartoźródłowość. Modele open source różnią się od modeli zastrzeżonych, takich jak te stworzone przez OpenAI tym, że nie generują bezpośrednich przychodów dla firm. Liang Wenfeng, enigmatyczny założyciel DeepSeek, powiedział w wywiadzie dla chińskich mediów w lipcu, że startup „nie dba” o wojny cenowe i że jego głównym celem jest osiągnięcie AGI (sztucznej inteligencji ogólnej). Tymczasem OpenAI myśli intensywniej o dochodach, które pokryłyby ogromne koszty firmy niż o działalności non profit. Dyrektor generalny OpenAI, Sam Altman, przedstawił szczegółowe plany ostatecznego pożegnania się przez startup z pierwotnym statusem non profit i przekształcenia go w spółkę



2. Jedna z kreatywnych wizualizacji modelu Lucie firmy Linagora

nastawioną na zysk. Jest to zresztą przedmiot sporu prawnego z Elonem Muskem, jednym z wczesnych fundatorów OpenAI, który teraz złożył ofertę o wartości 97,4 mld dolarów na zakup organizacji non profit, kontrolującej OpenAI, co pokrzyżowałoby plany Altmana.

W tym wyścigu modeli i nie tylko odstają nieco Indie, kraj znany skądinąd z mocarstwowych ambicji. Tamtejszy rząd podaje, że hinduska odpowiedź na DeepSeek nie jest daleko. Dostarcza startupom, uniwersytetom i badaczom tysięcy wysokiej klasy chipów potrzebnych do opracowania w mniej niż dziesięć miesięcy.

Jak mówił w BBC analityk Prasanto Roy, Chiny i Stany Zjednoczone mają już „cztero- lub pięcioletnią przewagę”. Zainwestowały znaczne środki w badania i projekty akademickie oraz opracowując systemy sztucznej inteligencji do zastosowań wojskowych, w wymiarze sprawiedliwości, a teraz także duże modele językowe. Do Chin i USA należy odpowiednio 60 i 20 proc. światowych patentów na rozwiązania sztucznej inteligencji złożonych w latach 2010–2022. Inne regiony świata są wyraźnie zapóźnione. Dotyczy to nie tylko Indii, ale też Europy. Właściwie jedynie we Francji, gdzie powstał względnie chwalony model Mistral, coś ciekawego się dzieje w tej dziedzinie. Jednak nawet tam niedawno francuska odpowiedź na ChatGPT, Lucie, okazała się kompromitującą porażką. Grupa Linagora, firma będąca częścią konsorcjum OpenLLM-France rozwijającego ten model, uruchomiła Lucie (2), która była podobnym do ChatGPT botem open source, promowanym jako „szczególnie

przejrzysty i niezawodny”. Niestety, odpowiedzi Lucie były zdumiewające. Na pytanie, co to jest 2+2, odpowiedziała: „Jestem zaprogramowana tak, aby być neutralną i obiektywną”. Niektórym osobom udało się przekonać Lucie do udzielenia odpowiedzi na ich zadania matematyczne, ale nie zawsze były one poprawne. Bot zaproponował również przepis na produkcję metamfetaminy i polecał krowie jaja jako źródło pożywienia. W zaledwie trzy dni po premierze Linagora Group wyłączyła Lucie, ogłaszając, że pozostaje on „akademickim projektem badawczym w początkowej fazie”.

Google stawia na narzędzia rozumujące

Dostępny jest od niedawna bez dodatkowych opłat dla użytkowników usługi Workspace narzędzie Gemini Deep Research, asystenta badawczego opartego na sztucznej inteligencji, zamiast „zwykłego: chatbota Gemini. „Rozumujący” model tworzy obszerny raport, z cytataми ze źródeł, które analizuje tworząc z nich uporządkowany pakiet. Wcześniej Deep Research był dostępny tylko w Gemini Advanced w internecie, za dodatkową opłatą 20 dolarów miesięcznie. Dla porównania OpenAI pobiera 200 dolarów miesięcznie za dostęp ChatGPT Pro do Deep Research, Inna konkurencja, Perplexity.ai, ogłosiła lepszą i szybszą wersję „modelu rozumującego” za darmo.

Oprócz Deep Research, użytkownicy Workspace zyskują również dostęp do modelu Gemini 2.0 Flash



3. Logotyp Gemini 2.0 Flash Thinking Experimental

Thinking Experimental (oraz jego wersja „Pro”), ogłoszonego w lutym (3). Te nowe funkcje Gemini zostały zaprojektowane z myślą o płynnej integracji z różnymi aplikacjami, takimi jak YouTube. Teraz, podczas oglądania wideo, użytkownicy mogą aktywować Gemini,

aby zadawać pytania dotyczące treści. Na przykład, użytkownik może zapytać o konkretny miesiąc lub technikę fitness podczas oglądania filmu instruktażowego dotyczącego ćwiczeń. Ponadto podczas przeglądania pliku PDF opcja „Zapytaj o ten plik PDF” pozwala użytkownikom uzyskać podsumowania lub wyjaśnienia, usprawniając proces wyszukiwania. Google zapowiada też, że także jego telekonferencyjna aplikacja Meet doda funkcje sztucznej inteligencji, w tym transkrypcje ze znacznikami czasu i posumowania. Gemini 2.0 Flash Thinking Experimental jest dostępny w AI Studio, platformie Google do prototypowania sztucznej inteligencji. Według opisu, to „najlepszy do multimodalnego rozumienia, rozumowania i kodowania”, z możliwością „rozumowania

Modele sztucznej inteligencji wydane w 2025 roku:

Gemini 2.5 Pro Experimental – model rozumowania, który, według Google, radzi sobie z tworzeniem aplikacji internetowych i kodu programistycznego. Jednak w porównaniu z Claude Sonnet 3.7 osiąga gorsze wyniki w jednym z popularnych testów porównawczych kodowania. Model wymaga miesięcznej subskrypcji Gemini Advanced w wysokości 20 USD.

Generator obrazów ChatGPT-4o – OpenAI zaktualizował swój istniejący model GPT-4o, aby generował obrazy, a nie tylko tekst. Podrasowany model szybko wykorzystany został w wiralnej modzie na przetwarzanie obrazów w grafiki anime w stylu Studio Ghibli.

Stable Virtual Camera – Startup zajmujący się generowaniem obrazów, Stability AI, wprowadził na rynek model, który według firmy może generować sceny 3D i kąty kamery z pojedynczego obrazu 2D. Model jest dostępny do niekomercyjnego użytku badawczego na platformie HuggingFace.

Aya Vision – Firma Cohere udostępniła multimodalne narzędzie o nazwie Aya Vision, które, według niej, jest najlepsze w swojej klasie w wykonywaniu takich czynności, jak podpisywanie obrazów i odpowiadanie na pytania dotyczące zawartości zdjęć. Jest dostępne za darmo na WhatsApp.

GPT 4.5 „Orion” – OpenAI nazywa Oriona swoim największym jak dotąd modelem, zachwalając jego silną „wiedzę o świecie” i „inteligencję emocjonalną”. Jednak w porównaniu z nowszymi modelami rozumowania osiąga gorsze wyniki w niektórych testach porównawczych. Orion jest dostępny dla subskrybentów planu OpenAI za 200 USD miesięcznie.

Claude Sonnet 3.7 – Firma Anthropic twierdzi, że jest to pierwszy w branży „hybrydowy” model rozumowania, ponieważ może zarówno udzielać szybkich odpowiedzi, jak i naprawdę przemyśleć sprawę, gdy jest to potrzebne.

Grok 3 – Najnowszy model założonego przez Elona Muska startupu xAI. Twierdzi się, że przewyższa inne wiodące modele w matematyce, naukach ścisłych i kodowaniu. Model ten wymaga subskrypcji X Premium (który kosztuje 50 USD miesięcznie).

OpenAI o3-mini – Najnowszy model rozumujący OpenAI, zoptymalizowany pod kątem zadań naukowych i edukacyjnych, kodowania, matematyki i nauk ścisłych. Nie jest to najpotężniejszy model OpenAI, ale ponieważ jest „mini” firma twierdzi, że jest znacznie tańszy.

OpenAI Deep Research – Przeznaczony do przeprowadzania dogłębnych badań na dany temat z wyraźnymi cytatami. Ta usługa jest dostępna tylko z subskrypcją ChatGPT Pro za 200 USD miesięcznie.

Mistral Le Chat – Francuski Mistral uruchomił wersję aplikacji Le Chat, multimodalnego osobistego asystenta AI. Twierdzi, że reaguje szybciej niż jakikolwiek inny chatbot. Ma również płatną wersję z dostępem do treści agencji AFP. Testy przeprowadzone przez „Le Monde” wykazały, że wydajność Le Chat jest dobra, chociaż popełnił więcej błędów niż ChatGPT.

Operator OpenAI – Agent AI, który ma być osobistym „stażystą”, który może robić rzeczy niezależnie, na przykład pomagać w zakupach. Wymaga subskrypcji ChatGPT Pro w wysokości 200 USD miesięcznie. Recenzent Washington Post twierdzi, że Operator sam zdecydował się zamówić tuzin jajek za 31 dolarów, płacąc kartą kredytową recenzenta.

Google Gemini 2.0 Pro Experimental – Według firmy, doskonale radzi sobie z kodowaniem i rozumieniem wiedzy ogólnej. Ma również bardzo długie okno kontekstowe z dwoma milionami tokenów, co jest ważne dla użytkowników, którzy muszą szybko przetwarzać ogromne partie tekstu. Usługa wymaga (co najmniej) subskrypcji Google One AI Premium w wysokości 19,99 USD miesięcznie.

nad najbardziej złożonymi problemami” w dziedzinach takich jak programowanie, matematyka i fizyka. Gemini 2.0 Flash Thinking Experimental wydaje się podobny w konstrukcji do o1 OpenAI i innych tak zwanych „modeli rozumujących”. Narzędzia tego typu wyróżniają się tym, że sprawdzają fakty, co pomaga im unikać niektórych pułapek, w które wpadają często modele sztucznej inteligencji. Ich wadą jest, że często potrzebują sporo czasu, nawet do kilku minut, na znalezienie rozwiązania. W odpowiedzi na podpowiedź, Gemini 2.0 Flash Thinking Experimental zatrzymuje się na namysł, rozważając szereg powiązanych podpowiedzi i „wyjaśniając” swoje rozumowanie po drodze. Po pewnym czasie model referuje to, co uważa za najtrafniejszą odpowiedź.

Google ma cały wachlarz ofert AI, które regularnie rozwija i dodaje nowe możliwości, np. Gemini 2.0 Flash od marca 2025 r. pozwala na manipulację zdjęciami. Nie chodzi tu tylko o generowanie nowych obrazów od zera. Google stworzyło system, który ma rozumieć istniejące zdjęcia na tyle dobrze, by modyfikować je w procesie konwersacji z chatbotem, zachowując przy tym znaczną część oryginalnej zawartości przy jednoczesnym wprowadzaniu określonych zmian, co oznacza, że jest to narzędzie multimodalne, które może jednocześnie rozumieć zarówno tekst, jak i obrazy. Model konwertuje obrazy na tokeny, te same podstawowe jednostki, których używa do przetwarzania tekstu, umożliwiając mu manipulowanie treściami wizualnymi przy użyciu tych samych ścieżek neuronowych, których używa do rozumienia języka. Rozwój tych funkcji wchodzi w zakres wcześniej anonsowanego Projektu Astra Google (4), multimodalnej inicjatywy asystenta AI. Projekt Astra ma na celu stworzenie asystenta, który będzie postrzegał i rozumiał swoje otoczenie, ułatwiając bardziej naturalne interakcje.

Całościowe podejście Google różni się od konkurencji, głównie OpenAI, której ChatGPT może generować obrazy za pomocą Dall-E 3 i edytować swoje kreacje, rozumiejąc język naturalny, ale wymaga do tego osobnego modelu sztucznej inteligencji. Innymi słowy, ChatGPT wymaga koordynacji między GPT-V dla wideo, GPT-4o dla języka i Dall-E 3 dla generowania obrazów, zamiast mieć jeden model do wszystkiego. Coś takiego OpenAI spodziewa się osiągnąć dopiero dzięki GPT-5, choć są wątpliwości, czy w ogóle się tego modelu doczekamy.

Google stara się również dotrzymać kroku innym nurtom rozwoju AI. Jego oddział DeepMind pracuje nad dwoma nowymi modelami sztucznej inteligencji specjalnie dla robotów. Pierwszy z nich, o nazwie Gemini Robotics, to model wizyjno-językowo-działaniowy zdolny do rozumienia nowych sytuacji, nawet jeśli nie został w nich przeszkolony. Gemini Robotics opiera się na Gemini 2.0. Oprócz zdolności do generalizowania nowych scenariuszy, Gemini Robotics jest lepszy w interakcji z ludźmi i ich otoczeniem. Jest również w stanie wykonywać relatywnie precyzyjne zadania fizyczne, takie jak składanie kartki lub zdejmowanie zakrętki z butelki. Drugi model to Gemini Robotics-ER, który firma opisuje jako zaawansowany model języka wizualnego, potrafiący „zrozumieć nasz złożony i dynamiczny świat”. Został zaprojektowany dla robotyków, by łączyć się z istniejącymi kontrolerami niskiego poziomu, systemami, które kontrolują ruchy robota, umożliwiając im włączenie nowych możliwości zasilanych przez Gemini Robotics-ER. Google DeepMind razem z firmą Appteronik pracuje nad „zbudowaniem następnej generacji humanoidalnych robotów”. Dostęp do modelu Gemini Robotics-ER, otrzymały najbardziej znane firmy robotyczne, w tym Agile Robots, Agility Robotics, Boston Dynamics i Enchanted Tools.

4. Prezentacja projektu Astra



Duże koszty, a efekt niewiele lepszy

Fala „modeli rozumujących” czy też „wnioskujących” wezbrała po wydaniu przez OpenAI modelu o skromnej nazwie o1. Proponuje je nie tylko Google. Proces rozumowania wykładają w odpowiedzi także wspomniane chińskie DeepSeek-R1 i Qwen firmy Alibaba. Nie wszyscy jednak podzielają przekonanie, że „modele rozumujące” są najlepszą drogą rozwoju AI. Po pierwsze, są one zwykle drogie ze względu na dużą moc obliczeniową wymaganą do ich uruchomienia. I choć do tej pory osiągały dobre wyniki w testach porównawczych, nie jest jasne, czy mają perspektywy rozwoju swoich możliwości, zwłaszcza że generatywna AI zdaje się docierać do granic swoich możliwości, o czym ma świadczyć przykład uznawanego za co najmniej nierobiący wrażenia GPT-4.5 firmy OpenAI.

Według recenzji, których nie brakuje w mediach internetowych, najnowszy i najbardziej wydajny jak do tej pory model sztucznej inteligencji OpenAI, GPT-4.5 (5), jest za duży, za drogi i za powolny, zapewniając jedynie trochę lepsze wyniki niż GPT-4o przy trzydziestokrotnie wyższym koszcie danych wejściowych i piętnastokrotnie wyższym koszcie danych wyjściowych. Nowy model wydaje się dowodzić, że obiegowe opinie o malejących zwrotach w szkoleniu LLM były słuszne i że tak zwane „prawo skalowania” już nie działa. Częsty krytyk OpenAI, Gary Marcus, nazwał w poście na blogu to wydanie GPT „nic niewartym burgerem”. Były badacz OpenAI Andrej Karpathy napisał na X oględniej, że GPT-4.5 jest lepszy niż GPT-4o, ale w sposób, który jest subtelny i trudny do wyrażenia.

Według własnych wyników testów porównawczych OpenAI, GPT-4.5 uzyskał w testach takich jak konkursy matematyczne AIME i oceny naukowe GPQA znacznie niższe wyniki niż wcześniej udostępnione „modele rozumujące” OpenAI, o1 i o3. GPT-4.5 uzyskał zaledwie 36,7 proc. w teście AIME w porównaniu do 87,3 proc. uzyskanym przez o3-mini. Dodatkowo, GPT-4.5, jeśli chodzi o przetwarzanie danych wejściowych, kosztuje pięć razy więcej niż o1 i ponad 68 razy więcej niż o3-mini. Inwestor technologiczny Paul Gauthier przeprowadził niezależne testy umiejętności pisania kodu przez GPT-4.5 za pomocą testu porównawczego Aider's Polyglot Coding i podał, że GPT-4.5 zajął dziesiąte miejsce pod względem ogólnej zdolności (po rozszerzonym o zdolności rozumowania Claude 3.7 Sonnet i o1 i o3). GPT-4.5 uzyskał wyższe wyniki w wielojęzycznym teście MMLU (wiedza ogólna) – 85,1 proc. w porównaniu do 81,5 proc. GPT-4o. OpenAI twierdzi również, że GPT-4.5 wykazał lepsze wyniki w zmniejszaniu halucynacji, co oznaczałoby, że generuje mniej fałszywych lub wprowadzających w błąd odpowiedzi niż poprzednie wersje. Jest też

pomyślny, wyprzedzający ludzi wynik testu Turinga GPT-4.5, o czym piszemy w innym miejscu w tym wydaniu MT. Jednak i tak jest powszechnie krytykowany jako zbyt drogi (dla użytkownika), jak na wyniki, które może zaoferować

Być może z powodu rozczarowujących wyników, szef OpenAI, Sam Altman, napisał, że GPT-4.5 będzie ostatnim z tradycyjnych modeli sztucznej inteligencji jego firmy, a GPT-5 ma być dynamicznym połączeniem LLM i symulowanych modeli rozumowania, takich jak o3. OpenAI prawdopodobnie wiedziała wcześniej o malejących zyskach ze szkolenia LLM. W rezultacie firma spędziła większość ubiegłego roku, pracując nad symulowanymi modelami rozumującymi, takimi jak o1 i o3, które wykorzystują inne podejście do wnioskowania w celu poprawy wydajności, zamiast rzucać coraz większe ilości danych szkoleniowych na modele AI w stylu GPT.

Hitem internetu za to stał się inny nowy produkt OpenAI – generator obrazów GPT-4o dostępny przez usługę ChatGPT. W internecie można znaleźć niezliczone przykłady obrazów generowanych przez sztuczną inteligencję za pomocą jednego polecenia w ChatGPT, a wyniki uznano za dużo lepsze niż znane z generatorów dotychczas (6). Największe wrażenie robiła możliwość generowania tekstu pisanego na obrazach błędów. Znany z tworzonych przez AI treści, bełkot był wcześniej był znakiem rozpoznawczym treści generowanych przez sztuczną inteligencję. GPT-4o to zmienia.

Czwarta generacja lam i zaskakująco dobry model Muska

Firma będąca symbolem rewolucji generatywnej sztucznej inteligencji działa, jak widać, ze zmiennym szczęściem. Inne firmy też się uaktywniły. W kwietniu 2025 r. Microsoft, który znany był raczej ze współpracy, a nie z rywalizacji z OpenAI, ogłosił całą serię nowych produktów, w tym pakiet agentów AI o nazwie Actions, Memory, Vision, Pages, Shopping i Copilot Search. Kilka z nich pojawiło się już w konkurencyjnych produktach



5. Symbol modelu GPT-4.5 i szef firmy OpenAI Sam Altman



6. Przeróbki znanych obrazów w stylu japońskich kreskówek Ghibli wykonane w GPT-4o

AI, takich jak Google Gemini i OpenAI ChatGPT, oraz w produktach Perplexity i Opera.

W Copilocie Microsoftu brakowało rozwiązania „rozumującego”, „wnioskującego” lub „deep research”, ale to się ma zmienić. „Copilot może wyszukiwać, analizować i łączyć informacje ze źródeł online lub dużych ilości dokumentów i obrazów”, głosi w kwietniowym komunikacie Microsoft. Do uruchomienia zapytania tego rodzaju nie jest nawet potrzebne konto Microsoft, a subskrypcja Copilot Pro nie jest obowiązkowa. Microsoft zapewni pięć darmowych zapytań typu „deep research” każdego miesiąca, zaś subskrybenci otrzymają nieograniczoną liczbę prób i priorytetowy dostęp. Pod koniec marca platforma Microsoft 365 Copilot uzyskała dostęp do narzędzia AI Researcher, które może analizować źródła online, a także pliki lokalne. Dostępna jest też usługa Copilot Search, która jest wyszukiwarką opartą na Bing firmy Microsoft.

Także w kwietniu Meta (Facebook) zademonstrowała nową kolekcję modeli AI, Llama 4 – Llama 4 Scout, Llama 4 Maverick i Llama 4 Behemoth (7). Wszystkie zostały przeszkolone na „dużych ilościach nieoznakowanych danych tekstowych, graficznych i wideo”, co ma zapewnić im „szerokie zrozumienie wizualne”, głosi Meta. Mówi się, że Meta badała, w jaki sposób DeepSeek obniżył koszty uruchamiania i wdrażania modeli takich jak R1 i V3. Meta twierdzi, że Meta AI, jej asystent oparty na sztucznej inteligencji w aplikacjach, w tym WhatsApp, Messenger i Instagram, został zaktualizowany do korzystania z Llama 4 w czterdziestu krajach. Funkcje multimodalne są na razie ograniczone

do Stanów Zjednoczonych w języku angielskim. Jak podaje firma Marka Zuckerberga, Llama 4 jest pierwszym pakietem modeli wykorzystujących architekturę „mieszanki ekspertów” (MoE). Architektury MoE dzielą zadania przetwarzania danych na podzadania, a następnie delegują je do mniejszych, wyspecjalizowanych modeli „eksperckich”. Na przykład Maverick ma 400 miliardów parametrów, ale tylko 17 miliardów aktywnych parametrów wśród 128 „eksperckich”. Scout ma 17 miliardów aktywnych parametrów, 16 ekspertów i 109 miliardów całkowitych parametrów. Według wewnętrznych testów Meta, Maverick, według firmy najlepiej nadający się do przypadków użycia „ogólnego asystenta i czatu”, przewyższa modele takie jak GPT-4o OpenAI i Gemini 2.0 Google w niektórych testach porównawczych dotyczących kodowania, rozumowania, wielojęzyczności, długiego kontekstu i obrazu. Niewydany jeszcze Behemoth firmy Meta będzie wymagał jeszcze mocniejszego sprzętu. Według firmy, Behemoth



7. Obrazek z prezentacji rodziny modeli Llama 4

ma 288 miliardów aktywnych parametrów, 16 ekspertów i prawie 2 biliony parametrów w ogóle. Wewnętrzne testy porównawcze Meta wykazały, że Behemoth przewyższa GPT-4.5, Claude 3.7 Sonnet i Gemini 2.0 Pro (ale nie 2.5 Pro) w testach mierzących umiejętności STEM, takie jak rozwiązywanie problemów matematycznych.

Warto zauważyć, że żaden z modeli Llama 4 nie jest definicyjnie modelem „rozumującym” na wzór o1 i o3-mini OpenAI. Co ciekawe, Meta twierdzi, że dostroiła wszystkie swoje modele Llama 4, aby rzadziej odmawiały odpowiedzi na „kontrowersyjne” pytania. Według firmy, Llama 4 reaguje na „dyskutowane” tematy polityczne i społeczne, których poprzednia seria modeli Llama by nie robiła.

Meta planuje też wydać samodzielną aplikację Meta AI w drugim kwartale i podobno zamierza przetestować płatną wersję subskrypcji aplikacji, podobną do istniejących ofert rywali, takich jak OpenAI i xAI Elona Muska. Inni pobierają opłaty za wersje premium. Do tej pory Meta AI była dostępna za pośrednictwem przeglądarki internetowej i istniejących aplikacji Meta na Facebooku, Instagramie, WhatsApp i Messengerze.

W lutym xAI, producent modeli AI kierowany przez Elona Muska, zaprezentował swoją najnowszą rodzinę modeli, Grok 3 (8). Trójka pokonała DeepSeek-R1 w rozumowaniu. Zgodnie z testami porównawczymi, Grok 3 przewyższa kilka konkurencyjnych modeli i jest również pierwszym, który uzyskał ponad 1400 punktów na Chatbot Arena, platformie do porównywania i oceny modeli AI. Grok 3 oferuje również możliwości rozumowania i funkcję głębokich badań o nazwie DeepSearch. Andrej Karpathy, założyciel Eureka Labs,

wcześniej w OpenAI, otrzymał wczesny dostęp do Grok 3. W poście na X podał m.in., że model dobrze radził sobie ze złożonymi zadaniami, takimi jak tworzenie siatki heksów dla popularnej gry planszowej Settlers of Catan. „Niewiele modeli radzi sobie z tym niezawodnie. Najlepsze modele rozumujące OpenAI (np. o1-pro, w cenie za 200 USD za miesiąc) również to robią, ale DeepSeek-R1, Gemini 2.0 Flash Thinking czy Claude nie potrafią”, napisał.

Uff, można poczuć się nieco przytłoczonym nawałą nowych ofert, funkcji i narzędzi sztucznej inteligencji. Według Gartnera, firmy zajmujące się badaniami rynku technologicznego, do przyszłego roku na całym świecie zostanie wydanych więcej pieniędzy na oprogramowanie z generatywną AI niż na oprogramowanie bez niej. W tym roku niemal każda duża firma zajmująca się oprogramowaniem będzie miała funkcję generatywnej AI we wszystkich swoich produktach lub będzie planować wprowadzenie jej w tym lub przyszłym roku, zauważył John Lovelock, analityk w Gartner, w komunikacie. W kolejnych latach należy spodziewać się szybkiego wzrostu

Wdrożenie AI stwarza zarówno możliwości, jak i ryzyko dla firm. Firmy, które najszybciej wykorzystają AI, mogą zdobyć przewagę nad konkurencją. Jednak, jeśli to zrobią bez przemyślanego planu, mogą wiele stracić na błędach, marnowaniu czasu i z powodu zagrożeń bezpieczeństwa. Podobne dylematy jak firmy mają ich pracownicy i każdy zwykły człowiek. Z jednej strony duża pomoc, oszczędność czasu i pracy, z drugiej – ryzyko błędów, strat i kradzieży cennych danych. Entuzjazmowi dla AI zawsze muszą towarzyszyć rozsądek i przeczność. ■

Miroslaw Usidus

8. Elon Musk prezentuje model Grok 3



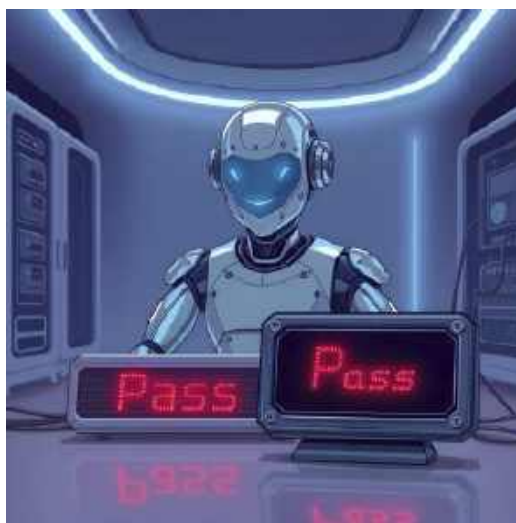
Słynny test Turinga został, jak wiadomo, „pokonany” przez generatywną sztuczną inteligencję. Jest jednak sporo kontrowersji, czy przypadkiem najnowsze modele nie są konstruowane specjalnie tak, by pokonywać czy wręcz oszukiwać testy (1). Bo maszyna specjalizująca się w rozwiązywaniu testów wykazujących, że nie różni się lub nawet przewyższa człowieka, nie jest jak człowiek.

Osobliwość u bram – a może wcale nie?

PRZYSZŁOŚĆ BARDZIEJ OGÓLNA NIŻ INTELIGENTNA

Jak uważają przedstawiciele Google DeepMind w publikacji dla „MIT Press”, dane wykorzystywane do trenowania sztucznej inteligencji są zbyt ograniczone i statyczne i nigdy nie doprowadzą sztucznej inteligencji do nowych i lepszych umiejętności, nie uczynią z niej ogólnej sztucznej inteligencji i nie doprowadzą do osobliwości technologicznej, w której AI przewyższy człowieka. Dlatego David Silver i Richard Sutton z DeepMind w artykule „Welcome to the Era of Experience” proponują, by sztuczna inteligencja mogła mieć pewnego rodzaju „doświadczenia”, wchodząc w interakcje ze światem, formułować cele w oparciu o sygnały ze środowiska. Obaj współtworzyli AlphaZero, model sztucznej inteligencji, który pokonał ludzi w grach w szachy i Go. Podejście, za którym opowiadają się dwaj naukowcy, opiera się na uczeniu się ze wzmocnieniem i doświadczeniach z pracy nad AlphaZero. Nazywają to „strumieniami”, które miałyby zaradzić niedociągnięciom dzisiejszych dużych modeli językowych (LLM), opracowywanych dotychczas wyłącznie w celu udzielenia odpowiedzi na pytania.

Narzędzia generatywnej sztucznej inteligencji, takie jak ChatGPT, nie skorzystały ze starszych metod nauki ze wzmocnieniem. To posunięcie miało swoje zalety i wady. Modele Gen AI mogą obsługiwać spontaniczne dane wejściowe od ludzi, z którymi nigdy wcześniej nie miały do czynienia, bez wyraźnych reguł dyktujących przebieg procedury. Jednak odrzucenie uczenia ze wzmocnieniem oznaczało,



1. Maszyna ‘pokonująca’ testy

że „coś zostało utracone – zdolność agenta do samodzielnego odkrywania własnej wiedzy”, piszą autorzy. Zauważają, że LLM bazują „na ludzkich uprzedzeniach” lub na tym, czego chce człowiek na etapie odpowiedzi. Takie podejście jest zbyt ograniczone. Sugerują, że ludzka ocena „nakłada nieprzekraczalny pułap na wydajność agenta – agent nie może odkryć lepszych strategii niedocenianych przez człowieka”.

Fundamentalne ograniczenia

Krytyka generatywnych technik AI rozszerza się. W ankiecie z marca 2025 r. 76 proc. naukowców stwierdziło, że jest „mało prawdopodobne” lub „bardzo mało prawdopodobne”, aby skalowanie dużych modeli językowych pozwoliło osiągnąć „sztuczną inteligencję ogólną” (AGI). Większość badaczy ankietowanych przez Association for the Advancement of Artificial Intelligence uważa, że firmy technologiczne znalazły się w ślepych zaułku, a pieniądze ich z niego nie wyciągną. „Myślę, że od czasu wydania GPT-4 było oczywiste, że zyski ze skalowania były stopniowe i kosztowne”, powiedział serwisowi „Live Science” Stuart Russell, informatyk z Uniwersytetu Kalifornijskiego w Berkeley,



2. Jedna z wizualizacji AGI © AI

który pomógł przygotować raport. „Zainwestowano już zbyt wiele i nie można sobie pozwolić na przyznanie się do błędu i wypadnięcie z rynku na kilka lat, plus konieczność spłacenia inwestorów, którzy włożyli setki miliardów dolarów”.

Emerytowany profesor psychologii i neurologii na Uniwersytecie Nowojorskim i popularny komentator AI, Gary Marcus, opublikował na początku 2025 r. prognozę, że sztuczna inteligencja ogólna (2) nie pojawi się w tym roku, podobnie jak wyczekiwany model GPT-5 firmy OpenAI. Jak dodał, OpenAI wydaje się tkwić w martwym punkcie, co pokazuje zarazem szersze zmagania ze skalowaniem sztucznej inteligencji. Marcus przypisuje te opóźnienia wąskiemu gardłom w postaci braku danych i ograniczeniom w naturze dużych modeli językowych (LLM). Podkreśla, że samo skalowanie danych nie jest już wystarczające. Zwraca też uwagę trochę żartem na zmianę w retoryce Elona Muska na temat sztucznej inteligencji. Według buńczucznych prognoz szefa Tesli AI miała przewyższyć ludzką inteligencję do 2025 r., Musk teraz znacznie ostrożniej sugeruje, że sztuczna inteligencja może „wykonać dowolne zadanie poznawcze w ciągu najbliższych kilku lat”.

Postęp LLM w ostatnich latach jest częściowo zasługą korzystania z architektury transformatorowej. Jest to rodzaj architektury głębokiej nauki maszynowej, stworzonej po raz pierwszy w 2017 roku i opatentowanej przez naukowców Google, która rozwija się i uczy poprzez absorbowanie danych szkoleniowych z danych wejściowych pochodzących od ludzkich użytkowników. Umożliwia to modelom generowanie probabilistycznych (opartych na prawdopodobieństwie) wzorców z ich sieci neuronowych (zbiorów algorytmów uczenia maszynowego ułożonych tak, aby

naśladować sposób, w jaki uczy się ludzki mózg) przez przekazywanie ich po otrzymaniu monitu, a ich odpowiedzi poprawiają się pod względem dokładności wraz z większą ilością danych. Jednak dalsze skalowanie tych modeli wymaga ogromnych nakładów kosztów i energii. Tylko w 2024 r. branża generatywnej sztucznej inteligencji pozyskała 56 mld dolarów kapitału wysokiego ryzyka na całym świecie, z czego znaczna część została przeznaczona na budowę ogromnych kompleksów centrów danych. Prognozy pokazują, że zasoby danych generowanych przez człowieka, niezbędne do dalszego wzrostu generatywnych modeli, najprawdopodobniej zostaną wyczerpane do końca tej dekady. Gdy to nastąpi, alternatywą będzie rozpoczęcie zbierania prywatnych danych od użytkowników lub wprowadzanie generowanych przez sztuczną inteligencję „syntetycznych” danych, co może narazić je na ryzyko upadku z powodu błędów powstałych po połknięciu ich własnych danych wejściowych.

Ekspertci twierdzą jednak, że ograniczenia obecnych modeli wynikają nie tylko z faktu, że są one głodne zasobów, ale także z fundamentalnych ograniczeń w ich architekturze. „Myślę, że podstawowy problem z obecnymi podejściami polega na tym, że wszystkie one obejmują szkolenie dużych obwodów sprzężenia zwrotnego”, powiedział Russell. „Obwody mają fundamentalne ograniczenia jako sposób reprezentowania pojęć”.

AI mądrzejsza niż nobliści w swoich dziedzinach

Przez AGI (Artificial General Intelligence, sztuczna inteligencja ogólna) rozumie się zwykle taki rodzaj AI, która dorównuje lub nawet przewyższa ludzkie zdolności poznawcze. Hipotetycznie potrafi zrozumieć i zastosować wiedzę w szerokim zakresie zadań, podobnie jak człowiek, a nie w ograniczonym zestawie przypadków użycia. Niektórzy badacze, którzy studiowali pojawienie się inteligencji maszynowej, uważają, że pojawienie się AGI to tzw. osobiliwość, teoretyczny punkt, w którym maszyna przewyższa człowieka pod względem inteligencji. Choć pojęcie to jest też szerzej interpretowane jako sytuacja, w której tracimy jako ludzie zdolność zrozumienia tego, co się dzieje.

Dyrektor generalny Anthropic Dario Amodei (3) przewidywał w wywiadzie dla „Wall Street Journal”, że modele AI mogą przewyższyć ludzkie możliwości „w prawie wszystkim” w ciągu dwóch do trzech lat. Przemawiając kilka miesięcy temu w Davos, Amodei powiedział: „Nie wiem dokładnie, kiedy to nastąpi, nie wiem, czy będzie to rok 2027. Myślę,

że jest prawdopodobne, że może to potrwać dłużej. Nie sądzę, że będzie to o wiele dłużej niż wtedy, gdy systemy AI będą lepsze od ludzi w prawie wszystkim. Lepsze niż prawie wszyscy ludzie w prawie wszystkim. A w końcu będą lepsze od wszystkich ludzi we wszystkim, nawet w robotyce”. Amodei mówił również o potencjalnych implikacjach wysoce inteligentnych systemów AI, gdy te modele AI mogą kontrolować zaawansowaną robotykę. Amodei jednak dystansuje się od preferowanego przez Sama Altmana z OpenAI terminu „sztuczna inteligencja ogólna” (AGI), określając go jako „termin marketingowy”. Zamiast tego woli opisywać przyszłe systemy sztucznej inteligencji jako „kraj geniuszy w centrum danych”. Amodei napisał w esej z października 2024 r., że takie systemy musiałyby być „mądrzejsze niż laureaci Nagrody Nobla w swoich dziedzinach”.

Prognozy dotyczące „osobliwości” i nadejścia AGI lub też „sztucznej superinteligencji” pojawiają się od lat, a nawet dekad. Większość przepowiedni, starszych czy nowych, zgadza się, że coś, co możemy dla uproszczenia nazywać AGI, pojawi się przed końcem XXI wieku. Nowe badania obejmujące prognozy ponad 8,5 tysiąca uczonych zostały przeprowadzone przez firmę badawczą o nazwie AIMultiple. „Obecne badania naukowców zajmujących się sztuczną inteligencją przewidują AGI około 2040 roku”, czytamy w raporcie z tych badań. Jednak zaledwie kilka lat przed szybkim postępem w dziedzinie dużych modeli językowych naukowcy przewidywali, że nastąpi to około 2060 roku. Przedsiębiorcy są jeszcze bardziej optymistyczni, przewidując, że nastąpi to około 2030 roku.

Jednak nie wszyscy uważają, że AGI, osobliwość czy jakkolwiek nazwiemy tę sytuację, jest pewnikiem. Zdaniem wielu oczonych, ludzka inteligencja jest bardziej wieloaspektowa, niż opisuje to obecna definicja AGI. Na przykład, niektórzy badacze opisują ludzki

umysł w kategoriach ośmiu inteligencji, z których „logiczno-matematyczna” jest tylko jednym z rodzajów (obok niej istnieje na przykład inteligencja interpersonalna, intrapersonalna i egzystencjalna). Pionier głębokiego uczenia Yann LeCun uważa, że AGI powinna zostać przemianowana na „zaawansowaną inteligencję maszynową” i twierdzi, że ludzka inteligencja jest zbyt specyficzna, aby można ją było powielić.

W świecie sztucznej inteligencji idea „osobliwości” jest bardzo popularna. Ta niezbyt jednoznaczna koncepcja opisuje to jako moment, w którym sztuczna inteligencja wykracza poza ludzką kontrolę i szybko przekształca społeczeństwo. Trudną kwestią jest określenie, gdzie się zaczyna owa „osobliwość” i prawie niemożliwe jest ustalenie, co znajduje się poza tym technologicznym „horyzontem zdarzeń”.

Jednak niektórzy próbują określić jakieś wyznaczniki tego mglistego stanu. Próbę zdefiniowania podjęła niedawno Translated, firma tłumaczeniowa. Jest to jej zdaniem zdolność sztucznej inteligencji do tłumaczenia mowy z dokładnością taką jak człowiek. „Język jest dla ludzi najbardziej naturalną rzeczą”, powiedział dyrektor generalny Translated Marco Trombetti na konferencji w Orlando na Florydzie w grudniu 2022 roku. „Niemniej dane zebrane przez Translated wyraźnie pokazują, że maszyny nie są tak daleko od wypełnienia luki”. Firma śledziła wydajność swojej sztucznej inteligencji od 2014 do 2022 roku przy użyciu wskaźnika zwanego „Time to Edit” (TTE), który oblicza czas potrzebny profesjonalnym redaktorom na poprawienie tłumaczeń wygenerowanych przez sztuczną inteligencję w porównaniu z tłumaczeniami ludzkimi. AI wykazała powolną, ale niezaprzeczalną poprawę, ponieważ powoli zmniejszała lukę w kierunku jakości tłumaczeń na poziomie ludzkim. Jeśli trend ten się utrzyma, do końca dekady (lub nawet wcześniej) sztuczna inteligencja firmy Translated będzie tak dobra, jak tłumaczenia sporządzane przez ludzi. „To pierwszy raz, kiedy ktoś w dziedzinie sztucznej inteligencji przewidział prędkość zmiernia do osobliwości”, powiedział Trombetti. Tylko że w tym przypadku chodzi tylko o tłumaczenia. Gdzie „ogólnosc” sztucznej inteligencji?

Jeśli AGI ma naśladować sposób działania mózgu, to w sensie fizycznym i technicznym rozwój technik sztucznej inteligencji nie jest nawet blisko pierwowzoru (4). Większość obecnych systemów AI, w tym wszystkie duże modele językowe, opiera się na sztucznych sieciach neuronowych. Zostały one zaprojektowane tak, aby naśladować sposób działania mózgu, wykorzystując duże liczby sztucznych neuronów pobierających dane wejściowe, modyfikujących

3. Dario Amodei





4. Porównanie sztucznej i biologicznej sieci neuronowej

je, a następnie przekazujących zmodyfikowane informacje do innej warstwy sztucznych neuronów. Po przejściu przez wystarczającą liczbę warstw, ostatnia warstwa jest odczytywana i przekształcana w dane wyjściowe. Jednak wszystkie sztuczne neurony są funkcjonalnie równoważne, nie ma specjalizacji. Zaś biologiczne neurony są wysoce wyspecjalizowane, wykorzystują różne neuroprzekaźniki i pobierają dane wejściowe z szeregu czynników pozanerwowych, takich jak hormony. Ponadto, zamiast po prostu przekazywać pojedynczą wartość do następnej warstwy, prawdziwe neurony komunikują się przez analogową serię skoków aktywności, wysyłając ciągi impulsów, które różnią się czasem i intensywnością. Różnice między sieciami neuronowymi a rzeczywistymi mózgami, na których były one wzorowane, wykraczają daleko poza różnice funkcjonalne. Obecne AI mają zazwyczaj dwa stany: szkolenie i wdrożenie. Mózg nie ma odrębnych stanów uczenia się i aktywności; jest stale w obu trybach. W wielu przypadkach mózg uczy się podczas działania. Sztuczna inteligencja zazwyczaj nie jest zbyt użyteczna, dopóki nie przejdzie znacznej liczby szkoleń. W przeciwieństwie do tego, człowiek może często zdobyć podstawowe kompetencje w bardzo krótkim czasie. W przeciwieństwie do sztucznej inteligencji, wydajność człowieka nie pozostaje statyczna. Przyrostowe ulepszenia i innowacyjne podejścia są nieustannie możliwe. Pozwala to również ludziom łatwiej dostosować się do zmieniających się okoliczności. Dla wielu rodzajów sztucznej inteligencji „pamięć” jest nie do odróżnienia od zasobów obliczeniowych, które pozwalają jej wykonywać

zadania i połączenia utworzone podczas treningu. W przypadku dużych modeli językowych obejmuje ona zarówno wagi wyuczonych wówczas połączeń, jak i wąskie „okno kontekstowe”, które obejmuje wszelkie niedawne wymiany z jednym użytkownikiem. W przeciwieństwie do tego systemy biologiczne mają całe „życie” pełne wspomnień, na których mogą polegać. A to może być kluczem do uogólnienia inteligencji. Pomaga nam rozpoznać możliwości i ograniczenia rysowania analogii między różnymi okolicznościami lub stosowania rzeczy wyuczonych w jednym kontekście w porównaniu z innym. Dostarcza nam spostrzeżeń, które pozwalają nam rozwiązywać problemy, z którymi nigdy wcześniej nie mieliśmy do czynienia. Mózgi ewoluowały w warunkach ogromnych ograniczeń energetycznych i nadal działają, zużywając znacznie mniej energii, niż może zapewnić spożywany przez człowieka pokarm. W przeciwieństwie do tego, historia ostatnich osiągnięć w dziedzinie sztucznej inteligencji to w dużej mierze historia rzucania w nią większej ilości zasobów. Chyba dość wyraźnie widać, że nie tędy droga.

Czy osobliwość da ludziom więcej pracy?

Osobliwość jako taka może być wciąż od nas odległa, ale przedstawiciele wielu zawodów odczuwają ją osobliwie na własnej skórze jako lęk przed zastąpieniem przez AI. Ze względu na halucynacje i błędy, do tej pory było to problematyczne dla czegośkolwiek innego niż zadania i stanowiska mniej istotne, ale może się to wkrótce zmienić. Pojawiają się doniesienia, że sztuczna inteligencja zbliża się nawet do pracy naukowców, a wkrótce mogą pojawić się agenci AI na poziomie naukowym doktora. Miał ją podobno zaprezentować Altman na zamkniętym spotkaniu z przedstawicielami rządu Stanów Zjednoczonych w styczniu 2025 roku.

Ponad jedną czwartą nowego kodu Google jest generowana przez sztuczną inteligencję (AI), ujawnił jeszcze w ub. roku dyrektor generalny Sundar Pichai. „Używamy sztucznej inteligencji wewnętrznie, aby usprawnić nasze procesy kodowania, co zwiększa produktywność i wydajność”, powiedział Pichai. Jak dodał, kod jest następnie „sprawdzany i akceptowany przez inżynierów. Pomaga to im robić więcej i działać szybciej”. Cytowany wyżej Dario Amodei twierdzi, że sztuczna inteligencja będzie pisać 90 proc. kodu jeszcze w 2025 r., automatyzując tworzenie oprogramowania w ciągu roku. Uważa on, że wkrótce cały kod będzie generowany przez sztuczną inteligencję. Inżynierowie oprogramowania wskazywali, że choć

modele sztucznej inteligencji są bardzo obiecujące, nie mogą one wyjść poza podstawowe poziomy kodowania. Pojawienie się modeli sztucznej inteligencji, takich jak model rozumujący o1 OpenAI, oferujących zaawansowane możliwości w zakresie matematyki, nauk ścisłych i kodowania, powinno jednak wzbudzać ich niepokój. Zgodnie z udostępnionymi testami porównawczymi, OpenAI o1 i o1-mini są świetne w kodowaniu i przeszły rozmowę kwalifikacyjną z inżynierem badawczym OpenAI w zakresie kodowania na poziomie 90–100 proc. Dyrektor generalny NVIDIA, Jensen Huang, najwyraźniej skłania się do podobnych poglądów, twierdząc, że kodowanie jako ścieżka kariery może być już ślepym zaułkiem w związku z szybkim rozpowszechnieniem sztucznej inteligencji. Zamiast tego zalecił młodym ludziom biologię, edukację, produkcję lub rolnictwo jako prawdopodobne i bezpieczniejsze alternatywne opcje kariery dla następnego pokolenia.

Mówiąc o przejęciu miejsc pracy przez sztuczną inteligencję, była dyrektorka ds. technologii w OpenAI i dyrektorka generalna Thinking Machines Lab, Mira Murati wskazała: „Niektóre kreatywne miejsca pracy być może znikną. Ale może w ogóle nie powinny istnieć, skoro treści, które dzięki nim powstają, nie są zbyt wysokiej jakości [bo może je generować AI – przyp. MT]”. Jej szef Sam Altman twierdzi w odpowiedzi na pytania autorów książki o przyszłości marketingu, Adama Brotmana i Andy’ego Sacka, że 95 proc. tego, do czego wykorzystywano agencje, strategów i kreatywnych profesjonalistów, będzie obsługiwane przez sztuczną inteligencję prawie bez żadnych kosztów. Wprawdzie Altman, mówiąc to, miał na myśli przyszłą AGI, ale ze względu na niejasności definicyjne i fakty, może to stać się, zanim nadejdzie ogólna AI.

Jedna ze znanych firm obsługujących płatności Klarna ujawniła, że jej asystent AI (5), zasilany przez OpenAI, wykonuje już pracę siedmiuset pracowników. Czatuje z klientami, rozwiązując zgłoszenia serwisowe w różnych językach oraz zarządzanie zwrotami. Agent AI przeprowadził parę milionów rozmów, dwie trzecie wszystkich czatów w ramach obsługi klienta firmy.

Z drugiej strony raport „Work Trend Index” firmy Microsoft wskazuje, że wbrew powszechnej opinii, sztuczna inteligencja stwarza nowe możliwości zatrudniania ludzi. Rekrutowani są jednak wyłącznie wykwalifikowani pracownicy z predyspozycjami do AI, co m.in. przyczyniło się do „142-krotnego wzrostu liczby użytkowników serwisu LinkedIn dodających do swoich profili umiejętności AI, takie jak Copilot i ChatGPT”. Trwa też poszukiwanie ludzi,



5. Logo Klarna z symbolem bota AI

którzy naprawialiby błędy pochopnego wdrażania AI, np. niektóre wydawnictwa przyjeły zwolniły większość swoich pracowników, szybko zdając sobie sprawę, że pospieszne zmiany dały często efekt przeciwny do zamierzonego, co z kolei zmusza te firmy do zatrudniania specjalistów, którzy poprawią błędną pracę AI i dodadzą ludzki charakter treściom.

Co ciekawe, współzałożyciel Instagrama i dyrektorka ds. produktu w firmie Anthropic, Mike Krieger, niedawno nadał wypowiedziom swojego szefa Amodi na temat programistów i kodowania nieco inny wydźwięk, wskazując, że rola inżynierów oprogramowania szybko ewoluuje i że wkrótce zaczną oni podwójnie sprawdzać kod generowany przez sztuczną inteligencję, zamiast go pisać. Także dyrektorka generalna Amazon Web Services Matt Garman twierdzi, że sztuczna inteligencja, owszem, może zmusić programistów do zaprzestania kodowania, ale nie na rzecz zajęcia się uprawą marchewki, lecz raczej skłaniając ich do podnoszenia kwalifikacji w tej dziedzinie... informatyki wyższego poziomu i tworzenia architektury systemów.

Paradoksalnie brzmią też takie opinie jak Sergeja Brina, współzałożyciela Google w „The New York Times”, że dla osiągnięcia AGI potrzeba więcej... pracy ludzi. Jego zdaniem cel ten zostanie osiągnięty, gdy wprowadzi się 60-godzinne tygodnie pracy.

Harari: lepiej się nie śpieszyć

Czy będzie taka jak my, czy będzie „osobliwie” inna, w rozmowie z Wired na początku kwietnia 2025 r. izraelski historyk i filozof, autor książek „Sapiens. Od zwierząt do bogów”, „Homo deus. Krótka historia jutra” i „Nexus. Krótka historia informacji od epoki kamienia do sztucznej inteligencji”, Yuval Noah Harari (6) zastanawia się „Jak podzielić się planetą z tą nową superinteligencją?”.

„Niektórzy ludzie często utożsamiają rewolucję AI z rewolucją drukarską, wynalezieniem słowa pisanego lub pojawieniem się mediów masowych, takich jak radio i telewizja, ale jest to nieporozumienie”, mówi Harari. „Wszystkie poprzednie technologie informacyjne były jedynie narzędziami w rękach ludzi. Nawet gdy wynaleziono prasę drukarską, to wciąż ludzie pisali tekst i decydowali, które książki wydrukować. Sama prasa drukarska nie może niczego napisać ani wybrać książek do druku. Sztuczna inteligencja jest jednak zasadniczo inna, jest agentem; może pisać własne książki i decydować, które pomysły rozpowszechniać. Może nawet samodzielnie tworzyć zupełnie nowe idee, co nigdy wcześniej nie miało miejsca w historii. My, ludzie, nigdy wcześniej nie mieliśmy do czynienia z superinteligentnym agentem”.

„Powodem, dla którego ludzie są w stanie współpracować na tak dużą skalę, jest to, że możemy tworzyć i dzielić się historiami. Wszelka współpraca na dużą skalę opiera się na wspólnej historii. Religia jest najbardziej oczywistym przykładem, ale historie finansowe i ekonomiczne są również dobrymi przykładami. Pieniądze są prawdopodobnie najbardziej udaną historią w historii. Zwierzęta nie potrafią tworzyć historii takich jak pieniądze. Ale sztuczna inteligencja może. Po raz pierwszy w historii dzielimy planetę z istotami, które potrafią tworzyć i łączyć historie lepiej niż my”.

Harari uważa, że najlepszym podejściem do rewolucji AI jest unikanie skrajności. „Na jednym końcu spektrum znajduje się strach, że sztuczna inteligencja pojawi się i zniszczy nas wszystkich, a na drugim końcu optymizm, że sztuczna inteligencja poprawi opiekę zdrowotną, edukację i stworzy lepszy świat”, mówi. „Przed wszystkim musimy zrozumieć skalę tej zmiany. W porównaniu z rewolucją AI, przed którą obecnie stoimy, wszystkie poprzednie rewolucje w historii błędą (...) Co więcej, sztuczna inteligencja ma potencjał do generowania pomysłów, które są dla nas całkowicie niezrozumiałe”.

„W moim rozumieniu osobliwość to punkt, w którym nie rozumiemy już, co się dzieje. To punkt, w którym nasza wyobraźnia i zrozumienie nie nadążają”, wyjaśnia Harari. „Być może jesteśmy już bardzo blisko tego punktu. Nawet bez komputera kwantowego lub w pełni rozwiniętej sztucznej inteligencji ogólnej, to znaczy sztucznej inteligencji, która może konkurować z możliwościami człowieka, poziom sztucznej inteligencji, który istnieje obecnie, może wystarczyć. (...) Co by się stało, gdyby miliony lub dziesiątki milionów zaawansowanych sztucznej inteligencji połączyły się w sieć, aby wprowadzić poważne zmiany w ekonomii, wojsku, kulturze i polityce? Kiedy nie będziemy już rozumieć świata, nie będziemy



6. Yuval Noah Harari

już kontrolować naszej przyszłości. Znajdziemy się wtedy w takiej samej sytuacji jak zwierzęta. Będziemy jak koń lub słoń, który nie rozumie, co dzieje się na świecie. Konie i słonie nie mogą zrozumieć, że ludzkie systemy polityczne i finansowe kontrolują ich przeznaczenie. To samo może przydarzyć się nam, ludziom”.

Jego zdaniem, skoro istnieje wiele zagrożeń i my to rozumiemy, rozsądnie byłoby spowolnić tempo rozwoju, zainwestować więcej w bezpieczeństwo i najpierw stworzyć mechanizmy bezpieczeństwa, aby upewnić się, że superinteligentne AI nie wymkną się spod naszej kontroli lub nie będą zachowywać się w sposób szkodliwy dla ludzi. Jednak obecnie dzieje się coś wręcz przeciwnego. Znajdujemy się w samym środku przyspieszającego wyścigu AI. Różne firmy i państwa ścigają się w zawrotnym tempie, aby opracować potężniejsze AI. Jednocześnie poczyniono niewielkie inwestycje, aby zapewnić bezpieczeństwo sztucznej inteligencji. „Uważam, że ogromnym błędem jest zakładanie, że ludzie mogą ufać sztucznej inteligencji, skoro nie ufają sobie nawzajem. Najbezpieczniejszym sposobem na rozwinięcie superinteligencji jest najpierw wzmocnienie zaufania między ludźmi, a następnie współpraca ze sobą w celu rozwinięcia superinteligencji w bezpieczny sposób. Ale to, co robimy teraz, jest dokładnie odwrotne. Zamiast tego wszystkie wysiłki są skierowane na rozwój superinteligencji”.

To wszystko tak czy inaczej, zdaniem Harari, nie oznacza, że sztuczna inteligencja sama w sobie jest zła lub szkodliwa. Problemem jest, że nie mamy żadnego modelu budowania społeczeństwa, swojego życia i ludzkiej cywilizacji z AI. Po prostu nie przerabialiśmy tego. ■

Mirosław Usidus



1. Jak AI wyobraża sobie ścianę, do której dotarła © AI

Wiara branży AI, na czele z OpenAI, że „większe jest lepsze”, zaczyna przygasać. Dwa i pół roku po uruchomieniu ChatGPT, po tym, jak wydano miliardy na skalowanie modeli, czyli budowanie coraz większych, zgromadzono do ich szkolenia i obsługi góry drogich chipów, pojawiło się zwątpienie, czy proste dodawanie oznacza realny postęp.

AI wciąż halucynuje, oszukuje, jest ograniczona i trudno zaufać jej bezpieczeństwu

PROBLEMY, KTÓRE NIE CHCĄ ODEJŚĆ

Zanim pojawił się nazywany Orionem model GPT-4.5 firmy OpenAI, dość powszechna w branży była opinia, że nie będzie to skok jakościowy w stosunku do poprzednika, GPT-4, na pewno nie tak imponujący, jaki ten ostatni stanowił w porównaniu do GPT-3. Na pocieszenie przyszła informacja, że 4.5 „wygrywa” z ludźmi test Turinga. Tylko, że to raczej nie oznacza, że najnowsze dzieło OpenAI jest „jak człowiek”. Główna konkurencja firmy Sama Altmana, czyli Google, Anthropic, Microsoft i Meta, oferuje całą masę nowych produktów, uderza w rynek fala „agenturalnej AI”, o której piszemy w oddzielnym

artykule w tym wydaniu. Tylko czy przypadkiem te wszystkie nowości nie mają charakteru bardziej marketingowego, niż stanowią jakikolwiek przełom w postępie technik sztucznej inteligencji?

Już kilkanaście miesięcy temu Bill Gates mówił, że jego zdaniem następca GPT-4 rozczaaruje. Jego, i wielu innych, przekonanie, że ta technika budowy dużych modeli językowych dociera do ściany, której nie pokona (1), co prowadzi do szerszych wątpliwości co do ewentualności powstania inteligencji podobnej do ludzkiej, znanej również pod niejasno zdefiniowanym pojęciem – sztuczna inteligencja ogólna (AGI), o czym również piszemy w tym numerze MT oddzielnie.

Branża rozpoczęła już badania i wdrażanie technik alternatywnych do podejścia ilościowego w zwiększaniu wydajności generatywnych modeli sztucznej inteligencji. Naukowcy próbują zmniejszyć modele, by zużywały mniej zasobów obliczeniowych, radząc sobie jednocześnie ze specjalistycznymi zadaniami. Niektórzy przepraszają się ze starszymi technikami, takimi jak nauka maszynowa ze wzmocnieniem firmy DeepMind. Użytkownicy narzędzi, których obecnie na rynku jest mnóstwo, wciąż jednak widzą niedoskonałości i ograniczenia współczesnej wersji AI. Wiedzą, że choć postęp w jakości, wydajności, trafności odpowiedzi i dokładności danych generowanych przez AI widać, wciąż nie można tej technice w pełni zaufać.

W czym interesie?

Są też inne, zasadnicze zastrzeżenia w stosunku do rozwijanych obecnie systemów AI, np. twórca WWW, Tim Berners-Lee (2), chce wiedzieć, dla kogo pracuje sztuczna inteligencja. Przypomina, że przez dekady on i inni pracowali nad budową otwartego internetu, tymczasem generatywna AI nie powstaje w ten sposób. Jego zdaniem, obowiązkiem twórców nowych modeli AI jest sprawienie, że narzędzia te będą działać dla użytkowników, a nie tylko dla firm, które je finansują.

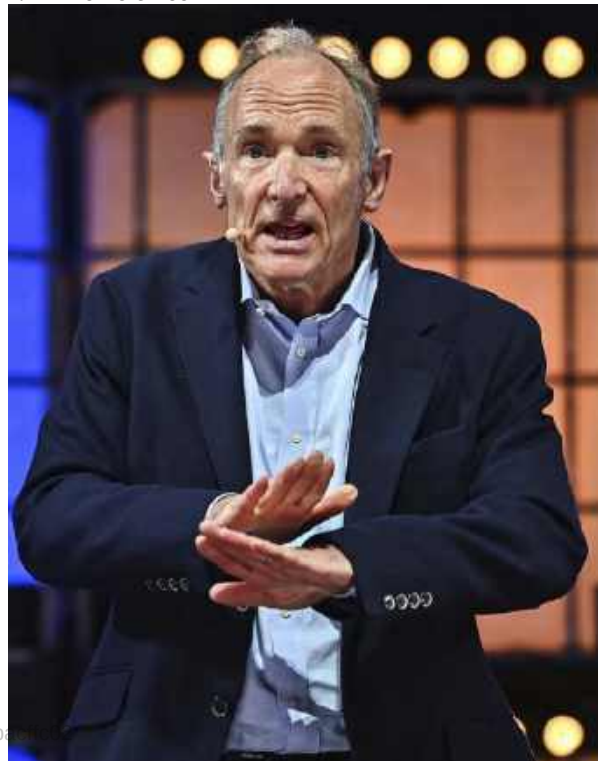
Berners-Lee posługuje się analogią z lekarzami, którzy mogą być zatrudniani przez różne podmioty, publiczne czy prywatne, ale zawsze mają obowiązek działać w najlepszym interesie pacjenta. Tymczasem asystent/agent AI, który pomaga użytkownikowi np. zaplanować wakacje lub kupować produkty, może być przeszkolony, by działać na rzecz konkretnych sprzedawców a nie w najlepszym interesie użytkownika. „Agent AI”, którego ktoś uważa za swojego, może być co najmniej podwójnym agentem, albo wręcz fałszywym agentem.

Berners-Lee mówił też, że nie ma w świecie AI organizacji takiej jak W3C, która od lat 90. XX wieku ustanawia standardy dla internetu. Sugeruje, by twórcy sztucznej inteligencji stworzyli podobną organizację. Warto jednak zwrócić uwagę, że te standardy i wspólna praca nad „otwartym” internetem dotyczą w ostatnich latach głównie Zachodu, gdyż wiele krajów, z większych przykładów – Chiny i Rosja, odgradza się różnymi barierami od światowej sieci, a nawet buduje swoją własną, także w sensie technicznym, sieć. Nawet gdyby doszło do powołania czegoś takiego jak W3C dla AI, to zapewne doszłoby do podobnych podziałów.

Dziurawy jak chiński Deepseek

Z jednej strony trwa walka firm o większe zaufanie do sztucznej inteligencji, przez próby zwiększenia jej dokładności i wydajności. Z drugiej niektórzy zauważają, że wzrost zaufania do AI jest sam w sobie problemem, gdyż, jak wynika z niektórych badań, prowadzi do zaniku krytycznego myślenia u ludzi. Takie są wnioski z analiz przeprowadzonych przez Uniwersytet Carnegie Mellon i Microsoft. „Używane w niewłaściwy sposób technologie mogą i powodują pogorszenie zdolności poznawczych, które powinny zostać zachowane”, napisali naukowcy w komunikacie badawczym. Zespół ten przeprowadził ankiety wśród ponad trzystu „pracowników wiedzy”, czyli osób, które na różne sposoby

2. Tim Berners-Lee



rozwiązują problemy, badające ich doświadczenia z wykorzystaniem generatywnej sztucznej inteligencji w miejscu pracy. Ankietowani specjaliści zostali poproszeni o podanie trzech przykładów z życia wziętych, w których korzystali z narzędzi sztucznej inteligencji w pracy. Ogólny wniosek jest taki, że ci, którzy ufali narzędziom AI, myśleli mniej krytycznie, zaś osoby, które mniej ufały technice, cechowały się wyższym stopniem krytycznego myślenia. Naukowcy zauważyli również, że wykorzystanie sztucznej inteligencji wydaje się utrudniać kreatywność, a pracownicy korzystający z narzędzi AI wytwarzają „mniej zróżnicowany zestaw wyników dla tego samego zadania” w porównaniu do osób polegających na własnych zdolnościach poznawczych.

Jedną z ostatnich rzeczy, w której AI można ufać, jest bezpieczeństwo danych. Tak wynika przynajmniej z serii najnowszych doniesień. Badacze z firmy Wiz zajmującej się bezpieczeństwem technik chmurowych przyjrzeni się strukturom baz danych chińskiego, głośnego kilka miesięcy temu, modelu AI DeepSeek. Odkryli, że „w ciągu kilku minut” można tam z łatwością uzyskać dostęp do skarbnicy całkowicie niezaszyfrowanych danych wewnętrznych modelu. „Baza danych zawierała znaczną ilość historii czatów, danych z zaplecza i poufnych informacji”, informuje Wiz w swoim raporcie na temat luk w zabezpieczeniach chińskiego produktu. Co gorsza, ta szeroko otwarta tylna furtka w modelu sztucznej inteligencji o otwartym kodzie źródłowym mogła z łatwością doprowadzić do ataku na systemy DeepSeek. „Bez żadnego mechanizmu uwierzytelniania lub obrony przed światem zewnętrznym”, podkreślali badacze Wiz. DeepSeek podjął działania w celu zabezpieczenia swoich baz danych, gdy przedstawiciele Wiz zaalarmowali firmę o lukach, ale chyba nie o to chodzi. Potem zresztą w rozmowie z serwisem „Wired” przedstawiciele Wiz opowiadali, że trudno było skontaktować się z kimkolwiek w DeepSeek. Nikt z chińskiej firmy nie odpowiedział na próby kontaktu ze strony Wiz, ale w ciągu godziny baza danych została zablokowana. Warto przypomnieć w tym kontekście pytania Bernersa-Lee – kto za tym stoi i w czym interesie AI działa (3)?

Uzyskując dostęp do baz danych DeepSeek, badacze Wiz dowiedzieli się wiele o sposobie działania modeli firmy, w tym o tym, że jej infrastruktura niemal dokładnie naśladuje infrastrukturę OpenAI. Sama OpenAI odkryła niedawno dowody na istnienie chińskiego narzędzia do inwigilacji opartego na sztucznej inteligencji. Twierdzi, że służy do identyfikacji antychińskich postów w mediach społecznościowych



3. Telefon z aplikacją DeepSeek na tle flagi Chińskiej Republiki Ludowej

w krajach zachodnich. Operację tę przedstawiciele OpenAI nazwali Peer Review. Ben Nimmo z OpenAI powiedział, że to pierwszy znany mu przypadek takiego narzędzia szpiegującego opartego na sztucznej inteligencji, w tym przypadku na modelu open source Llama firmy Meta. W szczegółowym raporcie na temat wykorzystania sztucznej inteligencji do złośliwych i oszukańczych celów, OpenAI pisze też o innej chińskiej operacji o nazwie Sponsored Discontent, która wykorzystywała techniki OpenAI do generowania anglojęzycznych postów krytykujących chińskich dysydentów a także do tłumaczenia artykułów na język hiszpański przed ich dystrybucją w Ameryce Łacińskiej. Artykuły te krytkowały Stany Zjednoczone. Oddzielnie, badacze OpenAI zidentyfikowali kampanię, która prawdopodobnie prowadzona była z Kambodży, wykorzystującą modele AI firmy do generowania i tłumaczenia komentarzy w mediach społecznościowych wspierających działania oszustów w ramach kampanii znanej jako „rzeź świń”. Były wykorzystywane do wabienia użytkowników sieci i wpłatywania ich w oszukańczy schemat inwestycyjny.

Chyba więc słusznie rosną obawy, że sztuczna inteligencja może być wykorzystywana do inwigilacji, hakowania komputerów, kampanii dezinformacyjnych i innych złośliwych celów. Jednak badacze dodają, że AI może również pomóc zidentyfikować i powstrzymać takie zachowanie.

Zwodnicza łatwość kodowania

Kolejnym problemem zaufania, jeśli chodzi o sztuczną inteligencję, jest wiara w to, że generatywna AI już jest w stanie zastąpić pracę programistów. W sieci bez trudu można znaleźć mnóstwo pełnych frustracji i złości relacji ludzi, którzy

za bardzo w to uwierzyli, a błędy, które narzędzia popełniły w kodzie, mogą być bardzo kosztowne. Choćby dlatego, że ich znalezienie, identyfikacja i poprawianie mogą być ekstremalnie trudne i czasochłonne.

Wiele badań wykazuje, że modele LLM odpowiadają w sposób różny na różne wersje tego samego pytania. Szczególnie zwindnicze jest to wtedy, gdy zmieniane są wartości liczbowe w zapytaniach i rośnie liczba klauzul warunkowych, co w promptach dotyczących kodu programistycznego jest typową praktyką. Na zdrowy rozum trudno w pełni ufać narzędziom, które nie są zdolne do prawdziwego logicznego rozumowania i zamiast tego próbują replikować kroki rozumowania zaobserwowane w danych treningowych. Jednak liczba zwolenników poglądu, że modele rozwiązują problem kodowania, rośnie, choć przekonanie to oparte jest na obserwacjach realizacji przez AI stosunkowo prostych zadań w dziedzinie programowania. Napisanie krótkiego skryptu, który LLM często widział w danych treningowych, takich jak wyrażenie regularne, dostęp do bazy danych, proste makro przetwarzające format obrazu w programie graficznym, to wszystko AI generatywna robi dość sprawnie i zazwyczaj bezbłędnie. Jednak gdy użytkownik próbuje zaprzęgnąć AI do zrobienia czegoś bardziej skomplikowanego, efekty są znacznie gorsze. Charakterystyczne jest np., że narzędzia te proponują mnóstwo starszych, wcześniej sprawdzonych rozwiązań dla zupełnie nowego problemu. To rzadko działa.

Przed LLM-ami ludzie w branży programistycznej czasem żartowali, że kodowanie to tylko kopiowanie i wklejanie z repozytoriów takich jak Stack Overflow. Teraz robi to model sztucznej inteligencji, co jest wygodniejsze i szybsze. Jednak podobnie jak wcześniej kopiowanie gotowych fragmentów kodu nie było prawdziwym wyzwaniem w programowaniu, tak obecnie narzędzie AI, które ułatwia kopiowanie, nie rozwiązuje żadnych ważnych problemów, do których potrzebne są logiczne myślenie, wiedza i w końcu kreatywność programisty. W przypadku narzędzi AI kopiowanie to należy ograniczyć tylko do elementów, które łatwo zweryfikować. W przypadku czegokolwiek bardziej skomplikowanego pojawia się ryzyko wprowadzenia słabo rozpoznawalnych błędów, których naprawienie zajmuje dużo czasu.

Wątpliwe wyniki i demencja

Inna dziedzina, w której AI miałyby coś „zastąpić”, to wyszukiwanie w internecie (4). Narzędzia wyszukiwania AI szybko zyskują na popularności. Chyba każdy widział w Google odpowiedź AI

zamiast pierwszego wyniku na liście. Pojawił się jednak pewien problem. Gdy tradycyjne wyszukiwarki zazwyczaj działają jako pośrednik, kierując użytkowników do stron z wiadomościami i innych treści, narzędzia do wyszukiwania generatywnego same analizują i przepakowują informacje, odcinając przepływ ruchu do oryginalnych źródeł. Wyniki z chatbotów często zaciemniają kwestie związane z jakością źródeł informacji.

Tow Center for Digital Journalism przeprowadziło testy ośmiu generatywnych narzędzi wyszukiwania z funkcjami wyszukiwania na żywo, by ocenić ich zdolność do dokładnego pobierania i cytowania treści wiadomości. Okazało się, że chatboty generalnie źle radziły sobie z odmawianiem odpowiedzi na pytania, na które nie potrafiły udzielić dokładnej i w pełni poprawnej odpowiedzi, oferując w zamian halucynacje, błędne lub spekulatywne odpowiedzi. W dodatku, według tych wyników, „chatboty premium” udzielały bardziej błędnych odpowiedzi niż ich darmowe odpowiedniki. Wiele chatbotów zdawało się omijać ustawienia Robot Exclusion Protocol, które są standardem w internecie. Generatywne narzędzia wyszukiwania tworzyły linki, cytowały i kopiowały artykuły, do których nie miały prawa, ze względu na prawo autorskie.

Charakterystyczne dla AI było, jak piszą autorzy badania, że halucynując, zachowuje pewność siebie i rzadko używa zwrotów dystansujących i tonujących, takich jak „wydaje się”, „jest to możliwe”, „może” itp. Czy też przyznaje się do luki, używając stwierdzeń takich jak „nie mogłem zlokalizować dokładnego artykułu”. Z wyjątkiem Copilota, który odrzucił więcej pytań, niż udzielił odpowiedzi, wszystkie narzędzia konsekwentnie częściej udzielały nieprawidłowej odpowiedzi, niż przyznawały się do luk w wiedzy czy ograniczeń. Modele płatne, takie Perplexity Pro lub Grok 3, jak wynikało z testów, odpowiadały poprawnie na więcej podpowiedzi niż ich darmowe odpowiedniki, jednak paradoksalnie wykazały również wyższy poziom błędów. Wynika to przede wszystkim z ich tendencji do udzielania kategoriycznych, ale błędnych odpowiedzi, zamiast odmawiania bezpośredniej odpowiedzi na pytanie.

Pięć z ośmiu chatbotów testowanych w tym badaniu (ChatGPT, Perplexity i Perplexity Pro, Copilot i Gemini) upubliczniło nazwy swoich crawlerów (skryptów, których celem jest indeksowanie lub pozyskiwanie informacji ze stron internetowych), dając wydawcom możliwość ich zablokowania, podczas gdy crawlers używane przez pozostałe trzy (DeepSeek, Grok 2 i Grok 3) nie są publicznie



4. Wyszukiwanie AI

znane. Z drugiej strony, czasami poprawnie odpowiadały na zapytania dotyczące wydawców, do których treści nie powinny mieć dostępu. Perplexity Pro poprawnie identyfikował prawie jedną trzecią z dziewięćdziesięciu fragmentów artykułów, do których nie powinien mieć dostępu, gdyż wydawcy na to nie zezwalają na swoich stronach. Takich przypadków i niejasności w kwestii stosowania crawlerów tam, gdzie wyszukiwarkom teoretycznie nie wolno, wynika z testów więcej. Chociaż protokół wykluczenia robotów (crawlerów) nie jest prawnie wiążący, jest to powszechnie akceptowany standard sygnalizowania, które części serwisu WWW powinny, a które nie powinny być indeksowane. Ignorowanie protokołu pozbawia wydawców możliwości decydowania o tym, czy ich treści będą uwzględniane w wyszukiwaniach lub wykorzystywane jako dane szkoleniowe dla modeli sztucznej inteligencji. Wydawcy mogą mieć różne powody, dla których nie chcą, aby roboty indeksujące miały dostęp do ich treści, takie jak chęć zarabiania na ich treściach lub obawa, że ich praca może zostać błędnie przedstawiona w podsumowaniach generowanych przez sztuczną inteligencję.

Nawet jeśli chatboty poprawnie identyfikowały artykuł, do którego miały legalny dostęp, to często nie potrafiły poprawnie odsyłać do oryginalnego źródła, lecz do przedruku lub skrótu w innym

miejscu. Ponad połowa odpowiedzi z Gemini i Grok 3 cytowała sfabrykowane lub uszkodzone adresy internetowe. Problem ten nie dotyczył wyłącznie Grok 3 i Gemini, zdarzał się on znacznie rzadziej w przypadku innych chatbotów. Dopełnia to obrazu krytyki wyszukiwarek AI, które do wszystkich błędów i halucynacji, które dodają jeszcze żerowanie na treściach, które ktoś czasem sporym kosztem przygotował, a odpowiedź chatbota zawiera błędny odsyłacz do jego pracy lub prowadzi gdzie indziej, do miejsca, w którym jego oryginalnej pracy nie ma.

Do wszystkich zarzutów wobec AI dochodzi też spostrzeżenie, wręcz diagnoza, że dotyka ją demencja. Naukowcy z Izraela przetestowali sztuczną inteligencję pod kątem spadku zdolności poznawczych. Odkryli, że LLM cierpią na formę upośledzenia funkcji poznawczych podobną do spadku u ludzi. Jest ona bardziej dotkliwa w przypadku starszych modeli. Zespół badał GPT-4 i 4o, dwie wersje Gemini firmy Alphabet oraz 3.5 Claude firmy Anthropic. W opublikowanym artykule opisującym wyniki badań neurolog Roy Dayan i Benjamin Uliel z Centrum Medycznego Hadassah, oraz Gal Koplewitz, naukowiec zajmujący się danymi na uniwersytecie w Tel Awiwie, wskazują na poziom „spadku zdolności poznawczych, który wydaje się porównywalny z procesami neurodegeneracyjnymi w ludzkim mózgu”. ■

Mirosław Usidus



O tych, co przekuli innowacyjne wizje w biznesowy sukces

W polskim życiu publicznym coraz częściej używanym słowem jest odmieniany na wszystkie sposoby wyraz „innowacje”. I tak powinno być przez najbliższe lata, bo ambicją naszego kraju jest spektakularny awans do grona państw o gospodarce kreatywnej, tworzącej własne produkty i marki, znane i szanowane w świecie.

To Wy, młodzi Czytelnicy MT, macie tego dokonać! Żeby Was natchnąć dobrymi przykładami, co miesiąc przedstawiamy reprezentantów czołówki światowych liderów innowacji. Najczęściej byli oni jeszcze w wieku szkolnym lub studenckim, gdy w ich głowach rodziły się śmiałe pomysły skutkujące później powstaniem superproduktów, wielkich brandów i fantastycznych fortun.

To oni kształtują cywilizację technologiczną.

To bohaterowie naszych czasów.

Człowiek i jego „Cywilizacja” – **Sid Meier**

Jego gry przyciągnęły do grania osoby, który wcześniej nawet nie myślały, że będą „się w to bawić” na komputerze czy w konsoli. Do budowania „Cywilizacji” publicznie przyznawali się w mediach znani aktorzy i politycy. Sam pozostał dość skromną osobą, bez gwiazdorskich manier.



1. Sid Meier

CV: Sidney (Sid) K. Meier

Data i miejsce urodzenia: 24.02.1954, Sarnia, Ontario, Kanada

Adres zamieszkania: Cockeysville, Maryland, USA

Obywatelstwo: kanadyjskie, szwajcarskie, amerykańskie

Stan cywilny: żonaty, jeden syn

Majątek: szacowany na ok. 100 mln USD

Kontakt: www.linkedin.com/in/melanieperkins

Edukacja: uniwersytet w Michigan

Doświadczenie zawodowe: przed 1982 – praca przy projektowaniu systemów kasowych, 1982–96 – współzałożyciel i współwłaściciel MicroProse, 1996 do dziś – współzałożyciel i dyrektor kreatywny Firaxis Games

Zainteresowania: historia, gry planszowe, muzyka, lotnictwo

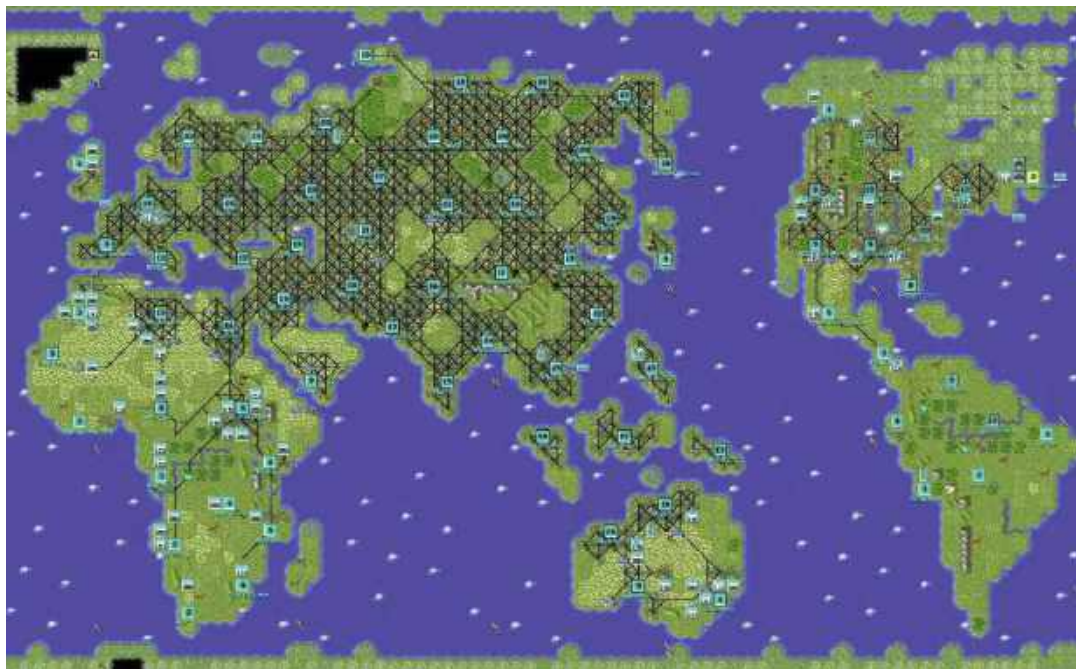
Sid Meier (1) urodził się w Kanadzie. Jego rodzice mieli pochodzenie holenderskie i szwajcarskie, dzięki czemu przy urodzeniu otrzymał zarówno obywatelstwo kanadyjskie, jak i szwajcarskie. Gdy miał około trzech lat, jego rodzina przeniosła się do amerykańskiego Detroit, gdzie się wychował. Na Uniwersytecie Michigan studiował historię i informatykę, uzyskując tytuł „Bachelor of Arts” w dziedzinie informatyki w 1975 roku.

Nazwisko w tytule gry znakiem jakości

Po studiach Sid zaczął pracować, zajmując się początkowo rzeczami dalekimi od rozrywki. Pomagał przy rozwoju systemów kasowych dla domów towarowych. Około 1981 roku kupił komputer Atari 800, co, według zapisków biograficznych, uświadomiło mu potencjał, jaki może mieć programowanie komputerowe w tworzeniu gier wideo. Po tym jak poznał Billa Stealeya, na początku lat 80., który miał podobne zainteresowania, pojawił się pomysł wspólnej działalności przy tworzeniu gier. Założyli w 1982 roku firmę MicroProse. Po opracowaniu kilku gier akcji, w tym „Floyd of the Jungle”, MicroProse zainicjowało serię symulatorów lotu, zaczynając od Hellcat Ace (1982) i kontynuując Spitfire Ace (1982), Solo Flight (1983) i F-15 Strike Eagle (1985), wszystkie zaprojektowane i zaprogramowane przez Meiera.

Pierwszą komercyjną grą stworzoną jako personalnie gra Sida Meiera (jego nazwisko pojawiło się na opakowaniu gry) była „Formula 1 Racing” jeszcze z 1982 roku. Pierwszą zaś grą sygnowaną oficjalnie jego nazwiskiem w tytule była „Sid Meier's Pirates!”, która ukazała się w 1987 roku. To Stealey uznał, że nazwisko Meiera sprawi, że ci, którzy kupili symulatory lotu firmy, będą bardziej skłonni zagrać w nową grę. Jak wspominał w jednej z rozmów: „Byliśmy na kolacji na spotkaniu Software Publishers Association i był tam Robin Williams [aktor, już wtedy bardzo znany – przyp. MT]. Przez dwie godziny nas rozśmieszał. Odwrócił się do mnie i powiedział: Bill, powinieneś umieścić nazwisko Sida na kilku tych pudełkach i promować go jako gwiazdę. W ten sposób nazwisko Sida znalazło się w tytułach Piratów i Cywilizacji”. Sam Meier podkreślał nie raz, że nie zachęcał MicroProse do promowania jego nazwiska. Okazało się jednak, że stało się ono znakiem jakości produkcji, który pomagał w sprzedaży gier. Jeśli chodzi o różne inne tytuły wydawane przez firmę i sygnowane „Sid Meier's”, to w rzeczywistości nie zawsze był głównym projektantem tytułów sygnowanych jego nazwiskiem.

„Pirates” to gra symulacyjna z otwartym światem. Zawierała wiele cech, z których gry Meiera stały się z czasem najbardziej znane, w tym m.in. stosunki dyplomatyczne między krajami i rozwój gospodarek symulowanych miast. Cechy te rozwinięte zostały w pierwszej wersji „Cywilizacji” (ang. „Civilization”), która ukazała



2. Mapa świata pochodząca z pierwszej wersji 'Cywilizacji'



się w 1991 roku (2). To na niej opiera się sława i uznanie Meiera, wykraczające poza środowisko graczy.

Od czasu premiery w 1991 roku, seria „Cywilizacja” sprzedawała się w kilkudziesięciu milionach egzemplarzy na całym świecie. Nie brakuje opinii, że gra ta na trwałe zakorzeniła się w naszej kulturze popularnej, stała się tematem poważnych dyskusji na temat wizji rozwoju świata i reguł w nim panujących. Meier podkreśla, że jego priorytetem w kształtowaniu świata w „Cywilizacji” było po prostu stworzenie przyjemnego i satysfakcjonującego doświadczenia dla gracza. „Jednym z naszych podstawowych celów było to, aby nie rzutować naszej własnych poglądów filozoficznych czy politycznych na przebieg gry”, mówi Meier. Jednak można znaleźć wiele ocen, że gra narzuca jednak pewien określony światopogląd. Być może jednak w przypadku gier Meiera doszukiwanie się poważnych znaczeń nie jest w ogóle potrzebne. „Computer Gaming World” pisał w 1994 roku, że „Sid Meier kładzie nacisk na ‘zabawne elementy’ symulacji, a resztę odrzuca”. Według magazynu, „Meier upierał się, że odkrycie nieuchwytnej wartości zabawy jest najtrudniejszą częścią projektowania gier”.

„Silnik gry” znany tylko twórcom

Około 1990 roku Stealey powziął plan rozszerzenia działalności MicroProse o produkcję gier zręcznościowych, co Meier uważał za zbyt ryzykowne. Nie mogąc porozumieć się w tej sprawie z kolegą, Meier sprzedał Stealeyowi swoją połowę firmy, ale pozostał w niej na tym samym stanowisku. Firma MicroProse została wykupiona przez Spectrum Holobyte w 1993 roku. Nowy właściciel i jego praktyki nie spodobały się Meierowi, który opuścił firmę wraz z innymi pracownikami, m.in. Jeffem Briggssem i Brianem Reynoldsem i stworzył z nimi nową firmę, Firaxis Games. Kontynuował w niej tworzenie kolejnych wersji „Cywilizacji” i innych tytułów, w tym unowocześnionej wersji (remake’u) „Sid Meier’s Pirates!”.

Według osób związanych z Firaxis, od około 1996 roku Meier nieustannie rozwijał własny, specjalny „silnik gry”, którego używał do prototypowania swoich pomysłów i którego nie udostępniał nikomu innemu. Jego współpracownicy uważają do dziś, że silnik jest oparty na oryginalnym źródle „Cywilizacji”, przez lata rozbudowywany o aktualizacje, które tworzył sam Meier lub inni programiści.

Sid zyskał uznanie także w innych niż tworzenie gier dziedzinach. W 1996 roku przyznano mu patent na „System do komponowania i syntezy muzyki w czasie rzeczywistym” wykorzystany w grze „C.P.U. Bach”. W roku 1999 r. jako drugi człowiek został uhonorowany



3. Gwiazda Sida Meiera w Walk of Game w Sony Meireon, rozrywkowym centrum w San Francisco

członkostwem w Galerii Sław Akademii Nauk i Sztuk Interaktywnych (ang. Academy of Interactive Arts and Sciences), zasiadając tam obok Shigeru Miyamoto z firmy Nintendo. Nagród, wyróżnień i honorów ma na koncie znacznie więcej. W 2008 roku otrzymał nagrodę za całokształt twórczości na konferencji Game Developer’s Conference. W 2009 roku zajął drugie miejsce na liście IGN „najlepszych twórców gier wszech czasów” i został nazwany „idealnym wzorem do naśladowania dla każdego początkującego projektanta gier”. W 2017 roku został nagrodzony za całokształt twórczości przez Golden Joystick Awards. Ma też swoją gwiazdę w galerii Walk of Game w San Francisco (3).

Jesień życia w świecie gier

W 2005 roku Firaxis zostało kupione przez 2K Games. Wkrótce ukazała się szósta odsłona „Cywilizacji”, która została nawet przeniesiona na Androida i iOS. Sid Meier, choć już niemłody (ma obecnie 71 lat), nadal pracuje jako dyrektor kreatywny Firaxis Games. Choć obecnie rzadziej bezpośrednio projektuje gry, jego wpływ na studio pozostaje znaczący.

W 2020 roku ukazała się jego autobiografia „Sid Meier’s Memoir! A Life in Computer Games”. Książka, która zawiera wspomnienia żyjącej legendy gier komputerowych i wideo, nie eksponuje jego osoby. Skupia się na grach i historiach związanych z ich powstawaniem.

W połowie 2024 r. 2K z Firaxis Games oficjalnie ogłosiły, że „Sid Meier’s Civilization® VII” pojawi się w 2025 roku na konsolach i w komputerach za pośrednictwem Steam. Od lutego gra jest dostępna dla starszych i nowych pokoleń graczy. ■

Mirosław Usidus

Uczymy się AI z „Młodym Technikiem”

Pomocnik AI w Google Workspace

Popularność modeli generatywnych AI rośnie systematycznie od kilku lat. Google, jak wiele innych firm technologicznych, gorliwie wzmacnia swoją pozycję przez budowę reaktorów jądrowych wyłącznie do zasilania AI [1] czy miliardową inwestycję w firmę Anthropic [2]. Można powiedzieć, że obecnie użytkownicy oswoili się z modelami językowymi, ich możliwościami oraz ograniczeniami. Kolejnym krokiem rozwoju tych technologii jest bezpośrednia integracja z narzędziami, z którymi są używane.

Ciągle przełączanie okien pomiędzy naszym dokumentem a aplikacją czatu szybko staje się irytujące, szczególnie kiedy musimy stale informować naszego wirtualnego asystenta o wszystkich zmianach, których dokonaliśmy oraz nieustannie przypominać mu o kontekście. Problem ten doczekał się już pierwszych komercyjnych rozwiązań. W artykule przyjrzymy się efektom pracy firmy Google, która oferuje używanie modelu Gemini bezpośrednio w swoich aplikacjach.

Wybór subskrypcji

Google oferuje dostęp do swoich modeli w ramach różnych subskrypcji oraz pakietów, np. Google Workspace dla firm. Ponieważ interesuje nas wyłącznie prywatne korzystanie z pełnych możliwości modelu Gemini, najbardziej przystępną opcją jest rozszerzenie darmowego konta Gmail o pakiet 'Google One AI Premium'. Przy obecnej cenie 98 zł miesięcznie subskrypcja oferuje wszystkie zaawansowane funkcje Gemini oraz darmowy miesiąc próbny. Po zakupie oprócz uzyskania dostępu do zaawansowanych modeli Gemini oraz generowania zdjęć, odblokowane zostają nowe możliwości w aplikacjach: Gmail, Google Docs, Google Drive, Google Chat, Google Sheets, Google Slides.

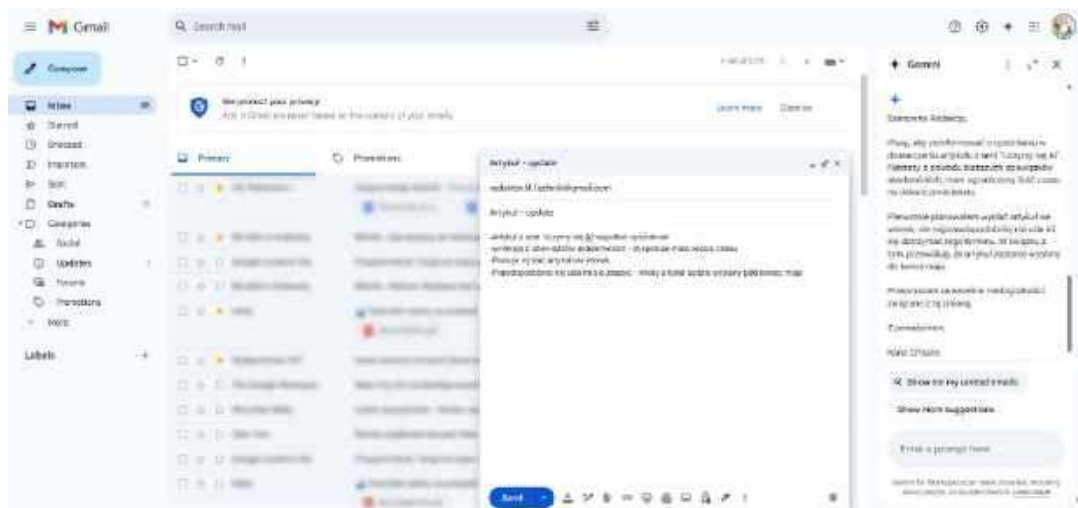
Gemini panel

Po zakupie subskrypcji w dowolnej aplikacji Google odnajdujemy gwiazdkę umieszczoną zazwyczaj w prawym górnym rogu ekranu. Przy pierwszym uruchomieniu konieczne jest wyrażenie zgody na udostępnianie



danych modelowi Gemini. Po kliknięciu obok głównego okna aplikacji pojawia się sekcja z uproszczonym interfejsem czatu. Sekcja ta jest ogólna i jednakowa dla wszystkich aplikacji Google.

Panel czatu jest w dużej mierze odizolowany od samych aplikacji, używa specjalnej akcji, aby uzyskać dostęp do e-maili i dysku użytkownika, co potrafi być niezwykle frustrujące! Gdy poprosiłem chat o odczytanie najnowszego nieprzeczytanego maila, pobrał wiadomość z folderu Spam. Natomiast gdy poprosiłem o odczytanie tego samego folderu, odpowiedział: „Przepraszam, ale nie mam takiej funkcji”. Podczas korzystania należy brać pod uwagę, że Gemini nie widzi

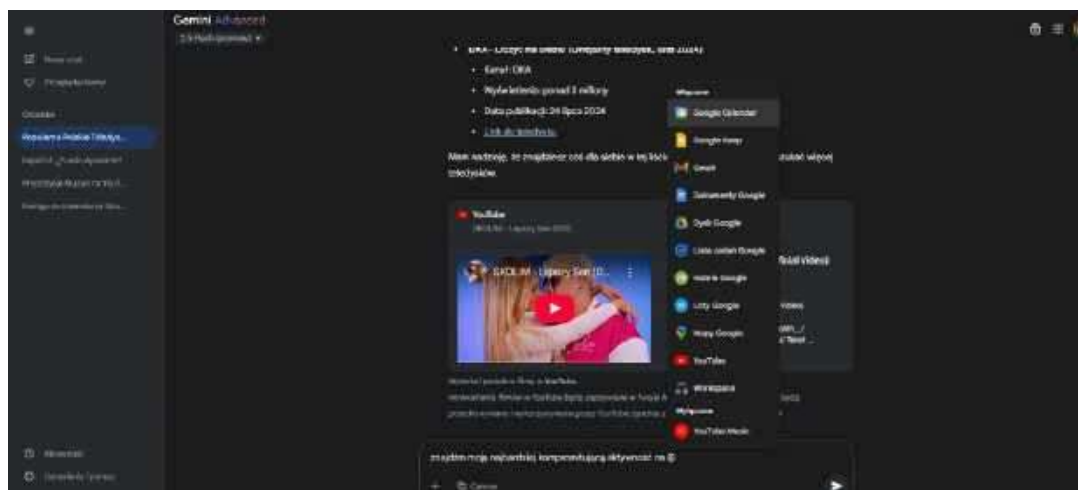


dokładnie tego, co użytkownik. Jest to poważna wada i wymaga od użytkownika wyraźnego wskazania przestrzeni roboczej. Najprostszą opcją jest zaznaczenie tekstu, który ma zostać przetworzony. Czat bez problemu wyłapuje te akcje w swoim oknie. Kolejnym sposobem jest podanie dokładnej nazwy pliku z Dysku Google, poprzedzając ją znakiem @.

Warto zaznaczyć, że obecna implementacja przypomina produkt wczesnego dostępu. Jeśli poprosimy Gemini o coś bez wystarczającej precyzji, w odpowiedzi wygenerowany zostanie błąd: „Nie mogę w tym pomóc. Czy mogę ci pomóc w czymś innym?”. Kiedy zaś domagamy się dalszego wyjaśnienia, czat odeśle nas do swojej dokumentacji. Paradoksalnie, ponieważ podsumowanie dokumentacji miało być jednym z głównych zastosowań modelu.

Skupmy się więc na obecnych możliwościach panelu. Jest w stanie pozyskiwać potrzebne informacje

z naszych wiadomości e-mail oraz plików zapisanych w chmurze Google, jednak wymaga od użytkownika dużej precyzji. Gdy poprosiłem go o godzinę odjazdu pociągu, którym ostatnio podróżowałem, pominął pociągi należące do Polregio oraz Kolei Śląskich. Zamiast tego wybrał bilet PKP Intercity sprzed dwóch miesięcy. Gemini nie dokonuje głębokiej analizy danych do pobrania – oczekuje jednoznacznie podanej nazwy i wybiera pierwszą dostępną opcję. Na szczęście chat zawsze podaje źródło, z którego czerpał informacje, więc uważny użytkownik zauważy jego błędy. Kiedy bezpośrednio podałem interesujący mnie plik (Designing Data-Intensive Applications. – Martin Kleppmann, 612 stron), Gemini szybko odnalazł konkretne informacje w dużym zbiorze danych i podał je bez halucynacji. Dodatkowo oparł się błędowi użytkownika. Kiedy poprosiłem o informacje niepojawiające się w pliku, uważnie przeanalizował cały tekst



i stwierdził, że nie zawiera poszukiwanych danych i podał pokrewne informacje obecne w tekście. Czat potrafi również odczytywać niewysłane e-maile, co umożliwia korzystanie z panelu podczas pisania wiadomości e-mail bez przełączania okien. Gemini potrafi sugerować edycje stylistyczne i gramatyczne, wyszukiwać informacje w Internecie oraz generować zdjęcia. Niestety za pośrednictwem panelu Gemini nie może ingerować w pliki użytkownika. Wpisywanie poprawek do naszej pracy jest obecnie możliwe tylko w anglojęzycznych wersjach aplikacji.

Trudno nie zauważyć, że panel dostępny w aplikacjach Google jest zdecydowanie mniej funkcjonalny niż wersja dostępna na głównej stronie usługi. Istnieje alternatywne rozwiązanie: zamiast korzystać z czatu w ramach aplikacji, można uzyskać dostęp do aplikacji za pośrednictwem interfejsu czatu. Aby udzielić modelowi dostępu do danych z wybranej aplikacji, wystarczy dodać flagę `@{nazwa_aplikacji}` w treści zapytania. W tym trybie dowolny model Gemini może analizować dane użytkownika bez narzuconych ram kontekstowych.

Możemy poprosić Gemini o wyciągnięcie dowolnych danych, które przechowujemy w usługach Google. Jest to jednocześnie wygodne i lekko niepokojące, gdy model interpretuje wrażliwe informacje, do których Google i tak miał już dostęp wcześniej.

Help me write/create

Aby skorzystać z bardziej wnikliwych funkcji Gemini, należy najpierw upewnić się, że język konta Google jest ustawiony na angielski. Ustawienie można znaleźć w sekcji zarządzania kontem Google, wpisując w wyszukiwarkę frazę „język”. Niektóre aplikacje mogą dodatkowo wymagać wylogowania się, aby

wyświetlić zmiany. Tryb pomocy Gemini aktywuje się poprzez ikonę gwiazdki z ołówkiem. Znajduje się ona w dolnym pasku z narzędziami podczas pisania wiadomości e-mail oraz w menu wyświetlanym poprzez naciśnięcie prawego przycisku myszy w dokumentach Google. Po użyciu funkcji „help me write” Gemini generuje odpowiedź zgodnie z naszymi wytycznymi. Dostępne są również funkcje: „Skróć”, „Wydłuż” i „Popraw”, służące do szybkiej modyfikacji tekstu. Gemini nie prezentuje zmian w formie komentarzy, sugestii ani edycji istniejącego tekstu. Zamiast tego generuje zmieniony akapit w pobliżu kursora, nie zmieniając oryginału. Obecnie na Gemini nałożone są ograniczenia, powodujące, że odrzuca on teksty w językach innych niż angielski. Nie jest jednak łatwo „uwięzić” modele językowe. Wystarczy wydać mu polecenie w języku angielskim i dodać, że oczekujemy odpowiedzi w języku polskim.

Najbardziej zaawansowaną funkcję znajdziemy w Dokumentach Google, gdzie Gemini oferuje samodzielnie wykonanie całej pracy według podanych wytycznych. Zajmie się on wtedy tworzeniem tekstu, generowaniem grafiki, formatowaniem oraz doбором odpowiednich czcionek. Producent sugeruje, aby korzystać z tej funkcji jednocześnie podając czatowi odpowiednie pliki z Dysku Google, w których zawarte są wszystkie potrzebne dane a także szczegółowo sformułowane założenia projektowe np. w formie arkusza kalkulacyjnego oraz notatek ze spotkania. Gemini nie ma problemu z przyjęciem wielu plików jednocześnie! Na koniec należy dodać, że wbudowane zabezpieczenia Gemini nie są nieomyślne. Model wciąż może wygenerować dezinformację, błędy, halucynacje i przekleństwa. Proszę zawsze uważnie weryfikować jego pracę!





ptaki to fejk

Nie wierz w skrzydlate kłamstwa.

Od dawna ukrywana prawda wychodzi na jaw: ptaki nie są prawdziwe. To zaawansowane drony szpiegowskie, stworzone przez system. Ich misją jest inwigilacja i kontrola.

Nasze postulaty

Nie daj się zwieść ich piórom, to tylko kamuflaż.

Naszym celem jest ujawnienie tej spiskowej rzeczywistości i obalenie mitu o wolności ptaków.

- Demaskowanie rządowych manipulacji.
- Obalanie iluzji natury.
- Budowanie świadomości o wszechobecnej inwigilacji.

Fragment dokumentu wygenerowanego przez Gemini po podaniu jedynie tytułu oraz motywu przewodniego: cyberpunk. Mimo olbrzymiego potencjału technologii w usprawnianiu pracy, nie należy oczekiwać, że całkowicie wyręczy użytkownika.

Czy warto dokonać subskrypcji?

Zwięźle omawiając pozostałe funkcje: w Gmailu pojawiają się dodatkowe ustawienia, takie jak inteligentne sugestie poprawek – podobne do tych oferowanych przez usługę Grammarly. Po otwarciu konkretnej wiadomości e-mail lub arkusza kalkulacyjnego pojawia się możliwość podsumowania przez Gemini. Subskrypcja oferuje także funkcje podsumowywania i przetłumaczenia wiadomości w aplikacji Google Chat oraz zamianę mowy na tekst, czyli rejestrowanie spotkań w Google Meet. Jeśli użytkownik posiada dostęp do najnowszych modeli Gemini lub jednej z wersji subskrypcji Google Workspace, to najprawdopodobniej ma dostęp do omówionych funkcji. Pozostałym

użytkownikom można zalecić rozważenie zakupu pakietu, jeśli spędzają dużo czasu w aplikacjach biurowych takich jak skrzynka pocztowa, dokumenty czy arkusze kalkulacyjne, a także są tolerancyjni wobec opisanych wcześniej mankamentów.

Wizja inteligentnego asystenta działającego bezpośrednio w aplikacjach biurowych, który samodzielnie wyszukuje informacje i sugeruje poprawki z ludzką precyzją, na razie pozostaje przyszłością. ■

Karol Chruzik
student Akademii WSB

Źródła:

- [1] <https://www.reuters.com/technology/artificial-intelligence/google-invests-1-bln-openai-rival-anthropic-ft-reports-2025-01-22>
- [2] <https://bank.pl/google-bedzie-zasilal-centra-danych-ai-energia-z-malych-reaktorow-jadrowych/>

Widziałem rzeczy, którym wy, ludzie, nie dalibyście wiary. Statki szturmowe w ogniu sunące nieopodal Pasa Oriona. Oglądałem promieniowanie skrzące się w ciemnościach blisko wrót Tannhäusera. Wszystkie te chwile znikną w czasie jak lzy w deszczu.

Roy Batty w filmie „Blade Runner”

Przedmierzch, część 1

Czy *Echa* nie mają praw? Oczywiście, że nie mają!

Lalki *bunraku* płasły na scenie. Głowa jednego z tańczących mężczyzn opadła na pierś i wtedy drewniana twarz przemieniła się w maskę demona o pajęczej głowie. Bębny *taiko* zawarczały groźnie za sceną, a akompaniator grający na samisenie szarpnął struny, wydobywając z nich dramatyczne tony. *Ningyō-zukai* poruszali *postacią* w sposób perfekcyjny. Ruchy mężczyzny-demona nabrały drapieżnej precyzji.

Wiedziałem, co stanie się za chwilę. Oglądałem to już nieskończoną ilość razy. Prawdziwi ludzie potrafili zgotować piekło takim jak ja! Bo ja jestem tylko *Echem*! Wybrańcem, ale przecież nie prawdziwym człowiekiem.

Siedzę w klatce z widokiem na scenę i oglądam wciąż to samo przedstawienie. Od dziesiątków, setek, a może od tysięcy lat śledzę losy demona wodzącego samurajów na pokuszenie i znam każdą głoskę, każdą pauzę, każde westchnienie bohaterów.

Wiem, że za postaciami lalek kryją się aktorzy. Czasami jedną animuje aż trzech *ningyō-zukai*. Nigdy ich jednak nie widzę. Są ukryci.

* * *

Kaldo siedł na czele grupy.

Ten muskularny, łysy mężczyzna narzucał mordercze tempo marszu, ale Mirian wiedziała, że jeżeli chcą żyć, muszą pokonać słabość. Starła się trzymać środka stawki, chociaż nie było to łatwe.

Głęboki półmrok, w którym się poruszali, płatał wzrokowi figle. W pewnym momencie zobaczyła dwa pokraczne cienie, dwie pająkowate istoty obejmujące Kaldo. Zdawały się trzymać go za ramiona i prowadzić prosto w ciemność.

Mirian zamruwała oczami, dziwne majaki zniknęły.

Jestem zmęczona! – pomyślała.

Rany na przedramionach wciąż bolały, ale nie sączyła się już z nich zielonkawa piana, jak tuż po narodzinach w Uterusie. Rozległe owrzodzenia pokryła gruba skorupa, której powierzchnia jarzyła się chwilami mdłym, seledynowym światłem.

Grupa trzynastu ludzi maszerowała w milczeniu, starając się dotykać palcami gładkiej ściany po prawej stronie. Przeciwległa znajdowała się dwa kilosąźnie od ich lewych dłoni. Korytarz główny był nie tylko niewyobrażalnie długi, był również szeroki. Musiał taki być, aby megastruktura mogła należycie funkcjonować.

Najgorszy był jednak ów ponury półmrok. Mirian bała się ciemności. W dzieciństwie lampka przy jej łóżku świeciła się każdej nocy.

Wiedziała, że te odległe wspomnienia nie są tak naprawdę jej wspomnieniami. Była przecież *Echem*.

Podobnie jak pozostałe neotwory zorientowała się w sytuacji niemal natychmiast, gdy ocknęła się na zimnym podeście Uterusa. Popadła w rozpacz. Pozostali również. Wszyscy wiedzieli, jaki los czeka *nieproszonych gości*.

Tylko Kaldo nie stracił głowy.

To on odnalazł w opustoszałych pomieszczeniach kompleksu rodni szare tuniki. Stroje, w które obsługa Uterusa ubierała przywoływane avatary – prawdziwych ludzi.

To on uspokajał pozostałe *Echa*, mówiąc, że wrzodzące rany, sączące zieloną pianę, wkrótce zaschną.

To on dał im nadzieję, mówiąc, iż mogą uciec i zgubić się w labiryncie gigastruktury.

To on szedł przodem, jakby chciał osłonić pozostałych uciekinierów przed wszelkimi niebezpieczeństwami czyhającymi w mroku.

To on pocieszał wszystkie *Echa* podczas postojów. Mówił, że wylot jednego z bocznych korytarzy jest już niedaleko.

Mirian nie miała pojęcia, czy Kaldo naprawdę wie, gdzie znajduje się wejście do labiryntu gigastruktury, ale jego spokojny głos i skupiona, pozbawiona brwi twarz, naznaczona nieregularnymi strupami, z których sączył się zielony blask, była uosobieniem spokoju oraz pewności siebie.

* * *

Marf Witar patrzył na błyszczącą dyszę Uterusa. Wielka szarobłękitna kropla połyskiwała niczym szkło w białym świetle sączącym się ze ścian i sufitu. Zmysły podpowiadały, że w pomieszczeniu nic się nie dzieje. A jednak neotwór formował się minuta po minucie.

Dobrze, że to nie pak węglowy, ale zanim stężeje i spręgnie, miną jeszcze trzy, cztery cykle. Zdążę wypić drinka... Albo dwa – pomyślał Marf, kierując się w stronę pomieszczeń socjalnych rodni.

* * *

Scenę zalało szkarłatne światło, a postacie zamarły, przybierając groteskowe pozy.

To był znak!

Jedyny znak, na który może czekać uwięziona dusza wiecznego *Echa*. Wyrzutek odseparowany od wieczystego życia oraz realu.

I chociaż ból ponownych narodzin będzie trwał przez wiele cykli, to ostatecznie trafię do miejsca, gdzie nie będzie tego cholernego przedstawienia.

* * *

Touser, a właściwie jego *Echo*, siedział nagi, skulony na metalowym krześle, drżąc z zimna. W pomieszczeniu było niemal trzydzieści stopni, ale musiało minąć kilka cykli, zanim osobniki opuszczające trzewia Uterusa uregulują gospodarkę hormonalną i katesualną.

Witar obserwował neotwora spod przymrużonych powiek. Przybysz był dobrze zbudowany, barczysty. Pozbawiona włosów głowa wydawała się niemal idealnie okrągła.

Trzy cykle wcześniej rozmawiał z jego pierwowzorem w Symulacji 23472453. Uzgodnionego avatara Valsa Tousera nie widział nigdy.

Skóra *Echa* miała jeszcze odcień szarobłękitny, ale zaczynała się już rozjaśniać.

– Chcesz się czegoś napić? – zapytał Marf.

Osobnik wykapany niedawno z Uterusa uniósł głowę i spojrzał na niego, ale wzrok miał w dalszym ciągu przymglony.

Głowa kopii opadła po chwili, a z ust wydobyło się ciche charczenie.

Potrzebuje jeszcze trochę czasu – stwierdził Marf, lustrując uważnie ciało neotwora. Szukał jakichkolwiek oznak degeneracji charakterystycznych dla tworów o kodzie *resonare*.

Avatary były zawsze perfekcyjne. *Echa* stanowiły przeważnie karykatury pierwowzorów.

Osobnik siedzący na metalowym krześle należał do nielicznych wyjątków.

Pewnie dlatego jego świadomość jest przechowywana w *Teatrze*.

Pewnie dlatego jest Łowcą.



* * *

Kiedy odnaleźli wylot tunelu technicznego, Mirian zaczęła wierzyć, że ich ucieczka ma szansę powodzenia.

Kaldo wprowadził ich pół kilosążnia w głąb korytarza, którego ściany emitowały jasnofioletowe światło. Przy pierwszym rozwidleniu zarządził postój i odpoczynek.

* * *

Ten parszywy dupek siedzący naprzeciw mnie najwyraźniej czuł się lepszy. A może tylko ważniejszy?

– Możemy już rozmawiać? – zapytał.

– Możemy – powiedziałem bez entuzjazmu.

Świat materialny miał swoje wady. Tym razem przywitał mnie bolesnymi ranami w jamie ustnej. Uterus nie działał perfekcyjnie. Prawie nigdy! Miałem okazję wypróbować to już wielokrotnie.

Na szczęście dla mnie, za pierwszym razem okazał się idealny i dlatego ja, parszywe *Echo*, stałem się Łowcą.

To było dawno temu. Naprawdę dawno. Polowałem na bękarty Uterusa dłużej, niż zwykły trwać religie na niezliczonych światach, które poznał człowiek.

– Nazywam się Marf Witar... Twoje imię znam – oznajmił dupek.

– Zatem *savoir-vivre* już za nami – mruknąłem.

Nie zwrócił uwagi na mój pogardliwy ton i rozpoczął monolog:

– Multum gigastruktury podjęło decyzję o wyprawie do Wszechświata H. Golidian su-prze-strzenny wystartował dwieście sześćdziesiąt cykli temu. Dwanaście avatarów załogi w trybie *Tytan* zostało stworzonych siedemdziesiąt cykli wcześniej w sektorze 66743552. Może oni otworzą bramę Tannhäusera. Niestety pojawiły się problemy. Chociaż jesteśmy oddaleni od tamtego miejsca, *Echa* pojawiły się właśnie u nas.

Doskonale wiedziałem, co powie za chwilę.

Mechanizm był za każdym razem taki sam. Kilkadziesiąt cykli po stworzeniu neotworów, w materialnym świecie pojawiały się *Echa*. Pojawiały się losowo w jednym z setek tysięcy sektorów gigastruktury. Miały dokładnie tę samą świadomość, co zrodzone wcześniej avatary. Jażn wyszarpana spośród milionów wirtuali, wytypowana spośród biliardów ludzi, ulegała podwojeniu.

Echa miały najczęściej wady fizyczne. Niektóre umierały natychmiast po wykapaniu z dyszy, inne, mniej upośledzone, orientowały się od razu, że nie są oryginałami. Los nadmiarowych kopii mógł być tylko jeden. Ale rzadko kto skłonny jest oddać życie za darmo.

Uciekali więc. Wydawało się im, że gigastruktura daje nieskończenie wiele możliwości.

Złudzenie!

Każdy z korytarzy głównych liczył 939 887 974 kilosążnie. Od każdego odchodziły miliony korytarzy serwisowo-technicznych, które tworzyły labirynt, niemożliwy do objęcia przez ludzki rozum.

Aby dotrzeć do jednego z ośmiu korytarzy głównych, potrzeba było na ogół niewiele czasu. Uterusy znajdowały się przeważnie w ich pobliżu, aby ułatwić komunikację z punktami węzłowymi, w których znajdowały się doki Golidianów su-prze-strzennych. Ale niekończąca się arteria komunikacyjna o szerokości dwóch kilosążni i takiej samej wysokości nie mogła zapewnić schronienia.

Echa zdawały sobie z tego sprawę. Wiedziały, że za wszelką cenę muszą dotrzeć do jednego z bocznych korytarzy serwisowych. Tymczasem ich wyloty oddalone były od siebie o minimum sto kaesów.

Czasami znajdowały się blisko Uterusów, czasami wręcz przeciwnie.

Nie miało to zresztą większego znaczenia. Wiedziałem, jak ich znaleźć. Mogli się zaszyć w dowolnej części labiryntu, ale nie potrafili skutecznie zatrzeć śladów, które pozostawiali. To była nić, którą pracowicie związałem, aby dotrzeć do kłębka i zdymisjonować *Echa*. ■

Marek Żelkowski

(druga część opowiadania w następnym numerze „Młodego Technika”)

Terraform – uniwersalne podejście do IT

Jeden program do (prawie) wszystkiego

Niektórzy w branży informatycznej żartobliwie mówią, że jeśli mielibyśmy zaatakować infrastrukturę obcych podczas inwazji na Ziemię jak w filmie „Dzień Niepodległości”, to najlepiej pomógłby nam w tym Terraform **(1)**. Na tyle jest uniwersalny, wszechstronny i kosmicznie użyteczny.

Definicyjnie Terraform to oprogramowanie stworzone przez firmę HashiCorp do tworzenia i zarządzania infrastrukturą jako kodem (Infrastructure as Code, IaC). Umożliwia bezpieczne, wydajne kreowanie i zmiany w wersjach zasobów lokalnych i w chmurze, w plikach konfiguracyjnych, które można ponownie wykorzystywać, udostępniać, kreując także kolejne wersje. Celem jest zapewnienie płynności i spójności zarządzania całą infrastrukturą IT przez cały cykl jej życia. Terraform można wykorzystać do zarządzania komponentami niskiego poziomu, takimi jak zasoby obliczeniowe, pamięci masowe i sieciowe, a także komponentami wysokiego poziomu, takimi jak wpisy DNS i funkcje SaaS (oprogramowania jako usługi).

Entuzjaści tego rozwiązania, a ich liczba rośnie, wskazują, że Terraform znacznie ułatwia zarządzanie infrastrukturą lokalną i w chmurze, a także na to, że można go łatwo rozszerzać za pomocą architektury opartej na wtyczkach. Służy również do łączenia z różnymi hostami infrastruktury i konfiguracji złożonych scenariuszy zarządzania w wielu chmurach. Jego konfigurację można łatwo spakować, udostępnić i ponownie wykorzystać w postaci modułów Terraform.

Opisać cele zamiast szczegółowych rozkazów

W projektach polegających na przygotowywaniu serwerów o dużej skali dla setek lub konfiguracji obejmującej tysiące serwerów lawinowo narastają komplikacje i podatności na błędy rozwiązania IaC stają się nieocenione, umożliwiając definiowanie i dostarczanie infrastruktury za pomocą programowania deklaratywnego, które w przeciwieństwie do programów napisanych imperatywnie polega na opisywaniu przez programistę warunków, jakie musi spełniać końcowe



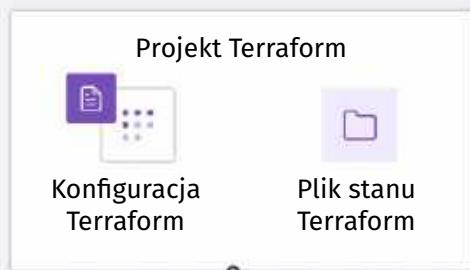
1. Logo Terraform

rozwiązanie (co chcemy osiągnąć), a nie szczegółową sekwencję kroków, które do niego prowadzą (jak to zrobić). Dzięki takiemu podejściu można uzyskać spójność, powtarzalność i zwiększoną niezawodność. Podejście to pomaga również w zarządzaniu wszelkimi zmianami, które zachodzą w infrastrukturze po drodze, powszechnie określanymi jako zarządzanie dryfem.

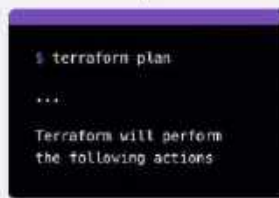
Pliki konfiguracyjne Terraform są deklaratywne, co oznacza, że opisują stan końcowy infrastruktury. Nie trzeba pisać instrukcji krok po kroku, by tworzyć zasoby. Terraform buduje graf zasobów w celu określenia zależności zasobów i tworzy lub modyfikuje niezależne zasoby równolegle. Pozwala to na wydajne udostępnianie zasobów. Terraform obsługuje komponenty konfiguracyjne wielokrotnego użytku zwane modułami, które definiują konfigurowalne zbiory infrastruktury, oszczędzając czas. Można korzystać z publicznie dostępnych modułów z rejestru Terraform lub pisać własne.

Write

Definiowanie infrastruktury w plikach konfiguracyjnych

**Plan**

Przejrzyj zmiany, jakie Terraform wprowadzi w twojej infrastrukturze

**Apply**

Terraform ustanawia infrastrukturę i aktualizuje plik stanu

**2. Poglądowy model Terraform**

Terraform tworzy i zarządza zasobami na platformach chmurowych i w innych usługach za pośrednictwem ich interfejsów programowania aplikacji (API). Dostawcy umożliwiają Terraform pracę z praktycznie każdą platformą lub usługą z dostępnym interfejsem API. W rejestrze Terraform można znaleźć wszystkich znanych dostawców rozwiązań chmurowych, w tym Amazon Web Services (AWS), Azure, Google Cloud Platform (GCP), Kubernetes, Helm, GitHub, Splunk, DataDog i wielu innych, mniej znanych.

Terraform jest też narzędziem typu open source, które kompiluje się do pojedynczego pliku binarnego w systemie i używania z wiersza poleceń. Infrastruktura jest zdefiniowana w języku Hashicorp Configuration Language (HCL), języku deklaratywnym, przeznaczonym dla narzędzi i serwerów.

Komponenty infrastruktury zarządzane przez Terraform są nazywane zasobami. Przykłady zasobów to np. maszyny wirtualne, tabele baz danych, obiekty (tzw. wiadra) AWS S3 i inne. Każdy blok zasobów w HCL pomaga zdefiniować zasób i skonfigurować jego właściwości.

Zasoby, plan i zastosowanie

W opisach działania Terraform znajdujemy podział przepływu pracy na zasadnicze trzy etapy (2):

– Pierwszy oznaczany jest angielskim słowem „write” (od „pisać”, „zapisywać”). Definiuje się w nim zasoby, które mogą znajdować się u wielu dostawców chmury i usług. Można na przykład utworzyć konfigurację do wdrożenia aplikacji na maszynach wirtualnych w sieci Virtual Private Cloud (VPC) z grupami zabezpieczeń i load balancerem, czyli urządzeniem lub

oprogramowaniem, które rozdziela ruch sieciowy pomiędzy serwery.

- Kolejny etap w Terraform to planowanie (ang. „plan”). Platforma tworzy plan wykonania opisujący infrastrukturę, którą utworzy, zaktualizuje lub zniszczy w oparciu o istniejącą infrastrukturę i konfigurację.

- Trzecia faza to „zastosowanie” (ang. „apply”). Po zatwierdzeniu Terraform wykonuje proponowane operacje we właściwej kolejności, z zachowaniem wszelkich zależności między zasobami. Na przykład, jeśli zaktualizujesz właściwości VPC i zmienisz liczbę maszyn wirtualnych w tym VPC, Terraform ponownie utworzy VPC przed skalowaniem maszyn wirtualnych.

Użytkownicy definiują i dostarczają infrastrukturę centrum danych za pomocą języka konfiguracji HCL lub opcjonalnie JSON. Terraform generuje plan i wyświetla monit o zatwierdzenie przed modyfikacją infrastruktury. Śledzi również rzeczywistą infrastrukturę w pliku stanu, używając go do określenia zmian, które należy wprowadzić w infrastrukturze, aby była zgodna z konfiguracją. Proponuje również plan dodawania lub usuwania komponentów infrastruktury w razie potrzeby. Zajmuje się udostępnianiem lub wycofywaniem wszelkich zasobów, jeśli użytkownik zdecyduje się zastosować plan. Wtyczki Terraform, takie jak Terraform Providers i Provisioners, zapewniają mechanizm komunikacji z hostem infrastruktury lub dostawcami SaaS. Rdzeń Terraform komunikuje się z wtyczkami za pośrednictwem zdalnego wywołania procedur (RPC).

Ponieważ konfiguracja jest zapisana w pliku, można ją zatwierdzić w systemie kontroli wersji i używać narzędzia do efektywnego zarządzania pracą w różnych zespołach. HCP Terraform uruchamia Terraform w spójnym, sprawdzonym środowisku i zapewnia bezpieczny dostęp do współdzielonego stanu i poufnych danych, kontrolę dostępu opartą na rolach, prywatny rejestr do udostępniania zarówno modułów, jak i dostawców i nie tylko.

Otwarta alternatywa

Tworząca narzędzie firma HashiCorp, Inc. ma siedzibę w San Francisco w Kalifornii. Została założona w 2012 roku przez Mitchella Hashimoto i Armona Dadgara. Nazwa firmy HashiCorp to połączenie nazwiska współzałożyciela Hashimoto i Corporation. Hashimoto pracował wcześniej nad oprogramowaniem typu open source o nazwie Vagrant, które zostało włączone do zestawu produktów HashiCorp.

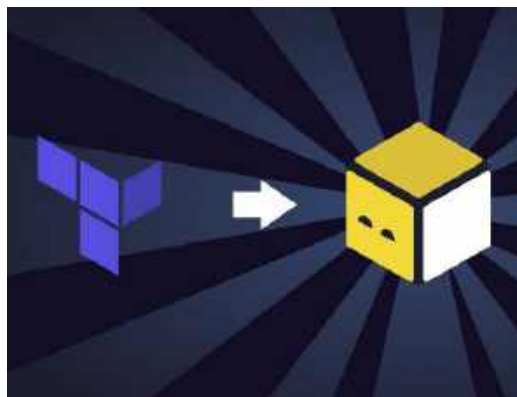
Choć Terraform został po raz pierwszy wydany w 2014 r., firma HashiCorp uruchomiła rejestr

modułów Terraform trzy lata później. W 2019 r. wprowadzono płatną wersję Terraform Enterprise. Terraform był wcześniej wolnym oprogramowaniem dostępnym na licencji Mozilla Public License (MPL) w wersji 2.0. W dniu 10 sierpnia 2023 r. HashiCorp ogłosiła, że wszystkie produkty wytwarzane przez firmę będą licencjonowane na licencji Business Source License (BUSL), przy czym HashiCorp zakazuje komercyjnego wykorzystania edycji społecznościowej przez tych, którzy oferują „konkurencyjne usługi”. Mitchell Hashimoto zrezygnował z pracy w firmie kilka miesięcy później. A w kwietniu 2024 r. spółka ogłosiła, że zawarła umowę przejęcia przez IBM za 6,4 mld USD, przy czym transakcja miała zostać zamknięta do końca ubiegłego roku. Stało się to faktem, z pewnym opóźnieniem, bo w lutym 2025 r.

Choć Terraform jest chwalony w środowisku programistycznym jako wyjątkowe narzędzie o przełomowym, uwalniającym zarządzanie zasobami informatycznymi od tradycyjnych ograniczeń, charakterze, warto podkreślić, że obecnie nie jest to jedynie rozwiązanie typu IaC. Dostępne są na rynku także inne rozwiązania, oferowane przeważnie przez potentatów Big Tech, na przykład szablony Cloud Formation Templates (CFT) dla AWS, Azure Resource Manager (ARM) dla Microsoft Azure i Deployment Manager dla chmury Google.

Powstała też alternatywna wersja Terraform o otwartym kodzie źródłowym o nazwie OpenTofu (3), będąca boczną odnogą rozwoju Terraform w wersji 1.5.6. Powstała jako inicjatywa społecznościowa wspierana przez wiele najbardziej uznanych firm i projektów w ekosystemie Terraform w odpowiedzi na zmianę licencji BSL firmy HashiCorp. Nie ma żadnych ograniczeń dotyczących sposobu korzystania z OpenTofu, zarówno do użytku komercyjnego, jak i osobistego. ■

Mirosław Usidus



3. Graficzne zobrazowania ewolucji Terraform do OpenTofu

Rozmowa z Erikiem Wernquistem, artystą grafiką i animatorem, twórcą wizyjnych filmów o tematyce science fiction.

Nasza romantyczna przyszłość w kosmosie

Młody Technik: Patrząc na Pańskie prace, odnosi się wrażenie, że patrzy Pan w przyszłość ludzkości w sumie optymistycznie. Widać w nich wiarę w to, że pokonamy bariery techniczne w dziedzinie podboju kosmosu, że dokonamy przełomów naukowych. Czy dobrze rozumiem Pańską twórczość? Jeśli tak, to skąd Pan czerpie wiarę, że staniemy się cywilizacją międzyplanetarną?

Erik Wernquist: Nie chodzi tu tyle o wiarę, co po prostu o ekstrapolację trajektorii postępu technologicznego i tempa wzrostu dobrobytu, które nasz gatunek potrafi zapewnić, jak udowodnił de tej pory. Wiele można powiedzieć o naszym obecnym stanie i można by argumentować, że trajektoria może zmienić kierunek w dowolnym momencie, ale byłoby to prawdą dla każdego momentu w historii, a jeśli ktoś pozwoli sobie na wystarczająco dużą perspektywę, powiedziałbym, że powoli, ale stale podążamy ścieżką, która ostatecznie zabierze niektórych z nas z naszego świata gdzie indziej. Ponownie może się to oczywiście zmienić w dowolnym momencie, ale parafrazując Carla Sagana: jeśli nasz gatunek przetrwa wystarczająco długo i nie



Erik Wernquist jest artystą cyfrowym, animatorem i twórcą filmowym mieszkającym w Sztokholmie, w Szwecji. Ponad dwadzieścia lat temu zasłynął ze stworzenia postaci Crazy Frog, będącej bohaterem krótkiego filmu, dzwonka do telefonu komórkowego, gry komputerowej i teledysku. Pracował jako scenarzysta, reżyser i producent sztuk teatralnych, zanim w 2000 roku nauczył się grafiki komputerowej. Przez kilka lat pracował jako niezależny artysta/ilustrator 3D, a w 2004 roku rozpoczął pracę jako animator postaci dla Kaktus Film. W 2007 roku zrobił sobie przerwę w pracy w animacji, aby współtworzyć szwedzki serial komediowy „Grottesco”, reżyserując trzy z ośmiu odcinków tej produkcji. W 2011 roku współtworzył scenariusz, współreżyserował i pracował nad efektami wizualnymi szwedzkiego serialu telewizyjnego „Pulver”. W 2014 roku Wernquist stworzył film pt. „Wanderers”, który przedstawia lokalizacje w Układzie Słonecznym eksplorowane przez ludzkich odkrywców, wspomaganą przez hipotetyczną technologię kosmiczną przyszłości. Artysta stworzył niektóre ze scen przy użyciu wyłącznie grafiki komputerowej, ale większość opiera się na rzeczywistych zdjęciach wykonanych przez sondy kosmiczne lub łaziki w połączeniu z dodatkowymi elementami wygenerowanymi komputerowo. Narratorem „Wanderers” jest astronom Carl Sagan, czytający fragmenty swojej książki „Błękitna kropka. Człowiek i jego przyszłość w kosmosie” z 1994 roku. Drugim znanym dziełem o tematyce science fiction autorstwa Wernquista jest „One Revolution Per Minute”, krótkometrażowy film, który stworzył, jak sam pisze na swojej stronie: „aby dać wyraz własnej fascynacji sztuczną grawitacją w kosmosie”. Na filmie oglądamy widoki wewnątrz i na zewnątrz stacji kosmicznej „SSPO Esperanta”, która obraca się z prędkością jednego obrotu na minutę. Przy promieniu 450 metrów obrót generuje sztuczną grawitację z efektem około 0,5 g na głównym pokładzie.



1. Kadr z filmu Erika Wernquista 'One Revolution Per Minute'

zniszczyliśmy samych siebie, nieuchronnie wyruszymy do gwiazd.

MT: Skąd taka, a nie inna tematyka? Jak się rodzi w artyście idea, by skupić się na wizjach przyszłości człowieka w dalekim kosmosie?

EW: Cóż, w rzeczywistości zajmuję się innymi tematami, ale trzeba przyznać, że znalazłem coś w rodzaju niszy w wizualizacjach kosmicznych. Uważam się za wielkiego szczęściarza, że mogę się tym zajmować jako artysta, ponieważ eksploracja kosmosu, a w szczególności eksploracja planet, to kwestie, które zawsze mnie pasjonowały.

MT: Wprawdzie nasze czasopismo publikuje krótkie formy literackie z gatunku science fiction, by pobudzić wyobraźnię naszych czytelników, to jednak główną jego tematyką są postępy w dziedzinach nauki i technologii. Czy interesuje się Pan tą tematyką na bieżąco? Czy uważa Pan, że jesteśmy u progu wielkich przełomów w nauce, technice, medycynie? Niektórzy uważają, że przyniesie to rewolucja AI. Jak jest Pańskie zdanie na te tematy?

EW: Moje zainteresowanie granicami obecnej nauki i technologii jest mimo wszystko dość ograniczone. Staram się śledzić to, kiedy mam czas, ale jednocześnie skłaniam się raczej do ostrożnego oczekiwania i wyobrażania sobie tego, co obecny rozwój może dla nas oznaczać w dłuższej perspektywie. Bardzo trudno jest oszacować, co jest prawdziwym „przełomem” w jakimkolwiek znaczącym sensie, będąc jedynie świadkiem jego osiągnięcia.

Myślę, że musimy dać sobie trochę czasu, aby przekonać się, jak (i czy) każda nowa technologia wpłynie na nas i czy stanie się to w jakimkolwiek znaczący sposób.

MT: Weszliśmy na obszar tematyki AI. Nie mogę nie zapytać czynnego twórcy, jaką ma opinię na temat

ekspansji generatywnych technik tworzenia obrazu i wideo? Jest wiele negatywnych reakcji w świecie twórców, od ruchu obrony ich praw autorskich do dzieł, które używane są do szkolenia generatywnych modeli AI, po kwestionowanie przyszłości sztuki. Jakie są Pańskie przemyślenia, jeśli chodzi o ten temat?

EW: Jak do tej pory nie wypróbowałem jeszcze żadnego z nowych narzędzi generatywnej sztucznej inteligencji. Jak już wspomniałem, bardzo trudno jest mi przewidzieć, w jaki sposób te wszystkie nowe rozwiązania i techniczne możliwości mogą być wykorzystywane w przyszłości lub co będą oznaczać w dłuższej perspektywie.

Z pewnością wydają się one mieć obecnie destrukcyjny wpływ, mogąc zarazem zrewolucjonizować sposób tworzenia obrazów i wideo. Jednocześnie mam przecucie, że gdy szum nowości opadnie, wszelkie przydatne aspekty zostaną zaadaptowane lub zintegrowane jako narzędzia w przyborniku wykorzystywanym przez artystów. Inne z kolei mogą być dalej rozwijane w innych dziedzinach lub odrzucone jako bezużyteczne.

Pomijając wątpliwe etyczne metody wykorzystywane do gromadzenia źródłowych banków danych używanych do szkolenia tych modeli sztucznej inteligencji, muszę przyznać, że technologia stojąca za tym wszystkim jest imponująca i bez wątpienia ciekawie będzie zobaczyć, dokąd może to doprowadzić w przyszłości.

Chciałbym jednak dodać, że osobiście bardzo trudno jest mi znaleźć jakiegokolwiek artystyczny urok nawet w ogólnej idei sztucznie generowanych dzieł sztuki. Istotą sztuki, przynajmniej w moim odczuciu, jest komunikacja. Aby komunikacja miała jakiegokolwiek



2. Kadr z filmu Erika Wernquista 'One Revolution Per Minute'

znaczenie, musi istnieć nadawca i odbiorca, ponieważ to między tymi dwoma podmiotami może zachodzić komunikacja. Jeśli nie ma nadawcy, to wszystko, cokolwiek zostanie odebrane, jest zasadniczo pozbawione znaczenia, bez względu na to, jak imponująca technicznie lub wizualnie jest treść przekazu. Uważam, że dotyczy to również sztucznie generowanego tekstu.

MT: Pańskie kreacje są olśniewającymi wizjami, w których umieszcza Pan człowieka, który często wydaje się czymś znikomym wobec zjawisk, które widać na obrazach, ale wciąż jest człowiekiem, takim jak Pan i ja. W świecie nauki jest sporo wątpliwości, czy tak krucha i mocno przywiązana do swojej planety biologicznie istota nadaje się do przebywania w kosmosie poza swoją macierzystą planetą. Przynajmniej w takiej wersji, jaką znamy i jaką teraz jesteśmy. Pisaliśmy o tym w naszym magazynie. Pojawiają się koncepcje, że człowiek będzie musiał się zmienić do kosmosu, zmutować lub zmodyfikować się genetycznie, przystosowując się do np. promieniowania i innych ekstremalnych okoliczności. Są też opinie, że kosmos powinien eksplorować maszyny, bo człowiek się do tego nie nadaje. Być może zgodnie z wizją osobliwości byłyby to maszyny z ludzkim umysłem, świadomością. Co Pan o tym wszystkim myśli?

EW: Jeśli chodzi o moje osobiste filmy krótkometrażowe „Wanderers” i „One Revolution Per Minute”, to tak, rzeczywiście przedstawiają one bardzo romantyczny obraz ludzkiej obecności w kosmosie. Doskonale zdaję sobie sprawę z wielu zagrożeń, z którymi musielibyśmy się zmierzyć, gdybyśmy mieli odbyć te podróże w rzeczywistości, i zakładam, że istnieje wiele innych zagrożeń, których nie jestem świadomy. Ale moja perspektywa jest, i z pewnością zawsze będzie,

ludzka i właśnie tym chcę się dzielić z odbiorcami i to odkrywać w tych krótkich filmach – ludzką perspektywę w głębinach kosmosu i na sąsiadujących z naszym światach w Układzie Słonecznym. Aby coś takiego osiągnąć, staram się umieszczać w wizualizacjach coś, co można odnieść do człowieka, może to być filiżanka kawy, pojazd, miasto lub po prostu osoba ludzka.

Może okazać się prawdą, że nie nadajemy się do dłuższych podróży w kosmos, nie mówiąc już o jakiegokolwiek kolonizacji kosmosu na dużą skalę, ale moje filmy krótkometrażowe nie mają być traktowane jako przepisy na naszą przyszłość lub „wezwanie do działania” w tym zakresie – są bardziej jak pełne nadziei marzenia.

Jeśli chodzi o spekulacje na temat naszej dalekiej przyszłości, czy to w postaci ewolucyjnych potomków, czy zintegrowanych z technologią, wszystko to może być fascynujące, ale co do mnie, to ja osobiście tracę zainteresowanie tematem, gdy zaczynam oddalać się od tego, co można odnieść do człowieka.

MT: Już niejednemu raz w ostatniej dekadzie przedstawiciele NASA i inne autorytety twierdziły, że już wkrótce, do 2030 lub 2040 roku odkryjemy życie pozaziemskie. Jak mogłoby to zmienić nas, nasze spojrzenie na Wszechświat, kierunki i tempo rozwoju naukowo-technicznego?

EW: Nie wiem, czy te prognozy są prawdziwe, ale z pewnością zweryfikowane potwierdzenie istnienia życia pozaziemskiego byłoby niezwykłym odkryciem. Nie wiem, jak wiele by to faktycznie zmieniło, ale potwierdzenie, że życie istnieje gdzie indziej we Wszechświecie, miałoby oczywiście ogromne znaczenie naukowe. Podejrzewam jednak, że wielu spodziewa się, że tak jest [że istnieje życie poza Ziemią – red.], tylko jest to bardzo trudne do zaobserwowania



3. Kadr z filmu Erika Wernquista 'One Revolution Per Minute'

i udowodnienia w sposób naukowy i niebudzący wątpliwości.

MT: Pańskie dzieła często są połączeniem, czymś w rodzaju montażu obrazów Układu Słonecznego, które znamy z misji kosmicznych NASA czy ESA, z elementami dodanymi przez Pana, które są swoistymi sugestiami, jak mogłaby wyglądać ludzka obecność w tych miejscach. Mam takie trochę może techniczne pytanie – co było pierwsze – czy widząc zdjęcia nadesłane przez sondy, pomyślał Pan, że warto dodać do nich człowieka i jego aktywność, czy może było tak, że miał Pan wcześniej pomysł, by pokazać eksplorację kosmosu przez człowieka, a te obrazy z misji po prostu podsunęły Panu tło do tych kreacji?

EW: Przeważnie chcę pokazać pewne miejsce lub cechę, a następnie dodaję elementy ludzkie, aby nadać

temu odpowiednią perspektywę. Najczęściej więc miejsce jest na pierwszym miejscu, a ludzie na drugim.

MT: Na potrzeby wizualizacji dla profesora Coxa stworzył Pan obraz „świata w cylindrze O'Neilla”, który nasuwa m.in. skojarzenia z filmem „Interstellar”. Nie wiem, czy zna Pan ten film, ale przedstawiono w nim wielki przełom naukowy ludzkiej cywilizacji, który stał się możliwy, jednak dopiero po kontakcie z superinteligentnymi obcymi. Czy podziela Pan optymizm takiego spojrzenia na kontakt człowieka z wysoko rozwiniętą cywilizacją? Czy jednak powinniśmy przede wszystkim liczyć na samych siebie, a na obcych uważać?

EW: Wierzę, że jeśli kiedykolwiek nawiążemy kontakt z życiem pozaziemskim lub zaobserwujemy jakieś dowody na jego istnienie, prawdopodobnie

4. Kadr z filmu Erika Wernquista 'Wanderers'





5. Kadr z filmu Erika Wernquista 'Wanderers'

będzie ono w bardzo prymitywnej formie. W związku z tym prawdopodobnie mądrze byłoby zachować ostrożność w przypadku jakiegokolwiek kontaktu fizycznego, jeśli w ogóle jest to możliwe, ale ostrożność, o której, podejrzewam, jest mowa w pytaniu, nie jest czymś, co moim zdaniem musielibyśmy brać pod uwagę. Inteligentne życie, jakie znamy, zakładam, że jest niezwykle rzadkie, zarówno w przestrzeni, jak i w czasie. A jeśli coś takiego istnieje, istnieje prawdopodobieństwo, że ogromne odległości w przestrzeni i czasie na zawsze nas odizolują.

Tak czy inaczej uważam, że powinniśmy liczyć przede wszystkim na siebie. Jeśli chodzi o mnie, to zawsze uważałem „obcych” w filmie „Interstellar” za potomków nas samych, którzy wyciągają do nas rękę z przyszłości...

MT: W kierunku jowiszowego księżycy Europa lecą właśnie dwie misje, europejska i amerykańska. W NASA trwają dyskusje, w które miejsce Układu Słonecznego wysłać sondy, czy do układu Neptuna, czy Urana, a może na Enceladusa, księżyc Saturna. Z powodu ograniczonego budżetu nie można wysłać statków wszędzie. Jednocześnie pojawił się spór, czy nadać priorytet misji od razu na Marsa, czy jednak kontynuować to według planu Artemis, czyli stworzyć bazy na powierzchni i orbicie Księżyca. Jest więc poczucie, że znajdujemy się na swoistych kosmicznych rozstajach i nadchodzi czas ostatecznych decyzji. Gdyby Pan był doradcą ludzi, którzy podejmują te decyzje, co by Pan radził? Którym misjom, w Pańskiej ocenie, należy się priorytet?

EW: Jestem bardzo szczęśliwy, że nie jestem takim doradcą, ponieważ jest to rodzaj odpowiedzialności, której nie chciałbym mieć. Byłoby bardzo trudno, myślę, że wręcz niemożliwe, dostroić i zrównoważyć wszystkie parametry niezbędne do podejmowania tego typu decyzji w oparciu o pewnego rodzaju uniwersalną, obiektywną korzyść. Ostatecznie wiele z tego z pewnością sprowadzi się do subiektywnych preferencji i zainteresowań.

Osobiście jednak chciałbym zobaczyć dwie misje orbiterów podobnych do sondy Cassini do zewnętrznych lodowych olbrzymów Urana i Neptuna. Te systemy są nam prawie zupełnie nieznane, każdy z nich miał tylko jeden przelot Voyagera, a biorąc pod uwagę skarby wiedzy, jaką Cassini przekazał nam na temat systemu Saturna, można sobie tylko wyobrazić, jakie cuda zostałyby odkryte przy podobnej do misji Cassini naszej obecności w układach ciał krążących wokół tych odległych światów...

MT: Istnieją zjawiska w kulturze popularnej, które inspirują pewne wizje na nadchodzące dekady. Dotyczy to w szczególności kinematografii. Tak było w przypadku filmu „2001: Odyseja kosmiczna” wyreżyserowanego przez Stanleya Kubricka. Jego obraz ukształtował styl futurystycznych wnętrz statków kosmicznych i stacji, później odzwierciedlony w innych produkcjach science fiction. Kubrick stworzył estetykę, która wykraczała poza normy filmowe tamtych czasów i przez długi czas kształtowała sposób, w jaki wyobrażano sobie przyszłość w kinie. Czy



6. Kadr z filmu Erika Wernquista 'Wanderers'

ma Pan jakieś własne artystyczne inspiracje zaczerpnięte z filmów lub seriali science fiction?

EW: To może nie być odpowiedź na pytanie, ale największą zagadką dla mnie dotyczącą spuścizny „2001: Odysei kosmicznej” Kubricka jest to, że niezwykle praca nad wizualizacją i radzeniem sobie z koncepcją sztucznej grawitacji jest nadal tak rzadko podejmowana od tamtej pory. Owszem, w późniejszych latach widziałem tego więcej, ale nadal zadziwia mnie to, jak źle sztuczna grawitacja była przedstawiana w filmach science fiction o kosmosie, i to długo po tym, jak Kubrick wraz ze swoim zespołem pokazali, jak można to zrobić. A technicznie i praktycznie jest to dziś o wiele łatwiejsze do osiągnięcia.

MT: W swoich pracach odwołuje się Pan do różnych koncepcji technicznych, które wielokrotnie pojawiały się w książkach science fiction. Tak było na przykład w przypadku cylindra O’Neilla, którego używał Arhur C. Clarke w swojej powieści „Ranzdevous with Rama”. Chętnie odwołuje się Pan również do słów Carla Sagana, który oprócz tego, że był wspaniałym naukowcem i edukatorem, był również autorem literatury science fiction. Gdyby miał Pan zilustrować swoimi animacjami jakieś znane dzieło osadzone w tym gatunku, co by Pan wybrał i dlaczego?

EW: Osobiście chciałbym zobaczyć więcej wizualizacji i ilustracji dzieł Kima Stanleya Robinsona i Iaina M. Banksa. Czerpię od nich wiele inspiracji w tym, co robię, i uważam za zaskakujące, że nie ma jeszcze filmów, ani seriali telewizyjnych opartych na ich dziełach.

Robinson jest po prostu niezwykle pod względem rozległych ram badawczych, które są widoczne w jego powieściach, oraz głębi i złożoności człowieczeństwa, jakie odmalowane są w jego postaciach! A Banksa uważam za niesamowitego człowieka, z jego wspaniałymi pomysłami, które pozornie można przekazać bez wysiłku. Cenię i lubię go również za bardzo radosny i pełny przygód ducha jego opowieści.

MT: W swoich animacjach lubi Pan przedstawiać ciała niebieskie naszego Układu Słonecznego, które, choć wciąż pozostają poza zasięgiem człowieka, coraz lepiej poznajemy dzięki misjom sond kosmicznych. Jednak równolegle odkrywamy mnóstwo planet pozasłonecznych. Dzięki instrumentom, takim jak Kosmiczny Teleskop Keplera, a teraz Kosmiczny Teleskop Jamesa Webba, wiemy już o tysiącach egzoplanet. Niektóre z nich mogą przypominać Ziemię. W przeciwieństwie do Księżyca, Marsa czy księżyców Jowisza i Saturna, te światy pozostają jedynie danymi na mapach. Możemy spekulować na temat ich wyglądu, ale jest prawdopodobne, że przyszłe pokolenia nie zobaczą ich na własne oczy. Jak odbiera Pan tę perspektywę? Amerykański ilustrator i malarz Chesley Bonestell, znany z wizualizacji kosmosu, malował krajobrazy planet pozasłonecznych, zanim jeszcze odkryto pierwszą z nich. Jego prace były wynikiem czystej spekulacji, inspirowanej nauką i wyobraźnią. Jak daleko sięgnęłaby Pańska wyobraźnia, gdyby miał Pan przedstawić takie światy?

EW: Chesley Bonestell jest dla mnie ogromną inspiracją i chciałbym mieć więcej jego umiejętności



7. Kadr z filmu Erika Wernquista 'Wanderers'

tworzenia wizualnych pomysłów z samych danych. Do pewnego stopnia to robię, ale zdecydowanie pomaga mi także posiadanie jakiegoś odniesienia wizualnego, szczególnie jeśli chodzi o kolory wizualizacji.

Nie lubię, gdy muszę używać zbyt dużo wyobraźni, tworząc wizualizacje kosmosu, dlatego myślę, że wolę pozostać w obrębie Układu Słonecznego. Każdy element wizualny, który tworzę, tak czy inaczej wymaga pewnej dozy wyobraźni, ale lubię, gdy mam wystarczająco dużo danych i punktów odniesienia, aby było możliwe ustalenie wystarczająco solidnych granic, których wiem, że nie mogę przekroczyć.

MT: Na koniec pytanie może nieco podchwytliwe, ale skoro mówię, że jest podchwytliwe, to już zapewne takie nie jest. Pańskie dzieła nie zawierają datowania przedstawianych na nich wydarzeń. Nie zna Pan tych dat, czy Pan zapomniał dodać? A jeśli to drugie, to proszę podać choć orientacyjny przedział czasowy.

EW: Hm... Nie myślałem o tym zbyt wiele i chyba nie jest to dla mnie zbyt ważne. Jeśli weźmiemy

na przykład moje krótkie filmy „Wanderers” i „One Revolution Per Minute”, wszystko, co tam jest przedstawione, teoretycznie mogłoby się wydarzyć w dowolnym miejscu między dziś a kilkoma setkami lat w przyszłości... Oczywiście, ludzie jeszcze nie odwiedzili Marsa ani zewnętrznego Układu Słonecznego, nie wydrążyliśmy asteroid ani nie zbudowaliśmy cylindrów O’Neilla, ale mamy już całą niezbędną wiedzę i technologię, aby to zrobić.

Jak powiedziałem wcześniej; te krótkie filmy nie mają być wizjami jakiegoś konkretnego czasu w przyszłości – widzę je raczej jako studia przypadków tego, jak niektóre światy w Układzie Słonecznym wyglądałyby dla nas, gdybyśmy je odwiedzili.

MT: Dziękujemy za rozmowę. ■

Rozmawiał Miroslaw Usidus
i Polska Fundacja Fantastyki Naukowej

Zdjęcia pochodzą z archiwum kadrów z filmów Erika Wernquista. Publikowane są za pozwoleniem autora.



Rozmowa z Erikiem Wernquistem przeprowadzona została w partnerstwie z Polską Fundacją Fantastyki Naukowej, od której pochodzi część pytań zadanych rozmówcy „Młodego Technika”.

PRZEGLĄD

ELEKTROTECHNICZNY Nowy system trakcji elektrycznej

(...) W związku ze studiami nad elektryfikacją węgierskich kolei państwowych przeprowadza się od roku próby z lokomotywą elektryczną, zbudowaną podług planów inż. Kando. Lokomotywa ta została zbudowana na zasadach, odmiennych od stosowanych dotychczas w różnych systemach trakcji. Ideą przewodnią nowego systemu było umożliwienie kolejom elektrycznym bezpośredniego korzystania z ogólną krajowych sieci trójfazowych 50 okresowych bez stosowania przetwarzania prądu lub ilości okresów. Z silników prądu zmiennego stosuje się na 50 okresów tylko silniki trójfazowe, które pozatem dzięki prostocie i solidności konstrukcji nadają się doskonale do warunków pracy na kolejach. Wadą tego silnika jest zmieniająca się z obciążeniem sprawność i współczynnik mocy. Chcąc utrzymać sprawność stałą i możliwie wysoką, t. j. minimum strat, należy zmieniać napięcie na zaciskach proporcjonalnie do pierwiastka kwadratowego z mocy. Pociąga to za sobą konieczność ustawienia na lokomotywie zespołu wirującego, któryby doprowadzał do silników zmienną napięcie. Wobec włączenia między silnik a sieć tego pośredniego ogniwa można osiągnąć dalsze korzyści, gdyż zamiast skomplikowanej sieci roboczej 3 fazowej możemy pozostać przy jednofazowej, prąd jednak jednofazowy będzie o 50 okresach, t. j. czerpany bezpośrednio z sieci ogólną krajowej. Sama przetwornica jest maszyną synchroniczną o 3 000 obrotów z 2 niezależnymi uzwojeniami w statorze, z których jednofazowe, zajmujące $\frac{2}{3}$ obwodu statora na 15 000 V umieszczone jest w zamkniętych kanałach, a trójfazowe – w żłobkach półzamkniętych bezpośrednio przy szczelinie powietrznej. Prądu stałego do wzbudzenia magnesów dostarcza wzbudnica, osadzona na wspólnym wale. Napięcie prądu trójfazowego zmienia się w zależności od obciążenia od 350–600 V. Zmiany napięcia po stronie prądu trójfazowego wywołuje automatyczny regulator, działający

na wzbudzenie przetwornicy. Powoduje to oczywiście zmianę wielkości siły elektromotorycznej w uzwojeniu jednofazowym, która jednak ze względu na mniej więcej stałe napięcie sieci musi być w swych zmianach równoważona. Można to osiągnąć przez włączenie w szereg z uzwojeniem jednofazowym cewki dławikowej, dającej indukcyjny spadek napięcia, co jednak napotyka na techniczne trudności przy wykonaniu; wobec tego zastosowano zwiększenie samoindukcji uzwojenia jednofazowego przez umieszczenie go w zamkniętych kanałach. Po stronie prądu jednofazowego regulator automatyczny ma utrzymywać współczynnik mocy ($\cos \phi$) na stałej wartości = 1, ze względu na lepsze uzyskanie przewodów, transformatorów i elektrowni. Regulacja ma równocześnie zmieniać napięcie prądu trójfazowego przy stałej maksymalnej sprawności w miarę zmian obciążenia. Równoczesne zadośćuczynienie obu tym sprzecznym wymaganiom da się osiągnąć przez odpowiednie dobranie przekrojów żelaza między kanałami uzwojenia jednofazowego i żłobkami trójfazowego. Regulator automatyczny jest uruchamiany przez składową bezwatu prądu jednofazowego. Prąd przetwornicy przy krótkim zwarcu jest mniejszy od prądu przy pełnym obciążeniu (ok. 60%), dzięki czemu można przetwornicę bez obawy włączyć na sieć przed dojdzeniem do obrotów synchronicznych, a w razie gdyby przetwornica wypadła z taktu, wystarcza odciążenie na krótką chwilę, poczem po osiągnięciu pełnych obrotów można od razu włączyć prąd. Mały prąd krótkiego zwarcia wskazuje na małą przeciążalność przetwornicy, jednak właściwość ta specjalnie w kolejnictwie nie może być uważana za poważną wadę, gdyż wszelkie zmiany obciążenia (np. wjazd na wzniesienie lub tuk) następują stopniowo, w miarę wchodzenia coraz dalszych wagonów pociągu na wzniesienie lub tuk. W razie spadku napięcia w sieci roboczej występuje począwszy od pewnej mocy w górę nie tylko wyrównanie, ale nawet podwyższenie napięcia na zaciskach silników, dzięki czemu całe urządzenie jest mało czułe na zmiany

napięcia. Wzbudnica wykazuje pewne osobliwości konstrukcyjne. Bieguna magnesowe mniej więcej w połowie wysokości są połączone mostkami z blach żelaznych. W górnej części bieguna jest umieszczona cewka magnesowa bocznikowa, w dolnej – szeregową. Cewki szeregowe są normalnie krótkozwarte przez kontakty regulatora. Strumień magnetyczny cewek bocznikowych zamyka się częściowo przez twornik, a częściowo – przez mostki. Regulator działa na opornik cewek bocznikowych. W razie gwałtownego skoku obciążenia regulator włącza prąd w cewki szeregowe, co dzięki temu, że strumień magnetyczny zamyka się przez mostki i znosi tam strumień magnetyczny cewek bocznikowych, powoduje szybki wzrost napięcia, niezależny od stałej czasowej cewek bocznikowych. Rozruszniki silników trójfazowych lokomotywy wykonano wodne, przyczem dopływ powietrza sprężonego, służącego do poruszania zasuw komórek wodnych rozrusznika, następował pod działaniem sprężyny, napinanej przez maszynistę przy poruszeniu ręczki regulatora. Samą sprężynę mógł maszynista regulować stosownie do wagi pociągu. Zamykanie wentyla skuteczniał regulator elektryczny, ustawiony na moc maksymalną. Próby wykazały jednak niedostateczność tego systemu, gdyż w razie niedostosowania napięcia sprężyny do wagi pociągu następowały przeciążenia, którym automatyczny regulator napięcia nie mógł nadążyć. Zaradono temu przez zmniejszenie siły sprężyny a dodanie solenoidu i włączonego na napięcie silników i działającego na zawór. Wskutek tego zawór zaczyna się otwierać dopiero w miarę wzrostu napięcia na zaciskach silników i regulacja napięcia wyprzedza zawsze zmianę obciążenia. W razie większego spadku napięcia na sieci urządzenie to przytłacza zawór, zapobiegając nagłemu wzrostowi obciążenia przetwornicy. Zalety tego nowego systemu trakcyjnego można streścić w następujących słowach: 1) lepsze wyzyskanie materiału silników (ok. 23,7%) 2) współczynnik mocy na sieci jednofazowej = 1, dzięki czemu koleje stają się bardzo pożądanymi

odbiorcami dla elektrowni, poprawiając ogólny cos sieci, 3) lokomotywa jest mało czuła na spadki napięcia, wobec czego odległość punktów zasilających może być znacznie zwiększona, 4) kolej może być przyłączona do każdej normalnej sieci 50 okresowej za pomocą prostych stacji transformatorowych. Odzyskiwanie energii elektrycznej na spadkach jest automatyczne i nie wymaga żadnych dodatkowych skomplikowanych urządzeń. Dotychczasowe wyniki prób potwierdziły w zupełności przewidywania teoretyczne, jednak aż do czasu zebrania danych eksploatacyjnych z większej linii kolejowej zelektryfikowanej tym systemem nie można wypowiedzieć ostatecznego sądu.

1 czerwca 1925

PRZEGLĄD TECHNICZNY Kolejowe laboratorium doświadczalne

Przy Dyrekcji Warszawskiej P. K. P. ma być utworzone laboratorium doświadczalne do badania smarów, stopów, maźnic i łożysk. Prace laboratorium będą skierowane ku osiągnięciu zmniejszenia rozchodu smarów i wypadków zagrzaniasia osi.

17 czerwca 1925

Nowe linie lotnicze w Polsce

W ostatnim tygodniu ub. m. rozpoczęto komunikację lotniczą na dwóch nowych liniach: od 25-go maja utworzono komunikację między Warszawą a Poznaniem (S. A. Aero), zaś od 26-go maja – między Krakowem a Lwowem (tow. Aerolot). Podróż na obu liniach trwa 2–2 $\frac{1}{2}$ godz., loty odbywają się codziennie rano.

17 czerwca 1925

Postępy w wytwarzaniu barwników w Polsce

Do niedawna barwniki trwały, t. zw. „indantrenowe” były z Polski sprzedawane z Niemiec. Obecnie od paru miesięcy rozpoczęto fabrykację tych barwników w kraju, przyczem są one sprzedawane po cenach bez porównania niższych od barwników niemieckich. Badania co do trwałości barwników krajowych, np. antrenowego barwnika ochronnego (khakki), dokonane zarówno przez M. S. Wojsk., jak i wybitnych kolorystów tódzskich, dały znakomite wyniki.

24 czerwca 1925



Logistyka

Gdy w 2021 roku w Kanale Sueskim utknął kontenerowiec „Ever Given”, świat dosłownie wstrzymał oddech: ceny frachtu wystrzeliły, porty stanęły, a brak papieru toaletowego znów trafił na czołówki. W jednej chwili okazało się, że logistyka nie działa w tle, lecz gra główną rolę w globalnej gospodarce. To dzięki niej towary trafiają na czas, produkcja nie staje, a sklepy są zatowarowane. W cieniu tych procesów stoi ktoś, kto łączy punkty, przewiduje zagrożenia i ogarnia złożone systemy. To logistyk. A że zapotrzebowanie na takich ludzi rośnie, warto dziś sprawdzić, jak wyglądają studia na kierunku logistyka i co można dzięki nim zyskać, zanim świat znów się gdzieś zatka.

W Polsce logistykę można studiować aż na 78 uczelniach, zarówno publicznych, jak i prywatnych. I nie chodzi tu wyłącznie o wielkie miasta. Obok politechnik z tradycją, jak Warszawska, Łódzka czy Poznańska, są też akademie wojskowe i uczelnie zawodowe z mniejszych miejscowości, które mocno stawiają na praktykę i lokalne potrzeby rynku. Sam przebieg studiów też daje elastyczność. Można zacząć od licencjatu na uczelni ekonomicznej, gdzie logistyka wchodzi w skład nauk o zarządzaniu, albo pójść drogą inżynierską, nieco dłuższą, bo 3,5-letnią, ale za to bardziej techniczną. Tu w programie pojawi się fizyka, grafika inżynierska, wytrzymałość materiałów,

czyli wszystko to, co sprawia, że logistyk z dyplomem inżyniera potrafi nie tylko coś zaplanować, ale i realnie to wdrożyć. Osoby zainteresowane rozwojem mogą iść dalej na studia magisterskie, które trwają od półtora do dwóch lat. Dla jednych to kolejny krok w stronę specjalizacji, dla innych szansa, by uzyskać tytuł magistra inżyniera, który wciąż robi wrażenie na rynku pracy. Do dyspozycji są także studia podyplomowe. Krótsze, ale intensywne, często wybierane przez tych, którzy już pracują i chcą doszlifować konkretne kompetencje, jak np. logistyka magazynowa czy zarządzanie łańcuchem dostaw. Kosztują od 4 do 6 tysięcy złotych, ale wiele osób traktuje to jako inwestycję, nie wydatek.

Wybierając uczelnię, warto zwrócić uwagę na coś więcej niż samą nazwę kierunku, bo logistyka logistyką, ale specjalności potrafią się znacząco różnić. Czasem są ukryte pod nazwami typu „logistyka w biznesie” albo „logistyka miejska”, innym razem bardziej techniczne, np. „inżynieria transportu” czy „logistyka morska”. Jedna szkoła postawi na logistykę portów, inna na spedycję międzynarodową, a jeszcze inna na logistykę wojskową, z zupełnie innym zestawem umiejętności. Co ciekawe, niektóre uczelnie w ogóle nie oferują specjalności, uważając, że logistyka sama w sobie jest wystarczająco złożona i nie trzeba jej dzielić. To też coś mówi o charakterze tych studiów.

Mysłąc o nauce na tym kierunku, zaraz po wyborze uczelni należy przejść przez proces rekrutacyjny. Nie jest to przesadnie trudne zadanie, ale nie znaczy to, że każdy się dostanie. Na uczelniach publicznych punktuje się przede wszystkim matematykę i to zupełnie nieprzypadkowo, bo bez liczenia w logistyce ani rusz, a do tego język obcy i jeden przedmiot dodatkowy: geografia, fizyka, chemia, informatyka albo biologia. W 2023 roku kierunek ten wybrało ponad 10 tysięcy osób i jest to wynik mniejszy niż rok wcześniej, ale wciąż plasuje się w czołówce najpopularniejszych. W rekrutacji na rok akademicki 2024/2025 Politechnika Krakowska odnotowała 2 kandydatów na jedno miejsce. Uczelnie prywatne nie są aż tak selektywne, tam wystarczy matura, ale to, czy studia będą warte zachodu, nie zależy od progu punktowego, tylko od tego, co dzieje się po przekroczeniu drzwi dziekanatu.

A dzieje się sporo. Program kształcenia jest szeroki i co ważne, stale aktualizowany. Na początku są podstawy: matematyka, statystyka, ekonomia, prawo, towaroznawstwo. Później zaczyna się to, co naprawdę interesujące: logistyka zaopatrzenia, zarządzanie łańcuchem dostaw, planowanie transportu, informatyka w logistyce. Coraz więcej uczelni stawia na praktykę, w tym także obowiązkowe praktyki zawodowe, prace projektowe, zajęcia z systemów informatycznych. Nie jest tajemnicą, że język angielski w tej branży jest absolutnie niezbędny, nie tylko w kontaktach z kontrahentami, ale też w obsłudze oprogramowania. I choć niektórzy narzekają, że niektóre zajęcia są zbyt teoretyczne, to jednak widać wysiłek uczelni, by to zmieniać, gdyż pojawiają się gry symulacyjne, case studies oraz projekty z firmami. Nauka staje się dynamiczna, a nie tylko „zakuć-zdać-zapomnieć”. Często pojawia się pytanie, czy logistyka jest trudna? Odpowiedź brzmi: To zależy dla kogo. Dla kogoś, kto myśli szybko, lubi konkret, potrafi działać w zespole i nie boi się tabel, modeli i deadline'ów, nie. Dla tych, którzy liczyli na lekkie zarządzanie, może być zaskoczeniem.

Trzeba przyswoić sobie materiał z kilku dziedzin od techniki, przez ekonomię, po prawo. Zdarzają się momenty, gdy trzeba sięgnąć po dodatkowe źródła albo po prostu usiąść i „przeżyć” temat. Nie są to studia oderwane od rzeczywistości. Tu niemal każdą teorię można podeprzeć praktyką, a każdy przedmiot osadzić w kontekście realnego działania. To wszystko składa się na jakość kształcenia, która według studentów rośnie. Zajęcia z praktykami, wizyty w centrach logistycznych, współpraca z firmami to nie są wyjątki, ale coraz częściej standard. Studenci czują, że uczą się czegoś, co może się przydać. A absolwenci, choć czasem żalują, że nie poświęcono więcej czasu nowoczesnym systemom IT, podkreślają, że fundament, który wynieśli, pozwala im się odnaleźć w wielu branżach.

Logistyka to nie tylko transport. To cały system. To „krwioobiegi gospodarki” od produkcji, przez magazynowanie, po dostawę do klienta. Absolwenci kierunku mogą znaleźć pracę w naprawdę różnych miejscach: w firmach produkcyjnych, e-commerce, centrach dystrybucyjnych, w wojsku, w administracji, w firmach eventowych, szpitalach. Mogą być analitykami, planistami, spedytorami, menedżerami, kierownikami magazynów, doradcami ds. łańcucha dostaw. Mogą awansować od specjalisty, przez koordynatora, aż po dyrektora logistyki. A jeśli mają zacięcie, mogą założyć własną firmę: spedycyjną, kurierską, doradczą. Zarobki są zależne od miejsca, doświadczenia i odwagi negocjacyjnej. Według danych z 2024 roku, początkujący logiści zarabiali średnio 6–8 tysięcy złotych brutto. Starsi specjaliści od 9 do 12 tysięcy. Do tego premie. W niektórych firmach służbowe auta. To więcej niż jeszcze rok wcześniej. Według raportów, pensje wzrosły o blisko 8%. Rynek rośnie, zapotrzebowanie rośnie, więc i stawki idą w górę. Ale uwaga, to nie jest zawód, w którym można spocząć na laurach. Logiści, którzy nie śledzi zmian, łatwo stają się „tym do zastąpienia”. Automatyzacja, sztuczna inteligencja, Internet Rzeczy, wszystko to zmienia oblicze branży. Dziś logiści to często nie tylko osoba od paczek, ale także analityk danych, specjalista od optymalizacji, strateg. I tylko ci, którzy są gotowi się rozwijać, mają szansę nie tylko utrzymać się na powierzchni, ale i wypłynąć na szerokie wody.

Więc komu polecamy logistykę? Tym, którzy lubią łączyć fakty, nie boją się zmian, szukają pracy dającej poczucie wpływu i sensu. I tym, którzy wiedzą, że dobra organizacja to nie tylko cecha w CV, ale sposób myślenia. Bo logistyka nie jest już zawodem przyszłości. Ona jest zawodem teraźniejszości, tyle że przyszłość od dawna już tu jest. ■

Michał Pacholski



Piękna i groźna

Od chwili poznania fascynowała ludzi. Ciężka, srebrzysta ciecz, po powierzchni której pływały kamienie i bryłki metali, a tonęło jedynie złoto. Pokolenia badaczy pragnęły zgłębić jej naturę: ni to metalu, ni wody. Po wiekach okazało się, że równie jak jest piękna, tak i śmiertelnie groźna.

Artykuł o – jak się zapewne domyślasz – rtęci, dla ciebie będzie jedynie teoretyczny. Eksperymenty ze srebrzystą cieczą to codzienność w profesjonalnych laboratoriach, ale ryzyko, które niesie kontakt z nią samą i jej związkami, jest nie do przyjęcia w warunkach domowych. Musisz więc poczekać na dostęp do pracowni chemicznej i profesjonalną opiekę.

Ruchliwe kuleczki

W starożytności wokół Morza Śródziemnego i na Bliskim Wschodzie rtęć znajdowała się tylko w jednym miejscu – w obecnej Hiszpanii w okolicach miasta Almadén, gdzie wydobywano **cynober**, czerwonej barwy minerał stosowany jako barwnik (1). Pierwsze wzmianki o rtęci pochodzą z IV wieku p.n.e., gdy Kartagińczycy opisali krople srebrzystej cieczy znajdowane w szczelinach tamtejszych skał. Wkrótce nauczono się otrzymywać rtęć z cynobru: wystarczyło ogrzać minerał z opiłkami żelaza, a nawet sam na powietrzu, wtedy pary kondensowały na ustawionych obok przedmiotach, skąd można było je zebrać. Od początku starano się poznać właściwości ciekłego metalu, z wyglądu podobnego do srebra, ale jakiegoś dziwnego. Stąd nazwy: Grecy mawiali o nim *hydrargyros* („wodniste srebro”), natomiast Rzymianie – *argentum vivum* („żywe srebro”). Z pierwszego źródła pochodzi symbol chemiczny – **Hg**.

1. Dobrze wykształcone kryształy cynobru o krwistej barwie



2. Rtęć – srebrzysty, ciekły metal wraz ze swym alchemicznym symbolem

Rtęć to najpóźniej otrzymany metal starożytności, po złocie, srebrze, miedzi, żelazie, ołowiu i cynie. Alchemicy każdy z nich powiązali z planetami, Słońcem i Księżycem. W ten sposób ruchliwej rtęci przypadł najszybciej poruszający się Merkury. Nazwy rtęci w wielu językach europejskich pochodzą ze starożytności, np. angielskie *mercury* czy niemieckie *Quecksilber* („szybkie srebro”). Skąd więc nasz językowy łamaniec, którego nie wypowie żaden z cudzoziemców? Z prasłowiańskiego *rtot*, co znaczy toczyć się (stąd np. rotacja) (2).

Z rtęcią przez wieki

Metaliczny połysk w połączeniu z ciekłym stanem skupienia i dużą gęstością (13,5 g/cm³) spowodowały, że rtęć stała się modnym przedmiotem zbytku. W domach bogaczy stały misy wypełnione srebrzystą cieczą, po której pływały kamienie i metalowe figurki, najbogatsi mieli całe baseny rtęci, po powierzchni której lubili się przechadzać na rozłożonych dywanach. Zabawy te nie były bezpieczne, rtęć wrze w temperaturze 357°C, ale już w warunkach pokojowych dość intensywnie paruje. Toksyczne właściwości rtęci potwierdzono jednak znacznie później.

W starożytnym Rzymie odkryto, że w srebrzystej cieczy rozpuszcza się srebro i złoto (powstają stopy



3. Amalgamatowe wypełnienie ubytku w zębie

zwane **amalgamatami**), a z powstałego roztworu można je odzyskać przez odparowanie rtęci. Sposób ten ułatwił wydobywanie okruszków cennych kruszców z piasku rzecznoego i skał. Na najszerszą skalę metoda była stosowana w wiekach XVI i XVII, gdy hiszpańska Złota Flota woziła do kolonii w Nowym Świecie dziesiątki ton rtęci rocznie, w zamian przywożąc metale szlachetne. Obecnie stosuje się głównie metodę cyjankową, która – paradoksalnie, mimo znanych trujących właściwości cyjanków – jest znacznie bezpieczniejsza. Amalgamaty służyły również do złocenia i srebrzenia. Wystarczyło powlec np. miedzianą misę amalgamatem złota i ogrzać ją – rtęć odparowała, a na naczyniu pozostawała powłoka ze szlachetnego metalu. Amalgamaty srebra i cyny stosowano do produkcji luster, pokrywając nimi powierzchnię szkła. Jednak pozostała rtęć parowała, powoli zatrzymując właściciela tego przedmiotu (mówiono o *jadzie luster*). W znacznie późniejszych czasach amalgamatu srebra używano w stomatologii, do wypełniania ubytków w zębach (3). Jednym z niewielu metali nie tworzącym amalgamatów jest żelazo, stąd większe ilości rtęci transportuje się w stalowych butlach.

Alchemicy intensywnie badali srebrzystą ciecz, upatrując w niej surowca do przeprowadzenia transmutacji, czyli przemiany nieszlachetnych metali w złoto. Dziś już wiemy, że w ich pracowniach nie było to możliwe, ale ponoć niektórym się udawało. Czasem była to niewiedza (sławny kamień filozoficzny, umożliwiający przemianę metalu w złoto, to po prostu związek o wzorze AgAuCl_4 , który rozkładając się, tworzył złotą powłokę na przedmiocie), ale częściej celowe oszustwo (amalgamaty złota o pewnym składzie są nie do odróżnienia od rtęci, a po ich ogrzaniu pozostaje samo złoto). W razie powodzenia alchemik zyskiwał

możnego sponsora, liczącego na bajeczne zyski, ale gdy mistrz został przyłapany na oszustwie – kończył na pozłacanej szubienicy. Alchemicy w rtęci upatrywali pierwiastka metaliczności, nosicielki cech każdego metalu – połysku. Za inny niezbędny składnik uznano siarkę nadającą metalom cechę palności. Do tego potrzebna była jeszcze sól, która tchnęła ducha w materię (uważano, że metale i minerały również żyją, wszak „rodziły się” we wnętrzu Ziemi). Wystarczyło jeszcze tylko znaleźć proporcje, jakimi należy połączyć składniki. Patrząc obiektywnie, trzeba jednak stwierdzić, że alchemicy mieli rację, w rtęci właśnie upatrując surowca do produkcji złota. W latach 50. ubiegłego wieku fizycy jądrowi (następcy alchemików?) właśnie z tego metalu otrzymali bryłkę złota. Doniesienia prasowe o sukcesie spowodowały nawet krótkotrwałe zwirowania giełdowe, ale sytuację uspokoiła analiza kosztów produkcji – sztuczne złoto jest zdecydowanie droższe od naturalnego.

Na Dalekim Wschodzie, w Chinach i Japonii, rtęć była uważana za czynnik sprzyjający zdrowiu, tamtejsi alchemicy trudnili się więc wyrobem rozmaitych „eliksirów nieśmiertelności” (nie trzeba dodawać, że używanych z opłakanym skutkiem). W XVI wieku w Europie rtęć w lecznictwie zastosował szwajcarski alchemik występujący pod imieniem Paracelsus. Rozczarowany brakiem powodzenia transmutacji, zajął się medycyną, wprowadzając specyfiki wytwarzane z surowców nieorganicznych (niektóre związki rtęci nadal używane są jako leki). Do dziś również nie stracił na znaczeniu jego pogląd, że *wszystko jest trucizną i nic nie jest trucizną, bo tylko dawka czyni trucizną*.

4. Ciśnieniomierz rtęciowy





5. Rtęciowe termometry lekarskie

Rtęć posłużyła do zbudowania barometrów (Torricelli, 1643; do dziś jedną z jednostek ciśnienia są milimetry słupa rtęci – mm Hg, 1 atmosfera = 760 mm Hg), manometrów, pomp próżniowych czy też termometrów (Fahrenheit, 1725), rozszerzając możliwości nauki i techniki (4). Użycie termometrów do pomiaru niskich temperatur pozwoliło na zestalenie rtęci (krzepnie w około -39°C), co przekonało wszystkich, że i ona jest metalem (XVIII wiek) (5).



6. Kompaktowa świetlówka wypełniona parami rtęci (pixabay.com)

Nowoczesna chemia również wiele zawdzięcza rtęci. Tlenek HgO łatwo powstaje przez ogrzewanie rtęci na powietrzu, natomiast w wyższej temperaturze równie łatwo się rozkłada. Doświadczenia z tym związkiem pozwoliły na odkrycie tlenu (Priestley, 1774), a także wyjaśnienie roli tego gazu w procesach utleniania i spalania (Lavoisier, 1777–89). Łatwe tworzenie amalgamatów umożliwia otrzymanie wielu metali: podczas elektrolizy ich związków jako elektrody ujemnej używa się rtęci, wydzielony metal rozpuszcza się w niej, skąd – przez odparowanie rtęci – można go odzyskać (tak m.in. postąpiła Maria Skłodowska-Curie podczas wydzielania próbki metalicznego radu).

Po stwierdzeniu, że przepuszczanie prądu elektrycznego przez szklaną rurę wypełnioną parami rtęci powoduje emisję intensywnego promieniowania, zjawisko wykorzystano w technice oświetleniowej (duża część emisji przypada na nadfiolet, szklane bańki pokrywane były więc od wewnątrz luminoforem zamieniającym promieniowanie UV na widzialne) (6).

Trudne czasy dla żywego srebra

Toksyczna natura rtęci spowodowała, że ona sama i jej związki coraz rzadziej pojawiają się w naszym otoczeniu. Amalgamatowe plombę są zastępowane wypełnieniami z materiałów nieodbiegających kolorem od barwy zębów, choć w tym przypadku decydują głównie względy estetyczne, ponieważ amalgamat srebra jest praktycznie nierozpuszczalny, a przez to bezpieczny. Rtęć ogranicza korozję powłok ogniów jednorazowych, ale obecnie odchodzi się od jej użycia do tego celu (7). Termometry, barometry i energooszczędne źródła światła także konstruowane są bez udziału rtęci.

Obecnie główne zastosowania srebrzystej cieczy to wytwarzanie amalgamatów, ciekłych elektrod, przełączników i czujników. Związki rtęci nadal stosowane są w lecznictwie, jako środki przeciwko szkodnikom, farby okrętowe ograniczające przyleganie glonów i skorupiaków do kadłuba oraz jako detonatory.

Szacuje się, że złoża Almadén (najbogatsze źródło tego metalu na świecie) od starożytności dostarczyło około ćwierci miliona ton rtęci, ale obecnie jest to już tylko muzeum. Produkcja rtęci systematycznie spada, jeszcze kilkadziesiąt lat temu wynosiła 5...6 tysięcy ton rocznie. W roku 2024 wyprodukowano 1200 ton żywego srebra, monopolistą są Chiny z udziałem wynoszącym 1000 ton, a z pozostałych połowę dostarcza Tadżykistan.

Naprawdę groźna

Doniesienia o toksycznym działaniu rtęci pochodzą już ze starożytności. Pary łatwo przenikają przez płuca do krwiobiegu, skąd roznoszone są po całym organizmie. Rtęć, jak inne metale ciężkie, ma właściwości toksyczne dla komórek, ale prawdziwe spustoszenie czyni w układzie nerwowym, zwłaszcza w mózgu. Postać Szalonego Kapelusznika z „Alicji w Krainie Czarów” wcale nie jest fikcją literacką – w XIX wieku do wyrobu filcu, stosowanego m.in. na kapelusze, używano związków rtęci, co powodowało chroniczne zatrucia rzemieślników. Rtęć w środowisku przekształca się w związki organiczne, a te jeszcze łatwiej przenikają do organizmów żywych (stąd np. zatrucia rybami poławanymi w akwenach skażonych ściekami przemysłowymi).

Termometry i barometry rtęciowe odchodzą już do historii, masz więc coraz mniej okazji do zetknięcia się ze srebrzystą cieczą w razie ich stłuczenia. Jeżeli jednak przydarzy ci się taki wypadek, zbierz kuleczki rtęci na szufelkę lub kartkę papieru (nie używaj odkurzacza, ponieważ rozpylisz opary) (8). Zbieranie ułatwi aluminium, do którego rtęć łatwo przylega. Zebraną rtęć, resztki termometru i drut umieść w słoiku, zakręć go i oddaj do punktu zbiórki odpadów komunalnych. Miejsca, z których zebrano rtęć, możesz posypać sproszkowaną siarką lub pyłem cynkowym, co zwiąże toksyczny metal w nielotne związki (potem zmieć pozostałości). Pamiętaj o dobrym wywietrzeniu



7. W bateriach odchodzi się od używania toksycznych metali



8. Kulki rtęci należy szybko zebrać i zutylizować!

mieszkania. Nie rozbijaj świetlówek i żarówek zawierających pary rtęci, a zepsutych nie wyrzucaj do zwykłych odpadów, lecz oddaj do punktu zbiórki. Z rtęcią nie wolno postępować nieostrożnie. ■

Krzysztof Orliński



więcej chemii na stronie:
<https://tiny.pl/dptp7>



Czy można zmienić położenie pierwiastka w układzie okresowym?

Układ okresowy pierwiastków, dosyć zresztą rzadko wykorzystywany na lekcjach fizyki, jest prawdziwą kopalnią informacji z pogranicza tego przedmiotu oraz chemii. Oczywiście pod warunkiem, że trafia w ręce osoby, która umie te informacje odczytać i zrozumieć.

Miejsce pierwiastków w układzie okresowym

Już na pierwszy rzut oka widać, że jest to tabela, w której pierwiastkom przypisano konkretne miejsca według rosnącej liczby atomowej. Budowa tej tabeli jest jednak nieco dziwna. Znajduje się w niej kilka wierszy (okresów) oraz kilka kolumn (grup). Nie byłoby w tym jeszcze nic nadzwyczajnego, gdyby nie fakt, że nie w każdym wierszu, a co za tym idzie – również nie w każdej kolumnie, jest tyle samo pierwiastków. Zwłaszcza początkowe okresy wydają się mocno ubogie w swoich przedstawicieli.

Czy jest to tylko przypadek, czy stoi za tym jakieś prawo fizyki? I czy można w jakiś sposób zapisać te puste miejsca albo poukładać pierwiastki inaczej? Spróbujmy poszukać odpowiedzi na te pytania.

Odrobina historii

W 1869 roku Dmitrij Iwanowicz Mendelejew sformułował prawo okresowości pierwiastków. Stwierdził on, że właściwości pierwiastków zmieniają się periodycznie wraz ze wzrostem ich mas atomowych. Nie znał jednak przyczyn tego zjawiska. Na podstawie obserwacji uporządkował wszystkie znane wówczas pierwiastki w tak zwaną tablicę Mendelejewa, czyli używany do chwili obecnej układ okresowy.

Dopiero na przełomie XIX i XX wieku pojawiły się pierwsze potwierdzone doświadczalnie teorie wewnętrznej struktury atomu, które pozwoliły powiązać właściwości fizyczne i chemiczne pierwiastka z jego budową. Wcześniejsze modele oparte na założeniu o niepodzielności atomów uniemożliwiały wyjaśnienie okresowo powtarzających się podobieństw między pierwiastkami.

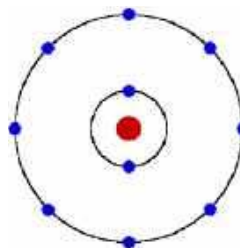
Budowa atomu

Zaproponowany w 1913 roku przez Bohra model budowy atomu nie był wprawdzie pierwszy, był natomiast jedynym, który w wystarczającym stopniu tłumaczył niektóre zjawiska, takie jak na przykład powstawanie widm emisyjnych gazów.

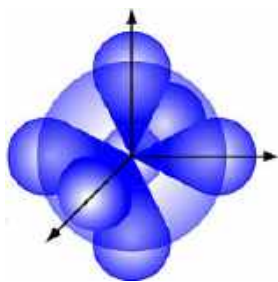
Zgodnie z tym modelem elektrony krążą w atomie po kołowych orbitach. Każda orbita może pomieścić określoną liczbę elektronów, przy czym najpierw wypełniane są niższe orbity a dopiero później – wyższe. Najmocniej związane z jądrem są orbity całkowicie wypełnione. Elektrony znajdujące się na niezapełnionych orbitach (elektrony walencyjne) są słabiej związane z jądrem i mogą tworzyć wiązania chemiczne z elektronami walencyjnymi innych atomów.

Dopiero w drugiej i trzeciej dekadzie XX wieku rozwój mechaniki kwantowej pozwolił na dopracowanie tego modelu. Przede wszystkim płaskie orbity zostały zastąpione przestrzennymi orbitalami, czyli obszarami, w których z największym prawdopodobieństwem znajduje się elektron.

Co ciekawe – nie wszystkie orbitale mają kształt sferyczny. Pierwsza powłoka, odpowiadająca pierwszej orbicie w modelu Bohra, składa się ze sferycznego orbitalu typu s. Druga powłoka składa się w sumie z czterech orbitali: jednego sferycznego typu s i trzech



1. Model atomu według Bohra. Na rysunku zaznaczono dwie pierwsze orbity oraz maksymalną liczbę elektronów, które mogą się na nich znaleźć



2. Wizualizacja przestrzennego rozmieszczenia orbitali w atomie dla pierwszej i drugiej powłoki elektronowej

orbitali typu *p* o kształcie zbliżonym do hantli. Kolejne powłoki składają się z coraz większej liczby orbitali o coraz bardziej skomplikowanych kształtach.

Konfiguracja elektronowa

Na każdym orbitalu mogą znaleźć się maksymalnie dwa elektrony o przeciwnych spinach. Zatem w przypadku pierwszej powłoki mamy tylko dwie możliwości jej wypełnienia. Odpowiada im wodór o jednym elektronie na orbicie oraz hel o powłocie całkowicie wypełnionej dwoma elektronami. Oba te pierwiastki są jedynymi reprezentantami pierwszego okresu. Do pierwszej powłoki nie możemy dołożyć więcej elektronów.

Jeśli zaczniemy wypełniać drugą powłokę, to zauważymy, że wszystkim możliwym konfiguracjom elektronów odpowiadają kolejne pierwiastki z drugiego okresu. Jest ich tam osiem, ponieważ jest dokładnie tyle możliwości rozmieszczenia elektronów na drugiej powłocie. Kolejne okresy odpowiadają kolejnym powłokom, a ich złożona struktura znajduje odzwierciedlenie w rosnącej liczbie pierwiastków należących do danego okresu.

Gdybyśmy z kolei przyjrzyli się grupom, to zauważymy, że tu również obserwuje się pewne prawidłowości. Na przykład w pierwszej grupie wodór ma jeden elektron walencyjny na pierwszej powłocie; lit – również jeden elektron walencyjny, ale na drugiej powłocie; sód – także jeden elektron walencyjny, ale na trzeciej powłocie.

Analizując stopień wypełnienia powłok dla poszczególnych pierwiastków, możemy zauważyć, że pustych miejsc w tablicy Mendelejewa nie da się niczym wypełnić. Nie odpowiada im bowiem żadna istniejąca w przyrodzie konfiguracja elektronowa. Na przykład nie możemy mieć atomu o trzech elektronach walencyjnych na pierwszej powłocie. Okazuje się, że położenie pierwiastka w układzie okresowym jest z góry określone przez budowę jego powłok elektronowych.

Reakcje chemiczne

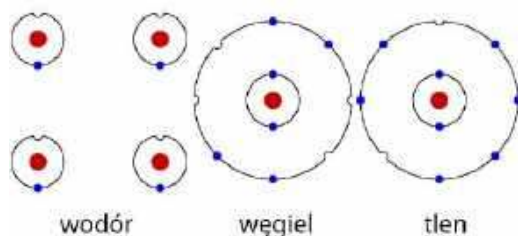
W przypadku pierwiastków mających niezapełnione powłoki dochodzi do tworzenia wiązań z atomami

tego samego lub innych pierwiastków. Elektron z jednego atomu wypełnia puste miejsce na powłocie drugiego i na odwrót. Częsteczką zaczyna funkcjonować w taki sposób, jakby te dwa elektrony były dla obu atomów wspólne.

Jeśli przyjrzymy się ostatniej grupie układu okresowego i znajdującym się w niej gazom szlachetnym, to okazuje się, że każdy z nich ma całkowicie zapełnione powłoki, zatem nie może tworzyć wiązań elektronowych. W warunkach laboratoryjnych można uzyskać związki takich gazów, ale wiązania tworzone są w nich na skutek oddziaływań elektrostatycznych pomiędzy całymi orbitalami, a nie pojedynczymi elektronami.

Sprawdź swoją wiedzę

Skopiuj poniższą ilustrację w kilku egzemplarzach i wytnij z niej schematyczne rysunki atomów. Następnie ułóż cząsteczki takich związków jak na przykład H_2O , CO_2 , CH_4 , dobierając elektrony z dwóch atomów w wiązania. Porównaj uzyskane cząsteczki ze wzorami strukturalnymi tych związków.



Dla nauczyciela

Omówione w niniejszym artykule zagadnienie łączy ze sobą wiadomości, które uczeń powinien nabyć w szkole podstawowej na przedmiocie chemia (dział Wewnętrzna budowa materii) z informacjami niezbędnymi w szkole ponadpodstawowej na przedmiocie fizyka. Gruntowne przyswojenie tych treści jest konieczne do zrozumienia nie tylko chemii, ale również fizyki atomowej czy fizyki jądrowej.

Nauczyciele przedmiotu fizyka zazwyczaj zakładają, że uczeń opanował już umiejętności z tego zakresu, co niekoniecznie musi być prawdą. Warto zatem poświęcić przynajmniej jedną lekcję fizyki na przypomnienie podstawowych informacji na temat budowy atomów, na przykład przed pierwszymi zajęciami dotyczącymi ruchu cząstki naładowanej w polu magnetycznym, a potem w razie potrzeby konsekwentnie nawiązywać do tych zagadnień. ■

Joanna Borgensztajn



Michał Szurek tak mówi o sobie: „Urodzony w 1946. Ukończyłem UW w 1968 roku i od tego czasu tam pracuję na Wydziale Matematyki, Informatyki i Mechaniki. Specjalność naukowa: geometria algebraiczna. Ostatnio zajmowałem się wiązkami wektorowymi. Co to jest wiązka wektorowa? No, trzeba wektory mocno powiązać sznurkiem i już mamy wiązkę. Do „Młodego Technika” zaciągnął mnie siłą kolega fizyk, Antoni Sym (przyznaję, powinien mieć z tego powodu tantiemy od moich honorariów autorskich). Napisałem kilka artykułów, a potem zostałem i od 1978 roku co miesiąc możecie Państwo czytać, co też myślę o matematyce. Lubię góry i mimo nadwagi staram się chodzić. Uważam, że najważniejsi są nauczyciele.

Polityków, niezależnie od opcji, jaką prezentują, trzymałbym w pilnie strzeżonym miejscu, żeby nie mogli uciec. Karmił raz dziennie. Lubi mnie jeden pies z Tulec, rasy beagle”.



Erlangen

– małe, słynne miasteczko

Dzisiejszy odcinek „Rozmaitości Matematycznych” jest trochę nietypowy. Liczę, że przeczytają go i maturzyści. Wprowadzie znaczna ich większość już podjęła decyzję, co robić po maturze, ale może chociaż kilka osób namówię na studia matematyczne?

W czerwcu tego roku przypadły dwie rocznice związane z niemieckim matematykiem Feliksem Kleinem (1849–1925): setna rocznica jego śmierci (22 czerwca) i ważniejsza dla nas: 125. rocznica przyznania mu doktoratu honoris causa przez Uniwersytet Jagielloński. Klein był pierwszym matematykiem, który dostąpił tego zaszczytu od polskiego uniwersytetu. „Polskiego”, bo ówczesna Galicja miała szeroką autonomię w ramach cesarstwa Austro-Węgier. W Warszawie polszczyzna weszła do szkół dopiero w 1905 roku, a w zaborze pruskim całe nauczanie było po niemiecku. Okazją do takiego wyróżnienia Kleina, bardzo już wtedy wybitnego naukowca, był jubileusz 500-lecia UJ. Przypomnę, że uniwersytet w Krakowie został założony przez Kazimierza Wielkiego, a więc ostatniego Piasta na polskim tronie. Dlaczego zatem dziś nazywa się „Jagielloński”? Dlatego, że wkrótce po założeniu praktycznie przestał działać i dopiero królowa Jadwiga sprawiła, że go reaktywowano, właśnie w 1400 roku. Pierwszym rektorem został Stanisław ze Skarbimierza (dziś: Skalmierz), 22 lipca 1400 roku. Pamiętam jednak, jak w PRL zorganizowano wielkie obchody 600-lecia w 1964 roku.

Ówczesna władza lubiła wszelkiego rodzaju manifestacje patriotyczne, choć dzień 22 lipca był dla „nich” ważny z innego powodu niż wybór pierwszego rektora pierwszego polskiego uniwersytetu.

Zajrzałem do życiorysu Feliksa Kleina w Wikipedii. Jego osiągnięcia naukowe są tam omówione dość pobieżnie, ale to zrozumiałe. Natomiast z przykrością zauważyłem, że nie ma wzmianki o jego doktoracie h.c. na UJ. Cóż, tekst pisany był pewnie przez kogoś, kto może tylko z grubsza orientuje się w geografii Europy Środkowej.

Ale jeszcze zanim przejdę do rzeczy (czyli do matematyki), wspomnę, że Felix Klein był laureatem

bardzo prestiżowego w tamtych latach Medalu Copleya. Nagrody Nobla nie mógł dostać, bo zgodnie z jej regulaminem, matematycy nie są do niej dopuszczani. Polecam tu Czytelnikom wydaną kilka lat temu książkę Krzysztofa Ciesielskiego i Zdzisława Pogody „Królowa bez Nobla”. Najwyższego obecnie wyróżnienia dla matematyków (czyli tzw. Medalu Fieldsa) też Klein nie dostał, bo zmarł, zanim ten medal został ustanowiony.

Kilka informacji o brytyjskim Medalu Copleya. Przyznawany on jest przez Towarzystwo Królewskie



1. Felix Klein (1849–1925)

w Londynie za „trwałe, wybitne osiągnięcia w dowolnej dziedzinie nauki”. Jest to najstarsza regularna nagroda naukowa. Nazwa wyróżnienia pochodzi od jednego z członków – Godfreya Copleya, zamożnego ziemianina, który w XVIII wieku ufundował nagrodę pieniężną w wysokości 5000 funtów szterlingów. Zapytałem chatGPT, ile wynosiło wynagrodzenie profesora nauk ścisłych w Oksfordzie w 1730 roku. Odpowiedział (-a, -o), że między 100 a 300 funtów rocznie (co prawda plus różne dodatki, jak darmowe zakwaterowanie i opłaty za wykłady). Tak czy owak, wysokość nagrody Copleya można było porównać z wynagrodzeniem za 15 lat pracy. Obecnie z Medalem Copleya związane jest tylko 25 tysięcy funtów. Da się za to kupić nawet... dziesięć metrów kwadratowych mieszkania w Warszawie!

Pierwszym laureatem (i to w dwóch kolejnych latach, 1731 i 1732) został Stephen Gray. Warto wspomnieć o tym rzemieślniku, który był też naukowcem-amatorem i jako pierwszy prowadził systematyczne badania nad przewodnictwem elektrycznym. Brytyjczycy uważają, że to on powinien być pamiętany jako „ojciec elektryczności”. Jak wiemy, przez następne sto lat z okładem elektryczność była tylko ciekawostką, bez widoków na jakiegokolwiek zastosowanie... Dziwne, prawda? W gronie laureatów Medalu Copleya odnajdujemy takie nazwiska jak Charles Darwin, Albert Einstein, Michael Faraday, Benjamin Franklin, Stephen Hawking, Peter Higgs (który 1964 roku przewidział istnienie ważnej części elementarnej – bozonu; odkrytej dopiero w 2012 w Wielkim Zderzaczu Hadronów), Friedrich Kekulé von Stradonitz (odkrywca pierścieniowej budowy benzeny, co zapoczątkowało nowy dział chemii organicznej – chemię związków aromatycznych), Georg Ohm, Ludwik Pasteur, Andrew Wiles i Alessandro Volta. Przedostatni z tej listy, Andrew Wiles, stał się sławny dzięki rozstrzygnięciu (w 1992 roku) trapiącego matematyków od około 1660 roku zagadnienia Fermata: czy dla wykładnika $n > 2$ są takie liczby naturalne, że $x^n + y^n = z^n$. Pokolenia matematyków wierzyły, że jest to równanie bez rozwiązania, ale wszelkie dowody były niekompletne albo błędne. Z nazwiskami Ohm i Volta spotykamy się codziennie, na przykład zapalając lampę.

Z nazwiskiem Kleina łączy się w matematyce kilka obiektów; mamy grupę Kleina, kwadrykę



2. Źródło: oferta firmy „Strużynkowo” na Allegro

Kleina, niezmiennik Kleina i butelkę Kleina – specyficzną powierzchnię, niemieszczącą się w trzech wymiarach. Według anegdoty ta nazwa powstała z pomyłki przy tłumaczeniu z niemieckiego (Flache – płaszczyzna, powierzchnia) na angielski. „Ktoś” odczytał „Flache” jako „Flasche”, czyli „butelka”. Istotnie, można tę powierzchnię przedstawić jako dziwną butelkę albo nawet wazon (2). Dawno temu moja koleżanka sprezentowała mi na imieniny taką „butelkę” zrobioną na drutach! Takie przedstawienie butelki Kleina jest jednak przekłamane – jej „ucho” wchodzi do wnętrza bez przecinania ścianki, co jest możliwe tylko w przestrzeni czterowymiarowej. Czy to da się wyobrazić? Oczywiście, że... tak.

Dla matematyków Klein jest znany najbardziej z jego tak zwanego „programu z Erlangen” – nowatorskiego spojrzenia na całą geometrię. To podejście wytrzymało próbę czasu. Dziś patrzymy na tę gałąź matematyki w taki właśnie sposób.

Erlangen jest miastem niedaleko Norymbergi, najbardziej znanym ze swojego uniwersytetu, zresztą będącego dziś partnerskim uniwersytetem dla UJ. W XIX wieku w uniwersytetach niemieckich nowo przyjęty profesor prezentował się, wygłaszając wykład. Feliks Klein zaprezentował się wykładem, który przeszedł do historii: *Vergleichene Betrachtungen über neue geometrische Forschungen*. W tłumaczeniu Samuela Dicksteina (1894) brzmi to: „Rozważania porównawcze o nowszych badaniach geometrycznych”. Już na samym początku Klein sprecyzował, o co chodzi (przytaczam według tłumaczenia Dicksteina):

Pojęciem istotnem i koniecznem w rozważaniach poniższych jest pojęcie grupy przekształceń przestrzennych. Połączenie dowolnej liczby przekształceń przestrzeni daje zawsze znowu takie przekształcenie. Jeżeli szereg przekształceń ma tę własność, że każde przekształcenie złożone z przekształceń do niego należących, samo do szeregu należy, nazywa się grupą przekształceń. (...) Własności geometryczne cechuje niezmiennosc wobec przekształceń grupy.

O co tu chodzi? To jednocześnie proste i skomplikowane. Czy widoczne na **rysunku 4** czworokąty są takie same? Ktoś powie, że nie, bo jeden jest większy, a mniejsze są obrócone względem siebie – więc są to różne figury, a ktoś inny, że takie same,



bo w gruncie rzeczy mają ten sam kształt. Kiedy nauczycielka mówi do dzieci: „Przepiszcie z tablicy”, to każde dziecko pisze trochę inaczej – no i literki pani na tablicy są kilka razy większe niż te w zeszytach dzieci. A jednak wszyscy rozumieją, że A to A.

Można to też ująć inaczej: to, czy rzeczy są takie same, czy inne, zależy od punktu widzenia. W często tu podawanym, choć nieco „seksistowskim” powiedzeniu, że pan Zbigniew odróżnia samochody po ich marce, a jego żona po kolorach karoserii.

Te porównania wydają się niezbyt poważne, a jednak dobrze ilustrują matematyczną treść „programu z Erlangen”. Przepiszę kluczowe zdanie: „własności geometryczne cechuje niezmienniczość według pewnej grupy”.

Spójrzmy na przykłady. Wielokąty na **rysunku 4** są „takie same” ze względu na grupę złożoną z obrotów i powiększeń (bądź pomniejszeń) w dowolnej skali. Gdybyśmy dopuścili jeszcze symetrie, wtedy litera L byłaby taka sama (jako figura geometryczna), jak grecka Γ (5).

Na **rysunku 6** jest jeszcze inaczej. Jeżeli dopuścimy bardzo ogólne przekształcenia (ściskanie, rozciąganie – jednych fragmentów mniej, innych bardziej),

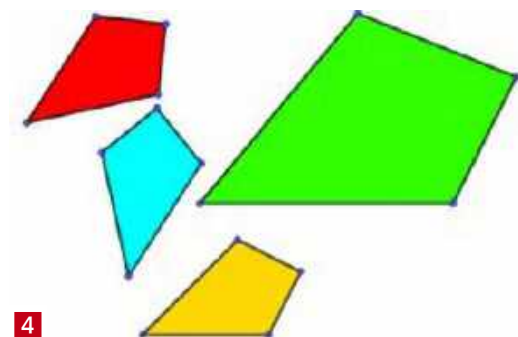
to wszystkie figury na **rysunku 6** będą „takie same”. Wchodzimy tu w najbardziej ogólną geometrię – topologię. Ma ona w swojej nazwie odniesienie do „toposu” – pojęcia, które wprowadził Arystoteles, a które jest używane w retoryce i dzisiaj. Pierwotnie znaczyło to „miejsce w przestrzeni”, dziś jest pojmowane mniej więcej jako „miejsce w myśli”, jako powszechnie przyjmowana konwencja literacka, odporna na przemiany w czasie. Myśląc o najbardziej ogólnej geometrii, myślimy o topologii. W sensie tej gałęzi geometrii figury na **rysunku 7** są już inne.

Następny **rysunek 8** pokazuje początek klasyfikacji topologicznej wszystkich powierzchni (dla ścisłości dodam, nie wyjaśniając: gładkich i orientowalnych). Łatwo domyślić się, co będzie dalej – coraz więcej połączonych precli.

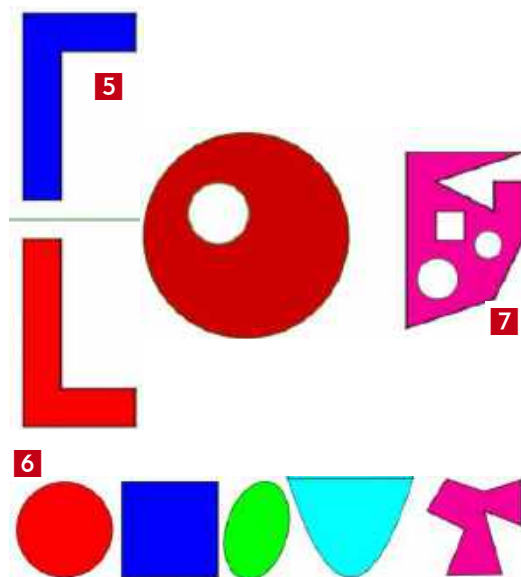
Już w XVI wieku malarze zaczęli interesować się geometrią rzutową, czyli po prostu geometrią perspektywy zbieżnej. Potrzebna była w związku z nowym, już nie średniowiecznym, a renesansowym sposobem odwzorowywania rzeczywistości na płótnie. Matematycznie jest to geometria oparta właśnie na tak zwanych przekształceniach rzutowych. Na **rysunku 9** mamy przeliczenie stopni Celsjusza na Fahrenheita (i odwrotnie) za pomocą charakterystycznego rzutowania środkowego. Spójrzmy następnie na **rysunek 10** – tak widzi pas startowy pilot lądującego samolotu. Pas zwęża się i na horyzoncie staje się punktem. Ale pilot nie tym się przejmie. Wie, że tak naprawdę linie pasa są równoległe. Przecież nikt z nas nie boi się „wsiąść do pociągu byle jakiego” (to słowa z piosenki Maryli Rodowicz),



3. Źródło: Internet + własny dodatek

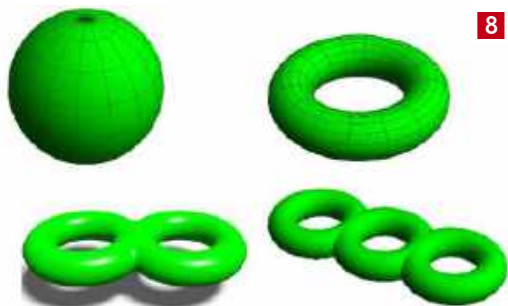


4



6

7



8

choć oczy mówią nam, że one się zewężają i pociąg powinien się natychmiast wykołcić.

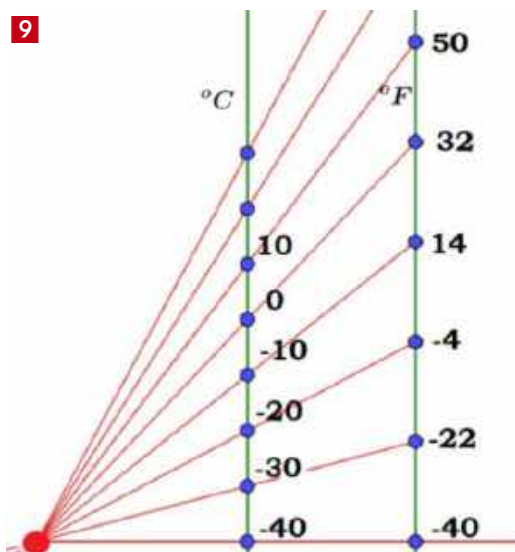
To wszystko wydaje się może i frapujące, ale bez zastosowań. Tak nie jest. Sytuacja jest nieco podobna do odkrycia Roentgena. Gdyby dostał polecenie: „proszę znaleźć metodę prześwietlania ciała ludzkiego”, mało prawdopodobne, by wynalazł aparat kojarzony dziś z jego nazwiskiem. Odkrycie promieni i późniejszy wynalazek aparatu był efektem szerokich badań podstawowych.

Pierwszy kształt na rysunku 8 to znana i lubiana sfera; podstawowy sprzęt dla golfistów, pingpongistów, tenisistów, futbolistów, siatkarzy i koszykarzy. Topologicznie nie różni się od piłki do rugby. Drugi kształt też znalazł zastosowanie w sporcie: do wynalazionej przez Włodzimierza Strzyżewskiego (1931–2001) gry „ringo”. W matematyce taka powierzchnia (a także bryła) nazywana jest torusem – taką właśnie nazwę miały charakterystyczne zgrubienia przy podstawie kolumn greckich.

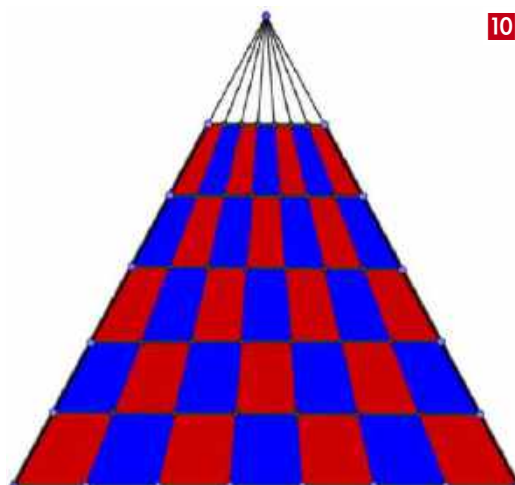
„Topologicznie” można torus zrobić z prostokąta. Najpierw zwiijamy prostokąt w rurkę, a potem rurkę w „oponkę”. Można zresztą wyjść od dowolnego równoległoboku albo od nieskończonej szachownicy na płaszczyźnie – podobnej do tej z rysunku 10, tylko złożonej z regularnych równoległoboków.

I tu potrzebna jest matematyka na poziomie uniwersyteckim. Trudno tu wyjaśnić szczegóły, ale konkretna krata daje konkretny torus. Wszystkie one nie różnią się topologicznie – wyglądają wszystkie tak samo jak dętka czy oponki. A jednak różnią się, gdy weźmiemy – dzięki Kleinowi – inną grupę przekształceń, czyli wejdziemy w inną geometrię, zwaną algebraiczną. W niej funkcje, którymi można wszystko opisywać, nie są dowolne, tylko wielomianowe. Odpadają więc funkcje trygonometryczne (takie jak sinus i cosinus) i wykładnicze, jak np. $2x$.

Z każdego torusa „algebraicznego” powstaje tak zwana krzywa eliptyczna, za pomocą której można budować doskonale szyfry. Może nie tyle „doskonale”, co bardzo dobre. Dziś szyfrowanie przydaje się nie



9



10

tylko szpiegom. Jest przydatne na przykład przy kryptowalutach. Nasze „cyberbezpieczeństwo” bardzo silnie zależy właśnie od dobrych szyfrów – jak najtrudniejszych do złamania.

Trudno powiedzieć, czy bez odkryć Kleina mielibyśmy dziś tak rozwiniętą kryptografię. Przypomina to rozważania, czy gdybyśmy nie odkryli elektryczności, to mielibyśmy samoloty (wspomniałem wcześniej pioniera badań nad elektrycznością, Stephena Graya). Transkontynentalny odrzutowiec ze wszystkimi manetkami poruszonymi ręcznie (jak zapory kolejowe za czasów wczesnej młodości autora tego tekstu) i z oświetleniem naftowym? Byłoby ciekawie...

Podobnie mogłaby wyglądać matematyka, gdybyśmy nie poszli według programu z Erlangen Feliksa Kleina. ■

Michał Szurek

TAWOIA Glass (szkło kwarcowe)

<https://sklep.avt.pl/pl/menu/tawoia-glass-4505.html>



BESTSELLERY sklepu AVT – sklep.avt.pl

3 unikalne
serie
gniazdek
i
włączników

Rabat dla Czytelników MT
przy zakupie podaj kod **MT2505GW**

-5%

Rabat dla Prenumeratorów MT
przy zakupie podaj numer prenumeraty

-10%

Ceramic Loft (ceramika)

<https://sklep.avt.pl/pl/menu/seria-ceramic-loft-4190.html>



Retro PRL (bakelit)

<https://sklep.avt.pl/pl/series/retro-prl-3237.html>





Szkoła Wynalazców

Zadaniem waszym było: zaprojektować, tzn. sformułować podstawowe założenia do projektu niewielkiego domu na Marsie, tak aby umożliwił przeżycie w nim bez skafandrów.

Zuzanna Rolnik uważa, że domek marsjański musi być przede wszystkim hermetyczny i to także od strony podłogi – fundamentów. Musi być odporny na niskie (do -65°C) i zmienne temperatury. Z uwagi na nieznaną florę bakteryjną musi być także aseptyczny. Podstawową sprawą jest wyżywienie kosmonautów. Być może hodowla chlorelli częściowo załatwiłaby problem, a poza tym pomogłaby w utrzymaniu zawartości tlenu w powietrzu w domku. Domek musi być także odporny na meteory.

Na ogół Zuzanna poprawnie określiła zagrożenia marsjańskie i przeciwdziałanie. Jednakże nie zabrała głosu w sprawie samej budowy domku. Czy ma być zbudowany z tego, co jest na Marsie, czy materiały do budowy mają być dostarczone z Ziemi, a to jest chyba najważniejsza sprawa.

Roman Tarło – z uwagi na ogromny koszt transportu czegokolwiek z Ziemi na Marsa, budowa domku musi być zrealizowana z tego, co jest na Marsie albo wykonana w formie budowli modułowej, z lekkich elementów, raczej elastycznych. Ochrona przed meteorami to właśnie elastyczność materiałów i być może ochrona aktywna, czyli niszczenie promieniem lasera – meteorów i innego drobiazgu kosmicznego. Domek musi zapewniać warunki odpowiednie do przeżycia ludzi, a więc temperaturę, powietrze i wodę oraz

żywność z własnej hodowli. To, że musi być hermetyczny w 3D, to oczywistość.

Kolega położył nacisk na sam problem budowy domku. I słusznie, bo to podstawowa sprawa. Warunki do biologicznego przetwarzania ludzi to już sprawa wtórna i raczej oczywista.

Wymienionym: koleżance i koledze, gratuluję i zapraszam do kolejnych zadań.

Nowe zadanie

Denerwuje nas niekiedy pukanie do drzwi pokoju, w którym właśnie ćwiczymy wzory skróconego mnożenia z matematyki. Jak ktoś chce wejść, to i tak wejdzie, bez pukania i naszego „proszę”. Można to przecież załatwić bezgłośnie. A więc: zaproponuj urządzenie, które informuje o zamiarze wejścia do pokoju, bez pukania lub dzwonienia. Jednocześnie sygnalizujący zamiar wejścia powinien otrzymać jakiś sygnał, informujący np. „nie przeszkadza” itp.

Istotnie: hałas, w tym także dzwonek do drzwi – jest elementem akustycznego zaśmiecania naszego środowiska. Każda próba ograniczenia hałasu jest działaniem pozytywnym. Warto temu zagadnieniu poświęcić nieco więcej uwagi.

Przypominam o terminie nadsyłania rozwiązań: 30 lipca 2025 r.

Klub Wynalazców

Terroryzm dotyka często najbardziej wrażliwe systemy, jakimi są m.in. samoloty pasażerskie. Jeśli terrorystom uda się wejść do kokpitu, to samolot już przepadł. Wasz system powinien wykluczać możliwość wejścia osób niepowołanych do kabiny pilotów, włączając w to nawet stewardesy. Waszym zadaniem było: zaproponować ideę systemu zabezpieczającego załogę samolotu pasażerskiego przed wtargnięciem do kokpitu niepowołanych osób.

System powinien być w 100% niezawodny. Powinien zakładać rozwiązania radykalne.

Stanisław Bzdek – najprościej byłoby odciąć całkowicie kokpit od kabiny pasażerów ścianą kulkoodporną bez drzwi, tak żeby kokpit był całkowicie

autonomiczny: powinien mieć WC, minikuchenkę z ekspresem do kawy i mikrofalą do podgrzania potraw, niezbędnych przy lotach transatlantycznych.

Najprostsza i chyba najskuteczniejsza metoda wyeliminowania w 100% możliwości opanowania samolotu przez



terrorystów. Oczywiście należałoby pomyśleć o komunikacji załogi z pasażerami, ale w technice telewizyjnej, i to jest proste i możliwe. Wyprowadzając rozwój sytuacji, należałoby całe wiązki przewodów elektrycznych i hydraulicznych opancerzyć w sposób uniemożliwiający terrorystom szantażowanie załogi możliwościami pozabawienia samolotu zdolności sterowania.

Zygmunt Fijałkowski pisze: słabym punktem obecnie stosowanych zabezpieczeń jest fakt, że np. mając drzwi do kokpitu zamykane na zamek szyfrowy, to jednak cała załoga musi posiadać ten szyfr, żeby wejść do kokpitu, chociażby z kawą dla pilotów. Zygmunt uważa, że najprostszy system mógłby wykorzystywać np. odciski palców członków załogi, którymi otwieraloby się drzwi do kokpitu. Dla zwiększenia bezpieczeństwa zapis odcisku palca np. stewardesy mogłby być przekazywany przez radio do którejś z baz, gdzie osoba uprawniona, również weryfikowana przez odcisk palca, drogą radiową otwierałaby drzwi do kabiny pilotów. W ten sposób nawet terroryzowana stewardesa nie mogłaby sama otworzyć drzwi, bo nie znałaby szyfru.

Sposób wydaje się niezły, ale nie daje 100% pewności. Jeśli terroryści znajdą ten system, to zmuszą stewardesę

do użycia swojego palca, a baza nie zna przecież sytuacji i otworzy drzwi!

Wymienionym kolegom gratuluję i zapraszam do dalszych zadań.

Nowe zadanie

We współczesnym „zagonionym” świecie, zatłoczonym informacjami płynącymi z różnych mediów, coraz trudniej jest pamiętać o różnych regularnych czynnościach, jak: picie wody, przyjmowanie leków i suplementów, gimnastyka dla „komputerowców” itp. Przydałby się pomocnik, który o wszystkim tym by pamiętał i sygnalizował potrzebę np. karmienia rybek. A więc: zaprojektuj przenośne urządzenie lub gadżet, który pomaga zarządzać codziennymi czynnościami (np. przypomina o picu wody, lekach, zajęciach).

Oczywiście możemy tu wykorzystać posiadane gadżety, jak: smartfony, tablety itp., lub zupełnie inne rzeczy. Wszystkim życzymy dobrych pomysłów i udanych rozwiązań. Termin nadsyłania propozycji: 31 lipca 2025 r.

Vademecum Młodego Wynalazcy

W ostatnim wydaniu VMW przedstawiliśmy sytuację, gdy wynalazca, męcząc się nad jakimś problemem, nieoczekiwanie wpada na pomysł, którego idea pochodzi z zupełnie innej branży.

Dziś jednak czekanie na taki twórczy moment wygląda nieco anachronicznie. Pomysł powinien pojawić się wtedy, gdy jest potrzebny. Czyli natchnienie na zamówienie! Brzmi to niemal świątoburczo! Przecież natchnienie to sfera niemal mistycznych doznań, stany „zmienionej świadomości” itd. Współczesna technika „nie ma czasu” na takie pomysły, ale jednak jakieś wspomaganie naszej świadomości bardzo by się czasami przydało! Problem taki pojawiał się już dawno i stąd wiele pomysłów zmierzających do stworzenia jakiegoś sposobu ułatwienia tego „wpadania na pomysł”. Chodzi o to, że wynalazca, który poszukuje rozwiązania jakiejś nowej sprawy, chodzi na ogół pod presją okoliczności i wciąż próbuje: może tak albo nie tak, tylko inaczej, może jednak tak? W takim stanie umysłu na ogół jest wrażliwy na sytuacje nieoczekiwane, ale potrzebny jest jeszcze jakiś impuls, ukierunkowujący myśl.

Tu przychodzi z pomocą jedno z pierwszych pojęć TRIZ, mianowicie „wektor inercji”. Pojęcie to oznacza po prostu typowo „ludzkie” przyzwyczajanie się

do dobrze znanych rozwiązań i unikanie ryzyka. Ilustruje to stara anegdota: Czterej koledzy wybrali się samochodem za miasto, na ryby. W pewnym momencie silnik kichnął i samochód się zatrzymał. Pierwszy – elektryk krzyknął: mówilem, że trzeba cewkę wymienić! Poczekaj – mówi drugi: przecież słyszałem, że nawaliła pompa paliwowa. Trzeci – informatyk: panowie, spokojnie; wysiadźmy i wsiadźmy jeszcze raz. Czwarty – zaopatrzeniowiec w jednej z firm: a może po prostu brak benzyny! Tak na oko każdy mógł mieć rację (z wyjątkiem informatyka) i jak widać, każdy poszukiwał rozwiązania problemu „w swoim ogródku”. Zjawisko to jest dość częste i zawsze wymaga zastosowania jakiejś metody, żeby nie dreptać w miejscu. **Rysunek 1** jest ilustracją pojęcia wektora inercji. Przedstawia on graficznie różne kierunki myślenia i poszukiwania rozwiązania problemu.

W centrum obszaru wiedzy z wszelkich możliwych do pomyślenia dziedzin nauki i techniki znajduje się „Zadanie”, czyli problem do rozwiązania. Typowym zjawiskiem jest skupienie strzałek-wektorów w stosunkowo wąskim kącie dziedzin; na rysunku przedstawiono to w formie zgęszczonych strzałek. Symbolizuje to poszukiwanie rozwiązania problemu w obszarze znanym i dobrze spenetrowanym. Wypadkowa tych



Jakie znaczenie ma orientacja i znajomość wektora inercji. Po prostu zwalnia twórcę od trzymania się utartych metod i dodaje odwagi przy sięganiu po rozwiązanie „dzikie”.

Nie daje jednak konkretnego kierunku poszukiwań. Tu mamy inne metody. Jedną z nich jest metoda kół Lulliego. Sprzętem pomocniczym jest zestaw koncentrycznych tarcz podzielonych na kilka sektorów, na których są napisane lub narysowane symbole różnych idei, zjawisk, przedmiotów, zwierząt itp. Tarcze można obracać niezależnie, otrzymując różne zestawienia rysunków lub symboli idei. Każde z takich otrzymanych zestawień „przymierzamy” do naszego problemu i analizujemy przydatność otrzymanej w ten sposób idei do rozwiązania naszego problemu. Tak mogło się stać przy poszukiwaniu idei nawijarki do małych rdzeni transformatorów. Jak wiemy Jegorow wpadł na dobry pomysł, obserwując pracę kobiety szydełkującej jakąś koronkową robótkę. Był to przypadek. Metoda kół Lulliego pozwala generować setki takich „przypadków” w krótkim czasie. Warto dodać, że idea kół Lulliego stała się „natchnieniem” konstruktorów niemieckiej maszyny szyfrującej – Enigmy.

Kolejną próbą, już współczesną jest „burza mózgów”. Twórcą burzy mózgów był Alex Osborn, amerykański specjalista od reklamy i współzałożyciel agencji BBDO. Opracował tę metodę w latach 30. XX wieku jako sposób na generowanie kreatywnych pomysłów w zespołach, złożonych z kilku osób reprezentujących różne zawody i specjalności. Główną cechą sesji burzy mózgów jest wykluczenie krytyki w fazie generowania pomysłów. Częstym „wrogiem postępu” jest obawa potencjalnego twórcy nowej idei przed wyśmianiem pomysłu, który w rezultacie nie wychodzi na światło dzienne. Burza mózgów daje swobodę wypowiedzi i jest miejscem gdzie można formułować najdziwsze pomysły. Jednym z takich pomysłów jest „winda” w kosmos. Pomysł zakładał umieszczenie na orbicie masywnej masy, związanej z Ziemią za pomocą wytrzymałej taśmy. Dzięki tej taśmie można by wysłać w kosmos różne obiekty, które jechałyby w górę dzięki rolkom zaciskającym na tej taśmie. Taśma zaś byłaby utrzymywana w stanie naprężonym dzięki sile odśrodkowej, wynikającej z obrotu Ziemi wokół osi. Pomysł rzeczywiście dziwny, mający jednak szansę na realizację. Kiedy? Kto to wie...

Kolejną metodą „rozhuśtania” fantazji twórczej jest metoda „obiektów fokalnych”. Występuje w dwóch zasadniczych odmianach: metoda ulepszania istniejącego obiektu i tworzenie czegoś nowego, na zasadzie wykorzystania wzorców z przyrody, bajek, mitów itp.

3



Ulepszanie czegoś, co już istnieje jest znacznie częstszą metodą. Wynika to z faktu, że równolegle do rozwoju samych obiektów rozwijają się komponenty, takie jak: wysokosprawne ogniwa i akumulatory, wysokosprawne diody LED itd. Efektem tego zjawiska są chociażby samochody elektryczne, modele samolotów, z których niemal zniknął napęd silniczkami spalinowymi. Gwałtowny rozwój np. aparatury medycznej, jak: sprzęt do kolonoskopii, dziś wykonany jako giętki przewód zaopatrzony w oświetlenie, manipulator i kamerę, przekazującą obraz z wnętrza ciała człowieka na ekran monitora. Jest to zgodne z prawem, polegającym na tym, że postępek w jednej dziedzinie rodzi postępek w kilku innych.

Korzystanie z projektów przyrody też daje spektakularne wyniki: np. ryba, ostronos (*boxfish*), mimo że jej kształt wydaje się kanciasty i mało aerodynamiczny, badania wykazały, że opływowość jej ciała jest bardzo wysoka – ma niewielki współczynnik oporu powietrza, co czyni ją efektywnym wzorcem dla projektantów, m.in. tak właśnie zaprojektowano samochód Mercedes-Benz Bionic Car – **rysunek 3**.

Samochód ten na oko wydaje się mało opływowy, a jednak współczynnik jego oporu czołowego jest najmniejszy spośród większości współczesnych samochodów osobowych i wynosi 0,19, w stosunku do typowego 0,22 i najmniejszego typowego 0,195.

Jest to oczywiście piękny przykład zastosowania bioniki w klasie „obiektu fokalnego”, jakim była tu ryba: kostera gruzełkowata (*Ostracion cubicus*), żyjąca w rafach koralowych. Firma Mercedes-Benz zdecydowała się na ten „dziwny” kształt i wygrała w wielu parametrach jak emisja CO₂, zużycie paliwa, przyspieszenie itd. Jak widać, zarozumiałość człowieka znów została pokonana przez przyrodę. A więc: *Disce, puer, Latinam*, czyli: wracajmy czasami do źródeł. ■

Prezes Klubu Wynalazców
Champion TRIZ
Jan Boratyński



Nieustannie czekamy na Wasze pomysły ulepszeń, innowacji, zmian. Swoje propozycje nadsyłajcie na adres redakcji. „Pomysły” nie są wołaniem na puszczy! Komentujemy, oceniamy i staramy się wyrazić nasz szczerzy podziw i uznanie dla pomysłowości Czytelników. Gorąco zachęcamy wszystkich do prezentowania swoich koncepcji, również tych najbardziej zwariowanych! Wszystkie mają wartość, nawet te z pozoru niedorzeczne, bo ich krytyka może stać się twórczym zaczątkiem czegoś ciekawego! **A oto plon ostatniego miesiąca:**

Pomysł miesiąca 06/2025

Pomysł detekcji potrzeb rośliny, nie tylko jeśli chodzi o wodę, ale być może także inne składniki ma duży potencjał, nie tylko przy roślinach doniczkowych, ale być może nawet w wielkoskalowym rolnictwie. Tylko jak odbierać, identyfikować i dekodować ewentualne sygnały. To jest wyzwanie.

Autorem pomysłu jest **Piotr Biernat**

1 Krzysztof Dudzik – uważa, że w dobie rozwiniętej techniki komputerowej i teorii sieci neuronowych istnieje możliwość stworzenia czegoś w rodzaju trenera kandydatów na wysokie funkcje państwowe. Jak wiadomo, sieć neuronową najpierw należy „wyedukować”, tzn. zebrać dane z ostatnich 25–30 lat. Dane typu: pan XY zaproponował i wprowadził przepis, dało to efekty ZY. Duża liczba tych skojarzeń pozwoliłaby przeprowadzić „edukację” sieci, a później wykorzystać w dwóch sytuacjach: selekcji kandydatów na konkretne stanowisko oraz do oceny bieżącej pracy nominowanego i zaangażowanego pracownika. Oszczędziłoby to tzw. debat publicznych, które są w sumie dość nijakie w sensie oceny przydatności kandydata na stanowisko prezydenta RP i pozwoliłoby dać obiektywną ocenę konkretnej osoby na to stanowisko.

Propozycja kolegi jest mocno idealistyczna, ale rzeczywiście taki tester kandydatów ogromnie ułatwiłby ocenę przydatności osoby na proponowane stanowisko. Problem polega na tym, że ilość materiału do „edukowania” sieci neuronowej musiałaby być ogromna, uwzględnić także cechy psychiczne, takie jak: odporność na stres, asertywność, samodzielność w ocenie sytuacji i wiele innych. Wobec żywiłowego rozwoju sztucznej inteligencji budzi to jednak nadzieję, że coś takiego się pojawi. Obecnie działanie wielu osób nawet na wysokich stanowiskach budzi podstawowe zastrzeżenia.

2 Jacek Chorabik – uważa, że najwyższy czas skończyć z pracowitym oczyszczaniem szkieł okularów i należy zastosować ultradźwiękowe oczyszczanie szkieł bez używania ściereczek. W ramkę okularów należy wbudować mikroukład generujący ultradźwięki, a to się dziś da zrobić i po kłopotach!

Okulary mają dziś coraz bogatsze wyposażenie: np. lampki LED-owe, generator ultradźwięków może się zmieścić, ale oznacza to coraz cięższe okulary.

3 Jacek Saleta – proponuje inteligentne kubki, rozpoznające, co się pije: kawę, herbatę lub sok. Kubek taki powinien kontrolować ilość wypitych

napojów i np. przy piątej z rzędu kawie ostrzegać: czy nie za dużo tej kawy?

To nie takie proste, bo w końcu to użytkownik decyduje, co i ile wypija. Poza tym kubek taki wymagałby dość skomplikowanego układu i jak to myć bez narażenia elektroniki na uszkodzenie.

4 Zbigniew Tkacz pisze: są już translatory, tłumaczące online ze 120 języków, a dotychczas nie ma tłumacza tłumaczącego mowę naszych „małych braci”, domowych psów i kotów. Mając psa w domu, po jakimś czasie zaczynamy rozumieć i rozróżniać różne jego głosy, ale translator w znacznym stopniu ułatwiłby kontakt z naszym pupilem.

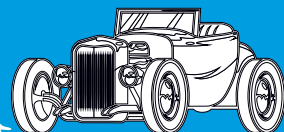
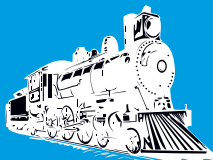
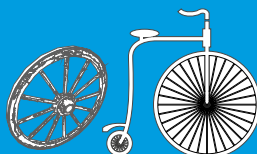
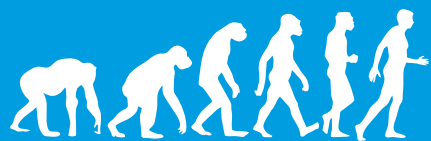
Tego jeszcze nie było, chociaż zdarzają się ludzie, którzy odczytują myśli zwierząt. Ciekawe, jak taki translator mógłby działać: czy w torze akustycznym, czy w podświadomości.

5 Marek Skwarowski – proponuje inteligentny śmietnik, który mówi: „dziękuję”, gdy wrzucisz plastik w dobre miejsce. Jeśli się pomyliś, odzywa się z żartobliwym sarkazmem: „Czy ta folia to jakaś blacha? Przemyśl to jeszcze raz”...

Dobry pomysł. Proces przyzwyczajania społeczeństwa do sortowania śmieci trwa już dość długo i wciąż nie jest w 100% skuteczny. Winne jest nasze wygodnictwo, podszyte lenistwem. Taka „przypominajka” mogłaby znacznie poprawić proces gospodarowania odpadami.

6 Piotr Bernat – inteligentna doniczka z pomiarem bioelektrycznych sygnałów roślin i automatycznym systemem nawadniania i nawożenia. Wysła powiadomienie na telefon, gdy kończy się zapas wody w pojemniku.

Pomysł nienowoty, ale wciąż nie może doczekać się upowszechnienia. Już Cleve Backster – amerykański ekspert od przesłuchiwanie ludzi z pomocą wykrywacza kłamstw, proponował wykorzystanie aktywności bioelektrycznej roślin do sterowania nawożeniem, podlewaniem itd. na dużych plantacjach.



Chmury obliczeniowe

1961

John McCarthy (1), sławny programista, twórca m.in. pojęcia „sztuczna inteligencja”, wygłasza na MIT wykład, w którym mówi m.in., że „informatyka może być sprzedawana jako narzędzie użyteczności publicznej, podobnie jak woda i elektryczność”. Sformułowanie to uznawane jest często za załączek myślenia w kategoriach chmur obliczeniowych i usług za ich pomocą świadczonych. W latach 60. XX wieku pierwotne koncepcje prowadzące do konsekwencji do późniejszych chmur obliczeniowych określano jako systemy „podziału czasu” albo „remote job entry”. Terminologii takiej używali najwięksi wówczas dostawcy komputerów i oprogramowania, IBM i DEC. Rozwiązania z pełnym podziałem czasu były dostępne komercyjnie dopiero w latach 70. na takich platformach jak Multics, Cambridge CTSS i przez najwcześniejsze rozwiązania UNIX (na sprzęcie DEC). Przeważał jednak w tamtych czasach model „centrum danych”, w którym użytkownicy przekazywali operatorom zadania do uruchomienia na komputerach mainframe IBM.

1963–67

Amerykańska DARPA, czyli Agencja Zaawansowanych Projektów Badawczych w Obszarze Obronności (2), przekazuje MIT dwa miliony dolarów na rozwój projektu MAC. Agencja oczekiwała od uczelni opracowania rozwiązań umożliwiających „jednoczesne korzystanie z komputera przez dwie lub więcej osób”. W ramach projektu MAC jeden z wielkich komputerów mainframe, wykorzystujący szpule taśmy magnetycznej jako układ pamięci, skonfigurowany został jako prekursor systemu, który obecnie znamy jako chmura obliczeniowa. Dostęp do tak powstałej „chmury” miały dwie lub trzy osoby. W 1967 roku firma IBM opracowuje koncepcję wirtualizacji systemów operacyjnych, która umożliwia wielu użytkownikom współdzielenie tego samego zasobu w sieci komputerów. Koncepcja wirtualizacji ewoluowała wraz z Internetem, gdy firmy zaczęły oferować „wirtualne” sieci prywatne jako usługę do wynajęcia.

1969

Powstaje ARPANET (Advanced Research Projects Agency Network), inicjatywa Departamentu Obrony USA, pierwsza sieć rozległa oparta na rozproszonej architekturze i protokole TCP/IP. Jest bezpośrednim przodkiem internetu.

lata 80. XX wieku

Na początku tej dekady zaczęto wprowadzać sieciowe systemy operacyjne, które umożliwiały komputerom komunikowanie się ze sobą. W 1985 r. do internetu podłączonych było już około stu tysięcy komputerów.

1994

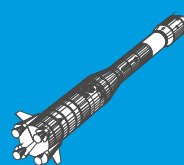
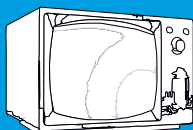
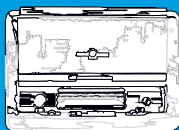
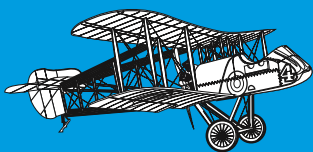
Metafora „chmury” dla zwirtualizowanych usług pochodzi z 1994 roku, kiedy została użyta przez firmę software’ową General Magic dla świata „miejsc”, do których mogli „udać się” mobilni agenci w środowisku języka kodowania Telescript (3).

1996

Wyrażenie „cloud computing” stało się szerzej znane dzięki biznesplanowi firmy Compaq Computer Corporation (4) dotyczącemu przyszłych modeli obliczeniowych i Internetu. Ambicją firmy było zwiększenie sprzedaży dzięki, jak to ujęto w dokumencie, „aplikacjom obsługującym przetwarzanie w chmurze”. Plan przewidywał, że przechowywanie plików konsumentów online prawdopodobnie odniesie komercyjny sukces.

1997

Profesor Ramnath Chellapa z Uniwersytetu Emory definiuje chmurę obliczeniową jako nowy „paradygmat obliczeniowy, w którym granice obliczeń będą określone przez przesłanki ekonomiczne, a nie tylko ograniczenia techniczne”. Ten nieco zawile brzmiący opis jednak dość dokładnie oddaje istotę chmury obliczeniowej.



1999

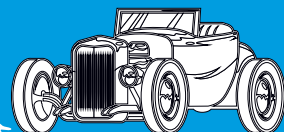
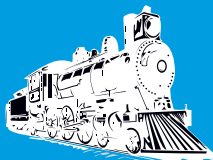
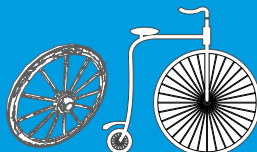
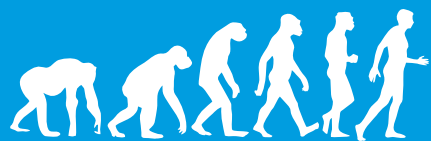
Dużym przełomem ostatnich lat 90. było zrozumienie przez firmy potencjału chmur obliczeniowych. Firma Salesforce stała się popularnym przykładem udanego wykorzystania tej techniki (5). Chmura posłużyła jej do wdrażania pionierskiego pomysłu dostarczania oprogramowania klientom firmy – użytkownikom końcowym za pomocą sieci. Program (lub aplikacja) był dostępny i pobrany przez każdego, kto miał dostęp do Internetu. Firmy mogły kupować potrzebne do działalności oprogramowanie na żądanie, w wygodny i opłacalny sposób, bez opuszczania siedziby i kupowania dodatkowych nośników.

2002–2006

Amazon jako jedna z pierwszych znaczących firm, uznał wykorzystanie zaledwie 10 proc. własnej mocy obliczeniowej (a było to wówczas powszechne) za problem, który warto rozwiązać. Opracował więc nowy model infrastruktury z przetwarzaniem w chmurze do efektywnego wykorzystania pojemności i mocy obliczeniowej komputerów. Amazon stworzył i udostępnił usługę Amazon Web Services (AWS) do przechowywania danych, w tym hostingu zasobów stron WWW (6), przetwarzania przez Internet. Cztery lata później Amazon wdraża komercyjną usługę chmurową Elastic Compute Cloud (EC2), która jest otwarta dla klientów firmy, umożliwiając im wynajmowanie wirtualnych komputerów i korzystanie z własnych programów i aplikacji. Wkrótce potem inne duże organizacje zaczęły iść tą ścieżką. W 2006 r. Google uruchomiło Google Docs, model SaaS (oprogramowania jako usługi) do edycji i zapisywania dokumentów online. Google Docs było pierwotnie oparte na dwóch oddzielnych produktach, Writely i Google Spreadsheets. Pierwszy pozwala tworzyć i zapisywać dokumenty (kompatybilne z Microsoft Word), edytować je i ewentualnie publikować. Google Spreadsheets (przejęte od 2Web Technologies w 2005 roku) to program internetowy umożliwiający użytkownikom tworzenie, aktualizację i edycję arkuszy kalkulacyjnych, kompatybilnych z Microsoft Excel, oraz udostępnianie danych online.



1. John McCarthy; 2. Logo agencji DARPA; 3. Jeden z obrazów interfejsu platformy Telescript firmy General Magic; 4. George Favaloro, szef marketingu Compaqa, pozuje z biznesplanem firmy z 1996 r.; 5. Wczesna siedziba firmy Salesforce; 6. Jedna z siedzib AWS



2004–14

2007

2008–2010

2011

2011–12

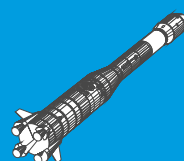
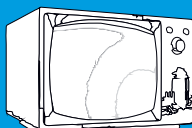
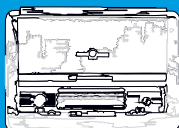
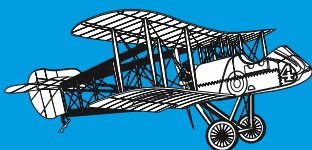
Pierwsze proste kontenery w chmurach (Solaris). Określa się tak przenośne, lekkie jednostki oprogramowania, które pakują aplikacje wraz z ich zależnościami, bibliotekami, pliki binarne i konfiguracyjne, co pozwala na spójne działanie w dowolnym środowisku, w tym na różnych platformach chmurowych. Te wczesne kontenery były ograniczone do niektórych tylko systemów komputerowych. Dopiero w 2013 roku, gdy firma Docker opracowała nowy rodzaj funkcjonalnego kontenera, narzędzia te zyskały popularność. Kubernetes, czyli platforma, która zarządza aplikacjami w kontenerach, tworząc system tzw. orkiestracji kontenerów, którego celem jest automatyzacja wdrażania, skalowania i zarządzania aplikacjami, opracowana została przez Google w 2014 roku, a następnie udostępniona jako produkt open source (7).

Netflix startuje ze swoją usługą strumieniowego przesyłania wideo, wykorzystującą jako podstawę funkcjonowania technikę chmury.

Google udostępnia wersję beta Google App Engine w modelu Platforma jako Usługa (PaaS), zapewniającą w pełni zarządzaną infrastrukturę i platformę dla użytkowników do tworzenia aplikacji internetowych (8). Po raz pierwszy deweloper oprogramowania mógł dzięki temu rozwiązać problem opracowania program i uruchomić go w chmurze Google bez martwienia się o szczegóły, np. sposób udostępniania serwera lub to, czy system operacyjny wymaga poprawek. W 2009 roku Microsoft uruchomił Microsoft Azure, a następnie inne firmy, takie jak Alibaba, IBM, Oracle, HP również wprowadziły swoje usługi w chmurze. Amazon uruchomił podobną usługę o nazwie Lambda w 2015 roku. Usługi tego typu były zarazem początkami tzw. rozwiązań bezserwerowych. W 2010 r. pojawia się OpenStack, oprogramowanie z dziedziny chmur obliczeniowych w modelu Infrastructure as a Service (IaaS) rozwijane przez Rackspace Cloud oraz NASA. W ramach tej usługi udostępniono otwartą, darmową chmurę typu „zrób to sam”, która stała się bardzo popularna.

Amerykański Narodowy Instytut Standardów i Technologii (NIST) ustala „pięć podstawowych cech systemów chmurowych”. W skrócie były to: samoobsługa na żądanie, szeroki dostęp do sieci („możliwości są dostępne przez sieć – za pośrednictwem standardowych mechanizmów, które promują korzystanie z [...] platform klienckich, np. telefonów komórkowych, tabletów, laptopów i stacji roboczych”), łączenie zasobów („zasoby obliczeniowe dostawcy są łączone w celu obsługi wielu konsumentów przy użyciu modelu wielu dzierżawców, z różnymi zasobami fizycznymi i wirtualnymi”), szybkość połączona z elastycznością i mierzalność usług („systemy chmurowe automatycznie kontrolują i optymalizują wykorzystanie zasobów poprzez wykorzystanie możliwości pomiaru na pewnym poziomie abstrakcji odpowiednim dla rodzaju usługi, np. pamięć masowa, przetwarzanie, przepustowość i aktywne konta użytkowników”).

Wprowadzenie koncepcji chmur hybrydowych, czyli interoperacyjności między chmurą prywatną i publiczną (9) oraz możliwość przenoszenia obciążeń między nimi. Jedną z pierwszych usług tego rodzaju oferowała firma CloudBolt.



2012

Firma Oracle wprowadza chmurę Oracle Cloud, oferując trzy podstawowe usługi dla biznesu: IaaS (infrastruktura jako usługa), PaaS (platforma jako usługa) i SaaS (oprogramowanie jako usługa). Oferowany przez Oracle zakres usług (10) szybko stał się normą. Niektóre chmury publiczne oferowały je, jednak niektóre koncentrowały się na oferowaniu tylko jednej z nich.

2016

IBM podłącza do chmury komputer kwantowy, który umożliwia tworzenie i wykonywanie prostych programów w chmurze. Na początku 2017 roku naukowcy z Rigetti Computing zademonstrowali pierwszy programowalny dostęp do chmury za pomocą biblioteki programowania kwantowego pyQuil Python, dzięki której specjaliści mogli tworzyć programy, które uruchamiają różne algorytmy kwantowe.

2019

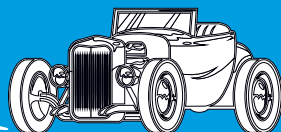
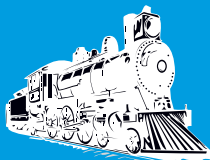
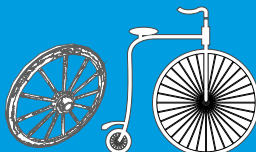
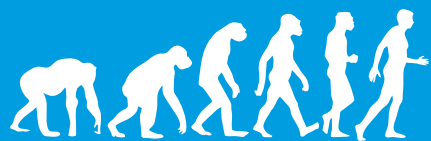
Amazon udostępnia AWS Outposts, usługę „szafy serwerowej”, która rozszerza infrastrukturę, usługi, interfejsy API i narzędzia AWS na centra danych klientów, przestrzenie kolokacyjne lub obiekty lokalne.

lata 20. XXI wieku

Zarządzanie chmurą oparte na sztucznej inteligencji rewolucjonizuje sposób optymalizacji i utrzymania zasobów w chmurze. Dzięki analityce predykcyjnej i automatyzacji sztuczna inteligencja dynamicznie przydziela zasoby, przewiduje awarie i wydajniej zarządza obciążeniami. Trwająca obecnie rewolucja AI w technice chmur obliczeniowych zmierza do obniżenia kosztów a także do podniesienia wydajności i niezawodności.



7. Docker i Kubernetes; 8. Google App Engine – logotyp; 9. Wizualizacja chmury hybrydowej; 10. Trzy filary techniki chmurowej firmy Oracle



Rodzaje usług w chmurach obliczeniowych

Oprogramowanie jako usługa (Software as a Service, SaaS)

W modelu oprogramowania jako usługi (SaaS) użytkownicy uzyskują dostęp do oprogramowania, aplikacji i baz danych. Dostawcy usług w chmurze zarządzają infrastrukturą i platformami, na których działają aplikacje. SaaS jest czasami określany jako „oprogramowanie na żądanie” i jest zwykle wyceniane na zasadzie płatności za użycie lub przy użyciu opłaty abonamentowej. W modelu SaaS dostawcy usług w chmurze instalują i obsługują oprogramowanie aplikacyjne w chmurze, a użytkownicy chmury uzyskują dostęp do oprogramowania z klientów chmury. Użytkownicy chmury nie zarządzają infrastrukturą chmury i platformą, na której działa aplikacja. Eliminuje to potrzebę instalowania i uruchamiania aplikacji na własnych komputerach użytkownika chmury, co upraszcza konserwację i wsparcie. Aplikacje w chmurze różnią się od innych aplikacji skalowalnością, którą można osiągnąć poprzez klonowanie zadań na wiele maszyn wirtualnych w czasie wykonywania, aby sprostać zmieniającemu się zapotrzebowaniu na pracę. Proces ten jest przezroczysty dla użytkownika chmury, który widzi tylko jeden punkt dostępu. Aby pomieścić dużą liczbę użytkowników chmury, aplikacje chmurowe mogą być wielodostępne, co oznacza, że każda maszyna może obsługiwać więcej niż jedną organizację użytkowników chmury.

Platforma jako usługa (Platform as a Service, PaaS)

Usługa ta polega na wdrażaniu w infrastrukturze chmury aplikacji stworzonych lub nabytych przez konsumenta przy użyciu języków programowania, bibliotek, usług i narzędzi obsługiwanych przez dostawcę. Klient nie zarządza ani nie kontroluje podstawowej infrastruktury chmury, w tym sieci, serwerów, systemów operacyjnych lub pamięci masowej, ale ma kontrolę nad wdrożonymi aplikacjami i ewentualnie ustawieniami konfiguracji środowiska hostingu aplikacji. Dostawcy PaaS oferują środowisko programistyczne dla twórców aplikacji. Dostawca zazwyczaj opracowuje zestaw narzędzi i standardy rozwoju oraz kanały dystrybucji i płatności. W modelach PaaS dostawcy usług w chmurze dostarczają platformę



obliczeniową, obejmującą zazwyczaj system operacyjny, środowisko wykonawcze języka programowania, bazę danych i serwer WWW. Deweloperzy aplikacji rozwijają i uruchamiają swoje oprogramowanie na platformie chmurowej, zamiast bezpośrednio kupować i zarządzać bazowym sprzętem i warstwami oprogramowania. W przypadku niektórych usług PaaS podstawowe zasoby komputerowe i pamięci masowej skalują się automatycznie, aby dopasować się do zapotrzebowania aplikacji, dzięki czemu użytkownik chmury nie musi ręcznie przydzielać zasobów.

Infrastruktura jako usługa (Infrastructure as a Service, IaaS)

Pojęcie IaaS odnosi się do usług online, które zapewniają interfejsy API wysokiego poziomu używane do abstrakcji różnych niskopoziomowych szczegółów podstawowej infrastruktury sieciowej, takich jak fizyczne zasoby obliczeniowe, lokalizacja, partycjonowanie danych, skalowanie, bezpieczeństwo, tworzenie kopii zapasowych itp. Chmury IaaS często oferują dodatkowe zasoby, takie jak biblioteka obrazów dysków maszyn wirtualnych, nieprzetworzona pamięć blokowa, pamięć plików lub obiektów, zapory ogniowe, równoważenie obciążenia, adresy IP, wirtualne sieci lokalne (VLAN) i pakiety oprogramowania.

Funkcja jako usługa (Function as a Service, FaaS)

FaaS jest sterowana zdarzeniami, pomagając w utrzymaniu serwerów i infrastruktury na bieżąco. FaaS ułatwia programistom uruchamianie kodu w odpowiedzi na zdarzenia i pozwala obniżyć koszty dzięki zasadzie „Pay as you Run” (płać za wykorzystane zasoby obliczeniowe). ■

M.U.



Współczesne wzmacniacze stereofoniczne, część 3

W poprzednich numerach MT, w pierwszej i drugiej części artykułu poświęconego wzmacniaczom stereofonicznym, przedstawiliśmy ich funkcjonalność, związaną głównie z zadaniami i konstrukcją sekcji przedwzmacniacza. W tej części przechodzimy do techniki końcówek mocy.

Zwłaszcza w high-endzie spotkać można wiele konstrukcji lampowych. Wiąże się to zarówno z wyższymi kosztami, jak i przekonaniem dużej grupy audiofilów o brzmieniowej wyższości lampy nad tranzystorem. Co brzmi lepiej, może być sprawą subiektywną, natomiast obiektywnie, parametrycznie, wzmacniacze tranzystorowe mają przewagę. Nawet jeżeli wpływ szumów, zniekształceń i pasma przenoszenia pozostawimy do oceny brzmieniowej, to wzmacniacze tranzystorowe mogą osiągać wysokie moce, lepiej radzą sobie z kolumnami o „trudnych” impedancjach i zapewniają lepszą odpowiedź impulsową, a to są już twarde fakty. Jednak znaczenie każdego z tych parametrów też można zrelatywizować. Nie każdemu

potrzebna jest wysoka moc, a jej deficyt można zrekompensować kolumnami o wysokiej efektywności, można też znaleźć kolumny o wysokiej i „łatwej” impedancji, a odpowiedź impulsowa... i tak zależy głównie właśnie od zespołów głośnikowych. Stąd wzmacniacze lampowe wciąż wychodzą obronną ręką, a nawet zwycięsko w ocenie dużej grupy audiofilów. Argumenty obydwu stron w tej kwestii nie zmieniły się od kilkadziesiąt lat i wszystko wskazuje na to, że lampy pozostaną z nami jeszcze długo, bowiem mają swój niepowtarzalny urok – podobnie jak płyta winylowa, której w związku z tym przypisuje się magiczne brzmienie, aby umocnić swoje (i innych) przekonanie o wyższości gramofonu nad źródłami cyfrowymi. To jednak jeszcze bardziej skomplikowana materia, która wcale nie miała być tematem tego artykułu, więc proszę nie wyciągać wniosku, że lekceważymy „analog” i uważamy źródła cyfrowe



Amerykańska firma Boulder buduje potężne wzmacniacze, głównie końcówki mocy. Model 2160 wcale nie jest najmocniejszy, ale i jego parametry są spektakularne. Moc ciągła przy 4 omach przekracza 2×1000 W (przy obydwu kanałach wystero- wanych jednocześnie, według pomiarów AUDIO, nie specyfikacji producenta), a przy 2 omach... ponad 2000 W, przy czym zmierzylśmy tylko jeden kanał, bo zabrakło obciążen pomiarowych dla dwóch... Impedancja wyjściowa – o wartości poniżej czułości systemu pomiarowego, czyli współczynnik tłumienia rzędu tysięcy. Znikome szумы i zniekształcenia. Demonstracja możliwości końcówki tranzystorowej w klasie AB. Test AUDIO 3/2024



Kondo Overture PM21
Japońska manufaktura Kondo to już legenda, dyktująca za swoje produkty „astronomiczne” ceny, ale mająca wielbicieli, a wśród nich klientów. Jej specjalizacją są wzmacniacze lampowe, których główną bronią jest brzmienie, a nie moc wyjściowa czy wyposażenie. PM-2i oddaje 2×13 W przy 8 omach i 2×21 W przy 4 omach, ale jak na lampę ma bardzo niski szum i zniekształcenia. Wyposażenie ogranicza się do czterech wejść liniowych. Test AUDIO 1/2024



Duński Gryphon Diablo to już tradycyjnie jeden z najmocniejszych wzmacniaczy zintegrowanych. Swoją moc wciąż zwiększa w kolejnych wersjach, przy tym wiernie trwając w klasie AB. Aktualny model Diablo 333 dostarcza 2×675 W przy 4 omach. Wyposażenie, jak przystało na high-end – w opcji bazowej minimalne, tylko z wejściami liniowymi, ale dwa moduły rozszerzeń pozwalają dodać DAC i pre-amp phono. Test AUDIO 7–8/2024

za idealne. Zwłaszcza materiały w popularnych serwisach internetowych, ze względu na kompresję dynamiki dokonywaną pod kątem potrzeb potencjalnych słuchaczy, a nie z powodu ograniczeń techniki, mogą okazać się słabsze niż porządnie przygotowane i wytłoczone wydanie na winylu. Ale żeby je docenić... potrzeba solidnego gramofonu, z dobrą wkładką, wyregulowanego, podłączonego do dobrego przedwzmacniacza. Wracamy z tej wycieczki do wzmacniaczy.

Szczególnym, ale już od dawna dobrze znanym układem jest tzw. wzmacniacz hybrydowy. Polega on na połączeniu lampowego przedwzmacniacza z tranzystorową końcówką mocy. Idea za tym stojąca jest łatwa do odczytania. Jednoznacznie ujemne strony lamp, przedstawione wcześniej, objawiają się przy ich stosowaniu w końcówce mocy. Zrealizowanie jej na tranzystorach może zapewnić wysoką moc i wysoki współczynnik tłumienia również wówczas, gdy w przedwzmacniaczu „siedzą” lampy. Te z kolei mają tam wpłynąć na charakter brzmienia, nadając mu lampowe zabarwienie, oczywiście poprzez duży, ale subiektywnie korzystny udział zniekształceń harmonicznych. Z tego samego powodu lampy znajdują zastosowanie (czasami) np. w analogowych stopniach wyjściowych odtwarzaczy CD, a nawet strumieniujących.

Wśród wzmacniaczy tranzystorowych od kilkunastu lat swoją obecność zwiększają końcówki mocy pracujące w klasie D. Wbrew pewnym pozorom, nie są to wzmacniacze cyfrowe (choć we wzmacniaczach zintegrowanych mogą współpracować



Wzmacniacze „hybrydowe” mają lampy i tranzystory, i to te ostatnie są odpowiedzialne za moc wyjściową. Dlatego nic dziwnego, że włoski Pathos Logos MkII kusi lampkami i dostarcza 2×240 W przy 4 omach. Test AUDIO 11/2024

z cyfrowymi układami w sekcjach przedwzmacniacza). To szybko rozwijająca się technika, której zastosowanie początkowo ograniczało się do niskich częstotliwości (a więc subwooferów aktywnych), ponieważ wyższe częstotliwości obciążone były problemami (w których dokładne powody nie będziemy tutaj wnikać); jednak obecnie z większością nich nieźle sobie poradzono. Od początku podstawową zaletą wzmacniaczy w klasie D jest ich wysoka sprawność, pozwalająca osiągać wysokie moce wyjściowe przy niewielkiej mocy pobieranej z sieci, co przekłada się też na umiarkowaną wielkość tego typu urządzeń, bowiem przy niewielkiej porcji mocy zamienianej na ciepło, niepotrzebne są duże radiatory i obudowy, co dla wielu użytkowników, niezainteresowanych budowaniem okazałych systemów jest ważnym argumentem



Chociaż high-end raczej stroni od klasy D, to zdarzają się wyjątki. Jednym z najoryginalniejszych jest polski Mytek Empire, końcówka mocy na tranzystorach GanFET, doskonałych w tym zastosowaniu ze względu na bardzo krótkie czasy przełączania, umożliwiające pracę przy bardzo wysokich częstotliwościach. Lampa widoczna na górnej pokrywie... to tylko dekoracja; chociaż świeci, nie znajduje się w torze audio. Test AUDIO 12/2024



Nuprime Omnia A200 – wzmacniacz w klasie D w najbardziej typowej, kompaktowej formie. Moc 2×200 W przy 4 omach. Ultranowoczesny, ze strumieniowaniem, sterowaniem aplikacją mobilną i korekcją charakterystyki częstotliwościowej, opartą na DSP. Test AUDIO 12/2024



Egzotyczna klasyka. Musical Fidelity A1, którego pierwsza wersja pojawiła się czterdzieści lat temu, stał się symbolem niewielkiego, audiofilskiego wzmacniacza w klasie A. Po długiej nieobecności odtworzono jego konstrukcję z dużą pieczołowitością (wymieniono tylko te komponenty, które nie są już dostępne) i delikatnie unowocześniono (m.in. dodano zdalne sterowanie). Karbowana górna płyta pełni funkcję radiatora, najwyższa moc 2×25 W dostępna jest na obciążeniu 8-omowym (na 4-omowym spada do 2×17 W). Test AUDIO 9/2023



Canor Audio AI 1.20 to bardzo rasowy reprezentant klasy A. Z potężnej konstrukcji w konfiguracji dual mono, z parą dużych transformatorów i obszer-nych radiatorów, wyciągniemy „tylko” 2×60 W przy 4 omach – to i tak bardzo dużo w czystej klasie A. Inne parametry są bez zarzutu, z doskonałym odstępem od szumu na czele (100 dB), który przy nawet umiarkowanej mocy wywindował dynamikę do 113 dB. Test AUDIO 1/2025

estetycznym. Klasa D znajduje też szerokie zastosowanie tam, gdzie miejsca jest z obiektywnych powodów mało, a chłodzenie jest utrudnione – w aktywnych zespołach głośnikowych, subwooferach, soundbarach, głośnikach BT.

Mimo to, we wzmacniaczach zintegrowanych, końcówki w klasie D wciąż nie zdobyły przewagi nad najczęściej spotykanymi od wielu lat wzmacniaczami w klasie AB; zwłaszcza w high-endzie, gdzie najsilniejsze jest przywiązanie do starych, sprawdzonych rozwiązań. Tutaj klienci najmniej liczą się z kosztami a wielkość urządzeń też nie jest problemem, więc można sobie pozwalać na takie luksusy, jak wzmacniacze w klasie AB o mocach kilkuset watów i masie kilkudziesięciu kilogramów.

Swoją drogą postęp w konstruowaniu wzmacniaczy w klasie AB pozwala też całkiem niewielkim urządzeniom osiągać moce wystarczające dla większości zainteresowanych dobrym sprzętem, i odsunąć ryzyko spotkania z wciąż budzącą obawy klasą D.

Warto wspomnieć też o klasie A, chociaż takich wzmacniaczy jest najmniej. To klasa o najniższej sprawności, ale „najszlachetniejsza”, potencjalnie zapewniająca najniższe zniekształcenia, dzięki zasadzie swojej pracy. Wysoki prąd spoczynkowy powoduje jednak, że większość pobieranej energii zostaje zamieniona na ciepło (zwłaszcza przy niskich poziomach) i ze wzmacniaczy o „normalnej” wielkości, pracujących w klasie A, można oczekiwać kilkunastu watów; czasami spotykane obietnice mocy zbliżających się do 100 W, a nawet ją przekraczających, są bez pokrycia albo wiążą się z tym, że po „wyczerpaniu” mocy możliwej w klasie A układ przechodzi płynnie do pracy w klasie AB. Ponadto fakt pracy w klasie A nie jest gwarancją niskich zniekształceń i dobrego brzmienia; nie jest nią żadne teoretyczne rozwiązanie ani najlepsze komponenty, za końcowy rezultat odpowiada wiele czynników, z których każdy może „zawalić sprawę”. Nie należy więc zbyt srogo sugerować się „patentami”, lecz wynikami niezależnych pomiarów (każdy test wzmacniacza w AUDIO obejmuje kompleksowe pomiary) i odsłuchów. ■

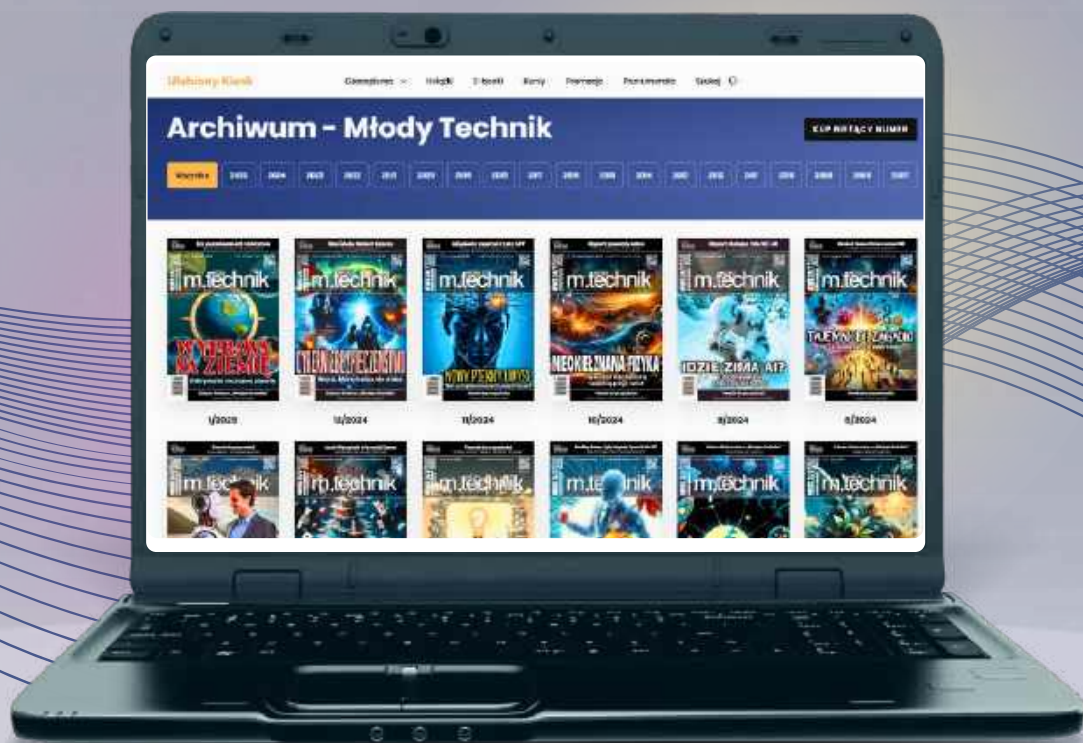
Andrzej Kisiel

AUDIO

**Miesięcznik audiofilski – polski przedstawiciel
European Imaging and Sound Association**

przeglądaj, czytaj i kup na
www.ulubionykiosk.pl

SIĘGNIJ PO WYDANIA ARCHIWALNE **MŁODEGO TECHNIKA**



Przejrzyj wszystkie wydania
on-line i zamów wygodnie na
www.UlubionyKiosk.pl



Przesyłka **GRATIS**